



Харківський національний університет радіоелектроніки
Державне підприємство
"Південний державний
проектно-конструкторський та науково-дослідний
інститут авіаційної промисловості"



Kharkiv National University of Radio Electronics
State Enterprise
"Southern National Design & Research Institute of Aerospace Industries"

Сучасний стан наукових досліджень та технологій в промисловості

Innovative technologies and scientific solutions for industries

Щоквартальний науковий журнал
№ 2 (36), 2026

Затверджений до друку
Науково-технічною Радою
Харківського національного університету радіоелектроніки
(Протокол № 6 від 21 травня 2026 р.)

ЗАСНОВНИКИ

Харківський національний університет радіоелектроніки,
Державне підприємство "Південний державний
проектно-конструкторський та науково-дослідний інститут
авіаційної промисловості"

АДРЕСА РЕДАКЦІЇ:

Україна, 61166, м. Харків, проспект Науки, 14
Інформаційний сайт: <http://itssi-journal.com>
<https://journals.uran.ua/itssi>
E-mail редколегії: journal.itssi@gmail.com

Quarterly scientific journal
No. 2 (36), 2026

Approved for publication
by the Scientific and Technical Council
of the Kharkiv National University of Radio Electronics
(Protocol No. 6 d.d. 21 May 2026)

ESTABLISHERS

Kharkiv National University of Radio Electronics,
State Enterprise
"National Design & Research Institute
of Aerospace Industries"

EDITORIAL OFFICE ADDRESS:

14 Nauky Ave., Kharkiv, Ukraine, 61166
Information site: <http://itssi-journal.com>
<https://journals.uran.ua/itssi>
E-mail of the editorial board: journal.itssi@gmail.com

Ідентифікатор медіа R30-03878
Витяг з реєстру суб'єктів у сфері медіа – реєстрантів від 25.04.2024 № 1410

*Журнал включено до Переліку наукових фахових видань України (категорія А), в яких можуть публікуватися
результати дисертаційних робіт на здобуття наукових ступенів доктора і кандидата наук
наказом Міністерства освіти і науки України від 23.12.2025 № 1683*

Харків – 2026

© Харківський національний університет радіоелектроніки,
Державне підприємство "Південний державний проектно-конструкторський
та науково-дослідний інститут авіаційної промисловості"

РЕДАКЦІЙНА КОЛЕГІЯ

Головний редактор

Бодянський Євгеній Володимирович,
д-р техн. наук, професор

Заступник головного редактора

Айзенберг Ігор Наумович,
канд. техн. наук, професор (США);
Шекер Серхат, д-р техн. наук, професор (Туреччина)

Члени редколегії:

Артюх Роман Володимирович, канд. техн. наук;
Бабенко Віталіна Олексіївна,
д-р екон. наук, канд. техн. наук, професор;
Безкоровайний Володимир Валентинович,
д-р техн. наук, професор;
Гасімов Юсіф, д-р мат. наук, професор (Азербайджан);
Гопєєнко Віктор, д-р техн. наук, професор (Латвія);
Го Цян, д-р техн. наук, професор (КНР);
Зайцева Єлена, д-р техн. наук, професор (Словаччина);
Зачко Олег Богданович, д-р техн. наук, професор;
Коваленко Андрій Анатолійович, д-р техн. наук, професор;
Косенко Віктор Васильович, д-р техн. наук, професор;
Костін Юрій Дмитрович, д-р екон. наук, професор;
Левашенко Віталій, д-р техн. наук, професор (Словаччина);
Лемешко Олександр Віталійович, д-р техн. наук, професор;
Малєєва Ольга Володимирівна, д-р техн. наук, професор;
Момот Тетяна, д-р екон. наук, професор (США);
Музика Катерина Миколаївна, д-р техн. наук, професор;
Назарова Галина Валентинівна, д-р екон. наук, професор;
Невлюдов Ігор Шакирович, д-р техн. наук, професор;
Опанасюк Анатолій Сергійович,
д-р фіз.-мат. наук, професор;
Павлов Сергій Володимирович, д-р техн. наук, професор;
Паржнн Юрій, д-р техн. наук, професор (США);
Перова Ірина Геннадіївна, д-р техн. наук, професор;
Петленков Едуард, канд. техн. наук (Естонія);
Петришин Любомир, д-р техн. наук, професор (Польща);
Пітерська Варвара Михайлівна, д-р техн. наук, професор;
Рубан Ігор Вікторович, д-р техн. наук, професор;
Семенець Валерій Васильович, д-р техн. наук, професор;
Семенов Сергій, д-р техн. наук, професор (Польща);
Сетлак Галина, д-р. техн. наук, професор (Польща);
Терзіян Ваган, д-р техн. наук, професор (Фінляндія);
Тєлєтов Олександр Сергійович, д-р екон. наук, професор;
Тімофєєв Володимир Олександрович,
д-р техн. наук, професор;
Філатов Валентин Олександрович, д-р техн. наук, професор;
Чумаченко Ігор Володимирович, д-р техн. наук, професор;
Чухрай Наталія Іванівна, д-р екон. наук, професор;
Юн Джин,
канд. фіз.-мат. наук, професор (КНР);
Ястремська Олена Миколаївна, д-р екон. наук, професор.

EDITORIAL BOARD

Editor in Chief

Bodyanskiy Yevgeniy,
Dr. Sc. (Engineering), Professor, Ukraine

Deputy Chief Editor

Igor Aizenberg,
PhD (Computer Science), Professor (United States)
Serhat Seker, Dr. Sc. (Engineering), Professor (Turkey)

Editorial Board Members:

Artiukh Roman, PhD (Engineering Sciences) (Ukraine);
Babenko Vitalina,
Dr. Sc. (Economics); PhD (Engineering Sciences), Professor (Ukraine);
Bezkorovainyi Volodymyr,
Dr. Sc. (Engineering), Professor (Ukraine);
Gasimov Yusif, Dr. Sc. (Mathematical), Professor (Azerbaijan);
Gopeyenko Vectors, Dr. Sc. (Engineering), Professor (Latvia);
Guo Qiang, Dr. Sc. (Engineering), Professor (P.R. of China);
Zaitseva Elena, Dr. Sc. (Engineering), Professor (Slovak Republic);
Zachko Oleh, Dr. Sc. (Engineering), Professor (Ukraine);
Kovalenko Andrey, Dr. Sc. (Engineering), Professor, (Ukraine);
Kosenko Viktor, Dr. Sc. (Engineering), Professor, (Ukraine);
Kostin Yuri, Dr. Sc. (Economics), Professor (Ukraine);
Levashenko Vitaly, Dr. Sc. (Engineering), Professor (Slovakia);
Lemeshko Oleksandr, Dr. Sc. (Engineering), Professor (Ukraine);
Malyeyeva Olga, Dr. Sc. (Engineering), Professor (Ukraine);
Momot Tatiana, Dr. Sc. (Economics), Professor (USA);
Muzyka Kateryna, Dr. Sc. (Engineering), Professor (Ukraine);
Nazarova Galina, Dr. Sc. (Economics), Professor (Ukraine);
Nevliudov Igor, Dr. Sc. (Engineering), Professor (Ukraine);
Opanasyuk Anatoliy,
Dr. Sc. (Physical and Mathematical), Professor (Ukraine);
Pavlov Sergii, Dr. Sc. (Engineering), Professor (Ukraine);
Parzhyn Yuriy, Dr. Sc. (Engineering), Professor (USA);
Perova Iryna, Dr. Sc. (Engineering), Professor (Ukraine);
Petlenkov Eduard, PhD (Engineering Sciences) (Estonia);
Petryshyn Lubomyr, Dr. Sc. (Engineering), Professor (Poland);
Piterska Varvara, Dr. Sc. (Engineering), Professor
Ruban Igor, Dr. Sc. (Engineering), Professor, (Ukraine);
Semenets Valery, Dr. Sc. (Engineering), Professor, (Ukraine);
Semenov Serhii, Dr. Sc. (Engineering), Professor (Poland);
Setlak Galina, Dr. Sc. (Engineering), Professor (Poland);
Terziyan Vagan, Dr. Sc. (Engineering), Professor, (Finland);
Teletov Aleksandr, Dr. Sc. (Economics), Professor (Ukraine);
Timofeyev Volodymyr,
Dr. Sc. (Engineering), Professor (Ukraine);
Filatov Valentin, Dr. Sc. (Engineering), Professor (Ukraine);
Chumachenko Igor, Dr. Sc. (Engineering), Professor (Ukraine);
Chukhray Nataliya, Dr. Sc. (Economics), Professor (Ukraine);
Yu Zheng,
PhD (Physico-Mathematical Sciences), Professor (P.R. of China);
Iastremska Olena, Dr. Sc. (Economics), Professor (Ukraine)

ЗМІСТ

Інформаційні технології

- 5 **Барковська О. Ю.**
Метод контекстно-обізнаної адаптації візуального користувацького інтерфейсу на основі оцінки відстані до користувача
- 15 **Гусєва Ю. Ю., Чумаченко І. В., Некрасов І. Б., Худяков І. О., Доценко Н. В.**
Концептуальна модель моніторингу портфельного управління в проєктних офісах на основі блокчейну
- 29 **Дубчак Л. О., Вівчар Н. Т., Савенко О. С., Дериш Б. Б., Заставний О. М.**
Програмний засіб реалізації інтелектуального методу аналізу дефектів лопатей вітрових турбін з обмеженою вибіркою даних
- 40 **Золотухін О. В., Бодяньський Є. В., Філатов В. О., Кудрявцева М. С., Василюк О. О.**
Стекова гібридна система обчислювального інтелекту на основі ядерної активаційної функції та її онлайн-навчання в задачі розпізнавання образів
- 52 **Кучук Г. А., Матвєєв М. І., Косенко Н. В., Лисиця Д. О., Левченко А. О.**
Оцінювання продуктивності завантаження сцени WebAR-застосунку на мобільному пристрої
- 70 **Ланін М. О., Паржун Ю. В., Бохан К. О., Перевозник К. М., Александрова Т. Є.**
Інваріантно-структурне навчання: формування концептів як динаміка гіперграфових атракторів
- 95 **Невінський Д. В., Вихлюк Я. І.**
Проєктно-орієнтована методологія проактивного управління епідеміями на основі штучного інтелекту й моделей SEIRS
- 109 **Пєвнєв В. Я., Юдін О. В.**
Класифікація та експериментальні дослідження ефективності алгоритмів факторизації
- 121 **Пиріг Н. Я., Удовенко С. Г., Гриньова О. Є., Чала Л. Е.**
Автоматизована система аналізу пошкоджень об'єктів міської інфраструктури з використанням моделей комп'ютерного зору
- 153 **Радюк П. М., Саченко А. О., Мельниченко О. В., Бруханський Р. Ф., Каштальян А. С.**
WEDD-RST: інтерпретовний підхід до виявлення дефектів фотоелектричних панелей у кіберфізичних системах
- 173 **Фандакова М., Охотницький Ш., Коваленко А. А.**
Інтерактивна платформа підтримки клінічних рішень: архітектура та порівняльна оцінка загальної та точно налаштованої діагностики за симптомами за методологією (LLMS)
- 185 **Чепурна І. С.**
Модель управління ресурсами та кворумом блокчейн-орієнтованої розподіленої інформаційної системи в мультимедійному середовищі

Сучасні технології управління підприємством

- 200 **Назарова Г. В., Дем'яненко А. А., Назаров Н. К., Іванісов О. В., Остапчук В. В.**
Детермінанти розвитку людського потенціалу в умовах структурної трансформації промисловості

Інженерія та промислові технології

- 214 **Невлюдов І. Ш., Євсєєв В. В., Максимова С. С., Пащенко О. С., Косенко В. В.**
Розроблення моделі синхронізації роботи колаборативних роботів-маніпуляторів у HRC-сценаріях
- 227 **Алфавітний покажчик**

CONTENTS

Information Technology

- 5 **Barkovska O.**
Method for context-aware adaptation of a visual user interface based on estimating the distance to the user
- 15 **Husieva Y., Chumachenko I., Nekrasov I., Khudiakov I., Dotsenko N.**
A conceptual model for monitoring portfolio management in project offices based on blockchain
- 29 **Dubchak L., Vivchar N., Savenko O., Derysh B., Zastavnyy O.**
Software tool for implementing an intelligent method for analyzing wind turbine blade defects with limited dataset
- 40 **Zolotukhin O., Bodyanskiy Y., Filatov V., Kudryavtseva M., Vasylets O.**
Stacking Hybrid System of Computational Intelligence Based on Kernel Activation-Membership Functions and its Online Learning in Pattern Recognition Task
- 52 **Kuchuk H., Matvieiev M., Kosenko N., Lysytsia D., Levchenko A.**
Evaluating the performance of a webar application's scene loading on a mobile device
- 70 **Lapin M., Parzhyn Y., Bokhan K., Perevoznyk K., Aleksandrova T.**
Invariant structural learning: concept formation as hypergraph attractor dynamics
- 95 **Nevinskyi D., Vyklyuk Y.**
Project-oriented methodology for proactive epidemic management based on artificial intelligence and SEIRS models
- 109 **Pevnev V., Yudin O.**
Classification and experimental studies of the factorization algorithms efficiency
- 121 **Pyrh N., Udovenko S., Grinyova O., Chala L.**
Automated system for damage analysis of urban infrastructure objects using computer vision models
- 153 **Radiuk P., Sachenko A., Melnychenko O., Brukhanskyi R., Kashtalian A.**
WEDD-RST: An interpretive approach to detecting defects in photovoltaic panels in cyber-physical systems
- 173 **Fandakova M., Ochotnickyy S., Kovalenko A.**
Interactive platform for supporting clinical decision-making using large language models (LLMs)
- 185 **Chepurna I.**
Model of resource and quorum management for a blockchain-oriented distributed information system in a multi-cloud environment

Modern Enterprise Management Technology

- 200 **Nazarova G., Demianenko A., Nazarov N., Ivanisov O., Ostapchuk V.**
Determinants of human potential development in the conditions of structural transformation of industry

Engineering & Industry Technologies

- 214 **Nevliudov I., Yevsieiev V., Maksymova S., Pashchenko O., Kosenko V.**
Development of a model for synchronizing the operation of collaborative robotic manipulators in HRC scenarios
- 227 **Alphabetical index**

The author is responsible for the accuracy of the facts, quotations and other information

Olesia Barkovska

METHOD FOR CONTEXT-AWARE ADAPTATION OF A VISUAL USER INTERFACE BASED ON ESTIMATING THE DISTANCE TO THE USER

The relevance of this research topic stems from the need to transition from static interfaces to flexible systems capable of dynamically adapting to the user's physical context, particularly to changes in viewing distance, thereby reducing cognitive load and enhancing the comfort of interaction. **The object of the study** is the process of context-aware adaptation of the user visual interface in real time depending on the distance to the user, specifically the typographic parameters of the user interface. **The subject of the study** is approaches to determining the distance to the user based on video stream data. **The aim of the study** is to develop a method for context-aware adaptation of the user interface based on determining the distance to the user, which ensures dynamic scaling of the typographic system to reduce cognitive load and enhance interaction comfort in real time. To achieve the set goal, the following **tasks** were addressed: existing approaches to determining the distance to an object were analyzed, and a method for dynamically changing the scale of the user visual interface was developed, based on the proportional relationship between the distance from the interface to the user and the size of the base typographic unit (point size). To experimentally verify the high accuracy of determining the distance from the visual interface to the user, the study examined the impact of various factors (head tilt in the frontal and sagittal planes, rotation in the horizontal plane), as well as the chosen approach to determining the distance to the object, on the accuracy of distance measurement. **The developed method** ensures the adaptation not only of font size but also of related typographic parameters (line spacing, character spacing, and paragraph indentation). **It has been established** that the approach based on a depth map with approximation provides the most balanced results in terms of accuracy and practical applicability, demonstrating an error of 1.83–3.67% in the range of 30–90 cm. It was also shown that the user's head orientation is a critical factor affecting measurement accuracy, and the maximum error at close range reaches 201.3%.

Keywords: interface; distance; computer vision; depth map; model; method.

1. Introduction

Human-computer interaction is a field of study that examines how people interact with computers and how effective that interaction is. The effectiveness of a human-computer interaction system is determined by high levels of system reliability; ensuring that the user remains productive over the long term; and facilitating the user's professional development.

From the user's perspective, the functional aspects include control (the user's initiative, actions, and commands directed at the system), the interface (the point of contact, which may be graphical, voice-based, tactile, etc.), and user information (data that the system collects and receives about the user through the interface. This may be explicit information or implicit behavioral data), and evaluation (the stage at which the user interprets the system's response received through the interface).

At the same time, traditional visual user interfaces largely remain static and do not account for changes in the physical context of interaction, particularly the

distance between the user and the screen. Under such conditions, even a well-designed interface can lose its ergonomic effectiveness, as fixed typographic parameters do not ensure adequate readability and ease of information perception when the user's position relative to the device changes. This is particularly relevant for modern digital systems that operate under variable usage conditions and involve prolonged visual interaction. Such adaptation is particularly significant for users with visual impairments, for whom the ease of perceiving visual information, the readability of text elements, and the ability to automatically adjust display settings directly impact the accessibility of digital services. In this context, dynamic interface adaptation can be viewed as a means of enhancing inclusivity and facilitating the integration of people with visual impairments into the digital environment.

One promising approach to solving this problem is the use of computer vision methods to estimate the distance to the user in real time, followed by dynamic adaptation of display parameters. This approach enables a transition from static interface solutions to context-

aware systems capable of automatically adjusting not only font size but also related typographic characteristics in accordance with perceptual conditions. In this context, the scientific interest lies in selecting a distance estimation method that ensures sufficient accuracy, resistance to changes in external conditions, and practical suitability for integration into adaptive user interfaces.

Based on the above, it is relevant to develop a method for context-aware adaptation of a visual user interface, in which the estimation of distance to the user is used as a contextual trigger for dynamic scaling of the typographic system.

Previous works dedicated to autonomy and increasing inclusivity for users with disabilities and various types of impairments include [1, 2]. The adaptation of visual interfaces is a logical continuation of these works.

This study focuses on context-aware adaptation of a visual interface based on estimating the distance to the user and dynamically scaling typographic parameters.

2. Analysis of the latest research and publications

Current scientific works propose various classifications of adaptive systems, justify the need for personalization – especially in educational environments – and focus on the user model (their knowledge, goals, and preferences) [3, 4]. Adaptation here refers to a response to the user's internal state (e.g., their level of knowledge). A number of articles propose system architectures that adapt to device characteristics (screen size, memory), providing a logical framework of rules [5–7]. These works also identify the primary beneficiaries of adaptation – people with disabilities (cognitive, visual impairments, etc.) [7–9]. Adaptation here (e.g., changing fonts, simplifying the UI) is viewed as a tool for overcoming accessibility barriers. The works [2, 7] describe adaptation rules but assume that the system already knows the context for activating the rules; however, no attention is given to the technology that activates these rules.

In a broad sense, an adaptive interface involves adjusting information presentation parameters to the user's characteristics and interaction conditions. For visual user interfaces, this specifically means the ability to dynamically change typographic parameters, whereas in a more general sense, adaptation may also encompass auditory, linguistic, or other assistive means of information presentation aimed at improving accessibility for various categories of users [10, 11].

This confirms the relevance of a dynamic approach, where, for example, a user may have no vision problems but be physically far from the screen, which also requires adapting the information displayed on the screen to the physical environment in real time.

Also relevant is the visual modeling of the context-adaptive system to demonstrate the location of the sensory trigger, as well as to prove that unstable adaptation is worse than no adaptation at all. Accurately determining the distance to an object and ensuring smooth font transitions is a critical task and will reduce cognitive load.

Thus, developing a method with increased accuracy in distance determination is not merely a technical task, but an important step toward transforming existing theoretical models of user interface adaptation into practically applicable and reliable tools.

3. Formulation of the research objective

The objective of this study is to develop a method for context-aware user interface adaptation based on determining the distance to the user, which enables dynamic scaling of the typographic system to reduce cognitive load and enhance the comfort of real-time interaction.

Let d be the estimated distance from the user to the screen, and $P(d)$ be the vector of typographic parameters of the interface, which includes:

- font size;
- line spacing;
- character spacing;
- paragraph indentation.

Then the adaptation problem consists in constructing a transformation $F: d \rightarrow P(d)$, where F is an adaptation function that establishes a relationship between the estimated distance d from the user to the screen and the vector of typographic interface parameters $P(d)$.

To achieve this goal, the following tasks must be solved:

- analyze existing approaches to determining the distance to an object (user);
- develop a method for dynamically scaling the user visual interface based on a proportional relationship between the distance from the interface to the user and the size of the base typographic unit (font size), from which all related parameters (line spacing, character spacing, and paragraph indents) are calculated;

– experimentally investigate the influence of variable factors (head tilt in the frontal and sagittal planes, rotation in the horizontal plane), as well as the chosen approach to determining the distance to the object, on the accuracy of distance measurement; analyze the obtained results.

4. Materials and methods

The mathematical basis of the proposed method is a linear scaling model, which establishes a directly proportional relationship between the physical distance to the user and the final font size (point size) of interface elements. The target font size S is calculated using the following formula:

$$S_{final} = S_{base} \times (L_{current} / L_{base}),$$

where S_{final} – the final (adapted) font size to be displayed to the user; S_{base} – the base, or reference, font size (e.g., 16pt). This is the size considered optimal for

reading at the reference distance; $L_{current}$ – the current distance to the user, dynamically measured by the system (e.g., in centimeters); L_{base} – the base (reference) distance (e.g., 50 cm) for which S_{base} is calculated.

Let's introduce a scaling factor $K = L_{current} / L_{base}$ to determine how much the current distance differs from the reference distance. If $K = 1$, the user is at the reference distance, and the font remains the base size. If $K > 1$, the user is farther away than expected. The system proportionally enlarges the font to maintain readability. If $K < 1$, the user is closer. The system proportionally reduces the font size to avoid excessive visual strain and fit more information on the screen.

In addition to distance, it is important to consider the lighting in the room where the user is located in order to reduce cognitive load (Table 1).

The proposed structural-functional diagram of the method for dynamically scaling the user interface based on the estimated distance to the user is shown in Figure 3.

Table 1. The dependence of sign legibility on distance and lighting

Distance from the eyes	Lighting	Minimum font size (body text) mm (pt)	Font size for important text mm (pt)
Up to 50 cm	200–500 lux	2,5–3 mm (7,1–8,5 pt)	3–5 mm (8,5–14,2 pt)
Up to 50 cm	>500 lux	2–3 mm (5,7–8,5 pt)	2,5–5 mm (7,1–14,2 pt)
50–100 cm	200–500 lux	3–5 mm (8,5–14,2 pt)	5–7 mm (14,2–19,8 pt)
50–100 cm	>500 lux	2–3 mm (5,7–8,5 pt)	3–5 mm (8,5–14,2 pt)

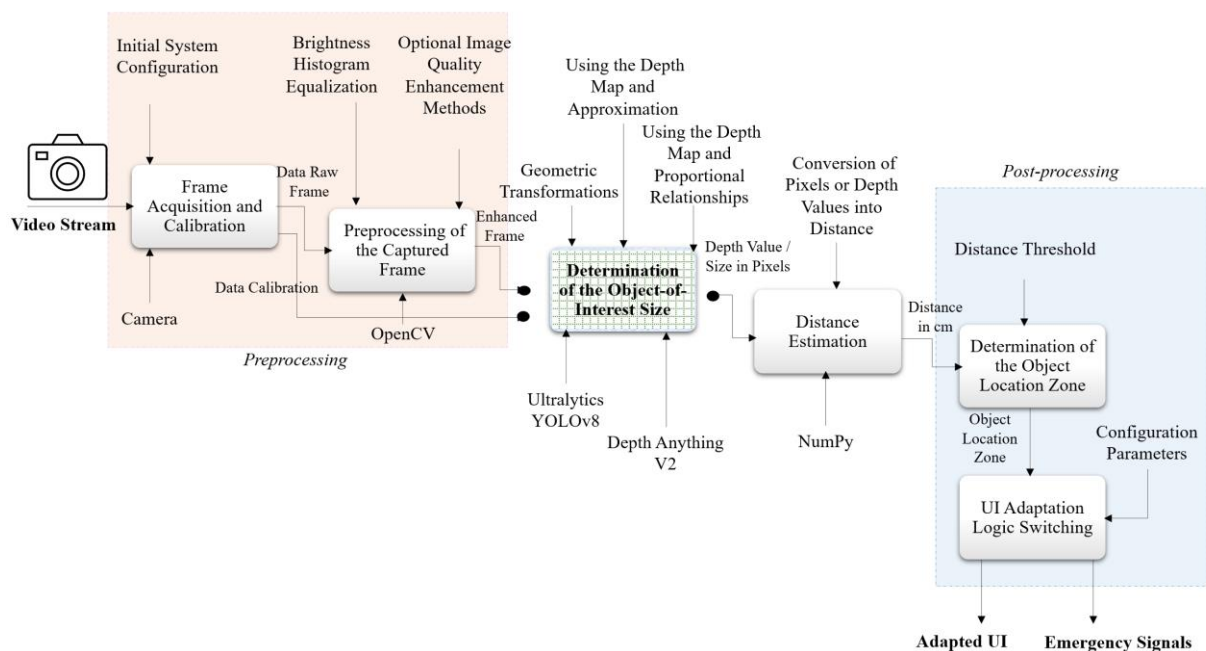


Fig. 1. Structural and functional diagram of the method for dynamically scaling a user interface based on the estimated distance to the user

The proposed method for dynamically scaling a user interface is based on estimating the distance to the user using data from the camera's video stream.

The input data for the described method consists of video frames from the camera, the initial configuration and calibration parameters of the video recorder, and image preprocessing rules for the subsequent main stage of the method – determining the distance to the user. This stage serves as the basis for deciding on further interface adaptation.

The output is an Adaptive User Interface (Adaptive UI), in which the scale and display logic of elements change according to the estimated distance between the user and the system.

The accuracy of the interface's adaptation to the distance from the user depends most on the object-of-interest sizing module, since its result serves as the input data for the distance estimation module. Errors arising at this stage lead to incorrect determination of the user's location and selection of the interface scaling mode. Thus, the more accurately the object's size and, accordingly, the distance to the user are determined, the more correctly Adaptive UI's dynamic scaling is performed. Therefore, further research focuses on refining algorithms to compensate for the effects of head rotations and tilts, particularly by using more accurate models for estimating a person's posture.

The introduced scaling factor K characterizes the deviation of the current user distance from the reference distance. Then the adapted font size is determined as

$$S_{final} = S_{base} \times K.$$

Other parameters of the interface's typographic system are scaled in the same way:

$$I_{final} = I_{base} \times K,$$

$$C_{final} = C_{base} \times K,$$

$$A_{final} = A_{base} \times K,$$

where I – line spacing, C – character spacing, A – paragraph indentation; I_{base} , C_{base} , A_{base} – their default values, and I_{final} , C_{final} , A_{final} – adapted values calculated based on the current distance from the user. This approach ensures consistent scaling of the entire typographic system of the interface and preserves the visual integrity of information presentation.

This paper presents an empirical study of various existing approaches for estimating the user's distance from the user interface to assess the accuracy of these approaches [12]:

- using geometric transformations [13, 14];
- using a depth map and approximation [15];
- using a depth map without approximation [16, 17].

The algorithm for the first approach studied is described in Figure 2. The simplified geometric approach determines the distance to an object based on the relationship between its size in pixels and the actual distance to the camera. After calibration in the range of 30–120 cm, the obtained data were approximated using a quadratic polynomial, which allowed for rapid estimation of the distance in centimeters. Despite its simplicity and speed, the method has limited accuracy due to its strong dependence on the stability of shooting conditions: camera position, viewing angle, and lighting.

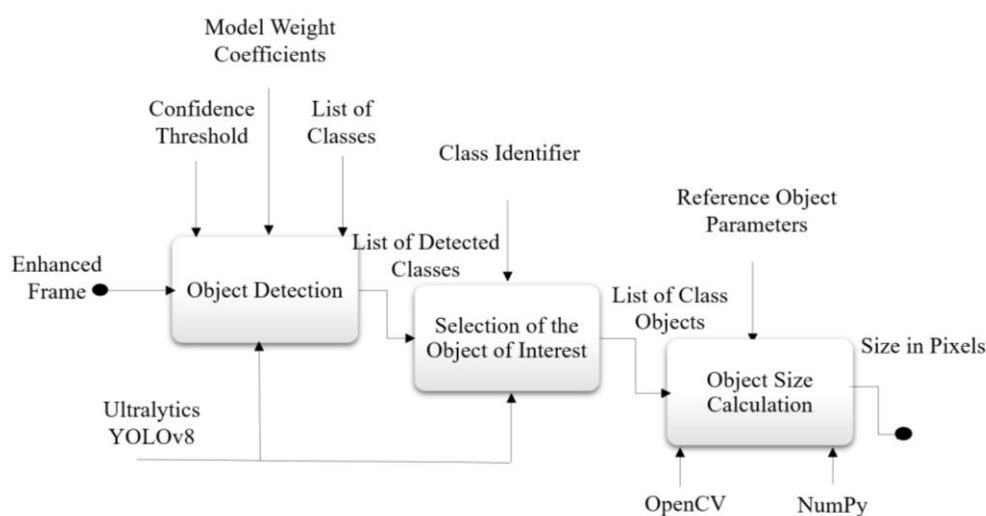


Fig. 2. Determining the size of the object of interest in the frame based on geometric transformations

The algorithm for the second approach studied is shown in Figure 3. It is based on a combination of object detection using YOLOv8 and analysis of depth maps from the Depth Anything V2 model. Unlike the first geometric method, which estimates distance based on the object’s size in pixels, this method directly measures the scene’s depth.

The algorithm includes the following steps:

1. Object detection in the frame (YOLOv8).
2. Selection of the target object of interest.

3. Determining the region of interest (ROI) for it.
4. Calculating the median depth value within this region from the depth map.
5. Converting relative depth values to absolute values (centimeters) through pre-calibration. This approach is independent of the object's size in the frame and is significantly less sensitive to changes in viewing angle, which potentially improves the accuracy of distance measurement in real-world conditions.

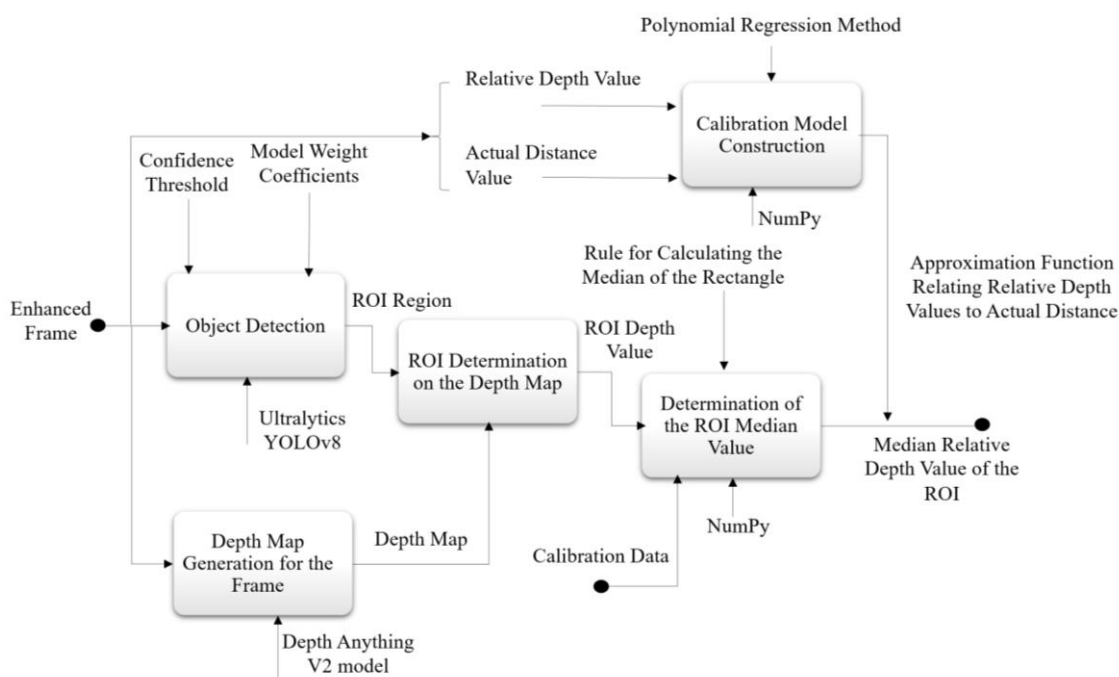


Fig. 3. Determining the size of the object of interest in the frame based on depth maps and approximations

The algorithm for the third approach studied (using a depth map without approximation) is similar to the second, as it is also based on constructing depth maps, but it includes a simplification: it omits the step of building a calibration model and does not use a function to approximate relative depth and actual distance. Instead, a reference depth is specified at a known distance.

5. Results and discussions

During the research, the influence of the user's actual distance from the screen, the lighting level in the room where the user is located, as well as the angle of head tilt in the frontal and sagittal planes was evaluated, corresponding to the real-world conditions of using the proposed and developed system. Figure 4 demonstrates the results of the studied methods, which use different

approaches and therefore have different computational complexity, different execution times, and different accuracy in distance determination.

The experimental validation of the proposed approach was conducted in two stages. In the first stage, a comparison was performed between three methods for determining the distance to the user: a method using geometric transformations, a method based on a depth map with approximation, and a method based on a depth map without approximation. The comparison was performed at four reference distances: 30, 60, 90, and 120 cm, corresponding to 12 controlled experimental configurations. In the second stage, for the method selected for further analysis, the effect of the user's head position was investigated: in the frontal plane, in the sagittal plane, and during rotation in the horizontal plane. For each type of orientation, the evaluation was

performed at three angles of deviation and four reference distances, resulting in an additional 36 configurations. Thus, the overall error assessment presented in this work is based on 48 controlled experimental configurations.

The first experiment focuses on determining the approach that yields the smallest error when measuring the user's distance from the screen. The experimental results are presented in Table 2.

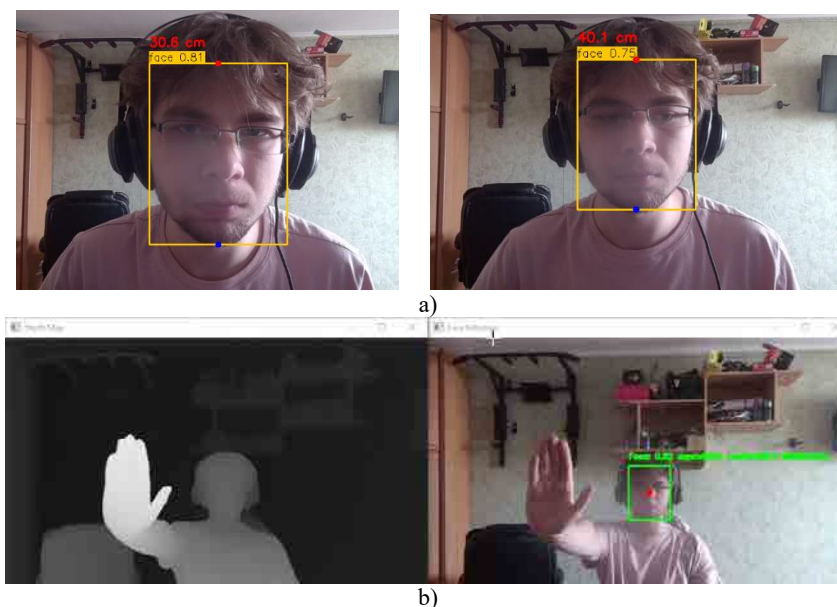


Fig. 4. Visual demonstration of the experimental results:

- a) display of the calculated distance at 30 and 40 cm using the geometric transformation method;
- b) display of the calculated distance at 120 cm using the method based on a depth map and approximation

Table 2. Error estimation for determining the user's distance from the screen depending on the selected approach and the actual distance from the screen

Method \ Distance	Real distance from the screen, cm			
	30	60	90	120
Method using geometric transformations	2%	0.33%	0.56%	8.92%
Method based on a depth map and approximation	3.67%	1.83%	2.11%	17.08%
Method based on a depth map without approximation	1.33%	1.83 %	13.67%	51.25%

Table 2 presents the results of an experimental evaluation of the error in determining the user's distance from the screen for three approaches: a method using geometric transformations, a method based on a depth map with approximation, and a method based on a depth map without approximation. The comparison was performed for four reference distances – 30, 60, 90, and 120 cm. The table shows the percentage error values for each approach at each of the specified distances, allowing the results of the first stage of experimental verification to be presented in a comparative format.

The aim of the next experiment is to assess the effect of head tilt in the frontal plane at various angles (0° , 15° , 45°), head tilt in the sagittal plane at various angles (0° , 15° , 30°), as well as head rotation in the

horizontal plane at angles (0° , 15° , 30°) on the accuracy of determining the distance to the screen (Figure 5). The results of the second experiment are presented in Table 3.

Table 3 presents the results of an experimental assessment of the error in determining the user's distance from the screen depending on the position of the head in space. The table shows the error values for three head orientation scenarios: tilt in the frontal plane, tilt in the sagittal plane, and rotation in the horizontal plane. For each option, several angles of deviation were considered, and the evaluation was performed at four reference distances – 30, 60, 90, and 120 cm. The percentage error values presented in the table allow for a general overview of the results of the second stage of the experimental study for various spatial positions of the user's head.

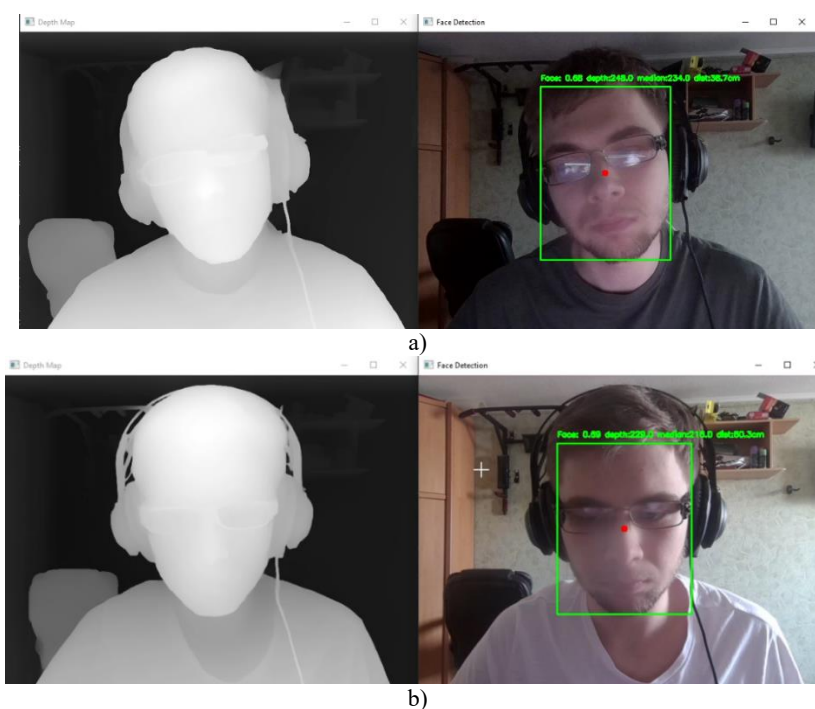


Fig. 5. Visual demonstration of the results of the experiments:
 a) distance measurement at 30 cm with the head tilted 15° in the frontal plane;
 b) distance measurement with the head tilted 15° in the sagittal plane

Table 3. Estimation of the error in determining the user's distance from the screen depending on head position in space and actual distance from the screen, %

Plane	Angle	30 cm	60 cm	90 cm	120 cm
Frontal plane, tilt	0°	25.3 %	10.7 %	3.6 %	1.1 %
	15°	19.0 %	20.0 %	10.8 %	4.9 %
	45°	40.0 %	8.8 %	2.8 %	2.3 %
Sagittal plane, tilt	0°	1.67 %	7.0 %	4.0 %	3.3 %
	15°	100.7 %	18.7 %	33.1 %	31.2 %
	30°	137.0 %	2.3 %	47.7 %	13.1 %
Rotation in the horizontal plane	0°	7.3 %	3.2 %	4.0 %	0.4 %
	15°	201.3 %	89.2 %	21.0 %	4.5 %
	30°	127.7 %	62.2 %	26.0 %	10.7 %

6. Discussion of Results

Analysis of the results of the first experiment (Table 2) demonstrates that for short and medium distances, all three approaches provide satisfactory accuracy. However, as the distance increases, the use of geometric transformations demonstrates the highest stability. A drawback of this approach is that it can only determine the distance for specific objects, for which the system was previously calibrated. Changing the focal length also reduces the reliability of this approach, as it requires a complete recalibration of the system. The third approach exhibits the highest error values. The second approach does not require a complete recalibration of the

system when the user changes, so it is recommended for subsequent experiments.

Analysis of the results (Table 3) showed a clear dependence of distance measurement accuracy on the user's head position. The most critical factor was a distance of 30 cm, where extreme error values were recorded – up to 201.3% with a 15° head rotation and 137.0% with a 30° sagittal tilt. Head rotations in the horizontal plane consistently produce the highest errors among all types of movements.

The lowest error values are generally observed at zero head angles compared with non-zero tilt and rotation conditions. At the same time, the results are not uniformly stable across all distances: for example, in the frontal plane at 30 cm the error reaches 25.3%,

while at 60–120 cm it decreases to 10.7%, 3.6%, and 1.1%, respectively. There is also a noticeable positive trend of decreasing error with increasing distance – at 120 cm, accuracy improves significantly even at significant angles of deviation.

Based on the results presented in Table 2 and Table 3, the second research objective, which was to determine the effect of the user's head position on distance measurement accuracy, was experimentally resolved.

The results obtained indicate the need for additional correction of measurement algorithms, especially for close distances, through the implementation of compensation mechanisms that account for the orientation of the head in space.

The developed methods and results of the experiments have the potential to be used in the creation of adaptive visual interfaces for information systems of various purposes, including educational platforms, medical services, banking and administrative terminals, mobile and desktop applications, as well as digital services for users with visual impairments. The development holds particular value for industry, as it can be integrated into HMI solutions, operator panels, smart displays, and systems for monitoring and controlling production processes, where readability, speed of information perception, and reduction of cognitive load on the operator are critical. In this context, the method can be viewed as a tool for improving the ergonomics, accessibility, and reliability of digital interaction in modern industrial and service environments.

7. Conclusions

During the course of this study, a method for context-aware interface adaptation was developed, the main functional element of which is determining the distance to the user, based on which the font size and related typographic parameters of the visual user interface (line spacing, character spacing, and paragraph indentation) are dynamically adjusted.

It has been experimentally proven that the approach based on depth maps and approximation (YOLOv8 + Depth Anything V2) is the most balanced solution. It combines acceptable accuracy (an error of 1.83–3.67% at distances of 30–90 cm) with resistance to changes in shooting conditions and the absence of the need for complete recalibration of the system when the user changes, unlike the geometric method.

A critical influence of the user's head orientation on measurement accuracy was identified. The largest errors (up to 201.3%) were recorded at close range (30 cm) when the head was rotated in the horizontal plane. This indicates the need to integrate compensation mechanisms that account for the user's posture in real time.

The scientific novelty lies in the development of a method for context-aware adaptation of a visual user interface, in which the estimation of the distance to the user is used as a contextual trigger for dynamically scaling not only the font size but also related typographic parameters, as well as in establishing the influence of the user's head orientation on the accuracy of distance measurement.

The practical significance lies in the possibility of applying the proposed method during the development of adaptive visual interfaces that automatically adjust display parameters depending on the distance to the user and can be used to improve readability, ease of information perception, and the accessibility of digital services.

Further research will focus on refining algorithms to compensate for the effects of head rotations and tilts, particularly by using more accurate models for estimating a person's posture. The integration of additional contextual factors, such as more accurate determination of ambient light levels and their automatic incorporation into the scaling model in the future, will also ensure greater system resilience to changing conditions.

Conflict of interest

The author declares that there is no conflict of interest, particularly of a financial, personal, authorial, or any other nature, that could influence the research or the results published in this article.

Funding

The research was conducted without financial support.

Data availability

The manuscript contains data in the form of electronic supplementary material.

Use of artificial intelligence

The authors confirm that they did not use artificial intelligence technologies to write this work.

References

1. Barkovska, O. (2025), "Context-aware, adaptive HMI technology for enhancing the autonomy of users with disabilities", *Scientific Notes of Taurida National V.I. Vernadsky University. Series: Technical Sciences*, Vol. 2, No. 6, pp. 32–38. DOI: <https://doi.org/10.32782/2663-5941/2025.6.2/06>
2. Barkovska, O. (2025), "Formal description of interaction and data flows in multimodal assistive systems for user autonomy support", *Visnyk of Kherson National Technical University*, Vol. 3, No. 4(95), pp. 15–20. DOI: <https://doi.org/10.35546/kntu2078-4481.2025.4.3.2>
3. Medina-Medina Nuria, Lina García-Cabrera (2016), "A Taxonomy for User Models in Adaptive Systems: Special Considerations for Learning Environments", *The Knowledge Engineering Review*, Vol. 31, No. 2, pp. 124–41. DOI: <https://doi.org/10.1017/S0269888916000035>
4. Kravcik, M., Angelova, G., Ceri, S., Cristea, A., Damjanović, V. et al. (2005), "Requirements and Solutions for Personalized Adaptive Learning", available at: [fihal-00590961](https://arxiv.org/abs/fihal-00590961)
5. Christian, S., Stephanidis, C. (2004), "User-Centered Interaction Paradigms for Universal Access in the Information Society", *8th ERCIM Workshop on User Interfaces for All*, Vol. 3196, 485 p. DOI: <https://doi.org/10.1007/b95185>
6. Silega Nemury, Gilberto Fernando Castro Aguilar, Inelda Anabelle Martillo Alcívar, Faggioni, K., Rogozov, Y., Lapshin, V. (2023), "An Ontology-Based Approach to Support the Knowledge Management of Software Quality Standards", *Enfoque UTE*, Vol. 14, No. 3, July 2023, pp. 49–56, DOI: <https://doi.org/10.29019/enfoqueute.946>
7. Miñón, R., Paternò, F., Arrue, M. et al. (2015), "Integrating adaptation rules for people with special needs in model-based UI development process", *Univ Access Inf Soc*. Vol. 15, pp. 153–168 DOI: <https://doi.org/10.1007/s10209-015-0406-3>
8. Heumader, P., Miesenberger, K., Murillo-Morales, T. (2020), "Adaptive User Interfaces for People with Cognitive Disabilities within the Easy Reading Framework", *Computers Helping People with Special Needs*. pp. 53–60. DOI: https://doi.org/10.1007/978-3-030-58805-2_7
9. Firmenich, S., Garrido, A., Paternò, F., Rossi, G. (2019). "User Interface Adaptation for Accessibility", *Web Accessibility. Human–Computer Interaction Series*, pp. 547–568. DOI: https://doi.org/10.1007/978-1-4471-7440-0_29
10. Rania, Hamdani, Inès, Chihi, (2025), "Adaptive human-computer interaction for industry 5.0: A novel concept, with comprehensive review and empirical validation", *Computers in Industry*, Vol. 168, 104268 p. DOI: <https://doi.org/10.1016/j.compind.2025.104268>
11. Kandemir, H, Kose, H. (2021), "Development of adaptive human–computer interaction games to evaluate attention". *Robotica*. 40(1), pp. 56–76. DOI: <https://doi.org/10.1017/S0263574721000370>
12. Kalampukatt, P., Ghosh A., Kumar, V. (2024), "Integrating Object Detection and Distance Estimation: A Comprehensive Review of Techniques and Applications", *SSRN*, 23 p. DOI: <http://dx.doi.org/10.2139/ssrn.4984837>
13. Temneanu, M.C., Donciu, C., Serea, E. (2025), "Distance Measurement Between a Camera and a Human Subject Using Statistically Determined Interpupillary Distance", *AppliedMath*, Vol. 5(3), 118 p. DOI: <https://doi.org/10.3390/appliedmath5030118>
14. Chou, K.S., Wong, T.L., Wong, K.L., Shen, L., Aguiari, D., Tse, R., Tang, S.-K., Pau, G. (2023), "A Lightweight Robust Distance Estimation Method for Navigation Aiding in Unsupervised Environment Using Monocular Camera", *Appl. Sci.*, Vol. 13(19), 11038 p. DOI: <https://doi.org/10.3390/app131911038>
15. Jiashi, Feng, Zilong, Huang, Bingyi, Kang, Xiaogang, Xu, Lihe, Yang, Hengshuang, Zhao, Zhen, Zhao (2024), "Depth anything v2", *Advances in Neural Information Processing Systems*, Vol. 37. pp. 21875–21911. DOI: <https://doi.org/10.52202/079017-0688>
16. Masoumian, A., Marei, D., Abdulwahab, S., Cristiano, J., Puig, D., Rashwan, H. (2021), "Absolute Distance Prediction Based on Deep Learning Object Detection and Monocular Depth Estimation Models", *Frontiers in Artificial Intelligence and Applications*, pp. 325–334. DOI: <https://doi.org/10.3233/FAIA210151>
17. Chernyshov, D., Koziuberda, M. (2025), "Comparative framework for analyzing distance metrics in high-dimensional spaces", *Innovative technologies and scientific solutions for industries*, Vol. 1, No. 31, pp. 143–150. DOI: <https://doi.org/10.30837/2522-9818.2025.1.143>

Received (Надійшла) 17.11.2025

Accepted for publication (Прийнята до друку) 21.05.2026

Publication date (Дата публікації) 27.06.2026

Барковська Олеся Юрївна – кандидат технічних наук, доцент, Харківський національний університет радіоелектроніки, доцент кафедри електронних обчислювальних машин, Харків, Україна;

Olesia Barkovska – Candidate of Technical Sciences, Associate Professor, Kharkiv National University of Radio Electronics, Associate Professor of the Electronic Computers Department, Kharkiv, Ukraine;

e-mail: olesia.barkovska@nure.ua

ORCID ID: <https://orcid.org/0000-0001-7496-4353>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=24482907700>

МЕТОД КОНТЕКСТНО-ОБІЗНАНОЇ АДАПТАЦІЇ ВІЗУАЛЬНОГО КОРИСТУВАЦЬКОГО ІНТЕРФЕЙСУ НА ОСНОВІ ОЦІНКИ ВІДСТАНІ ДО КОРИСТУВАЧА

Актуальність теми дослідження зумовлена необхідністю переходу від статичних інтерфейсів до гнучких систем, здатних динамічно адаптуватися до фізичного контексту користувача, зокрема до зміни відстані від екрана, що дозволяє знизити когнітивне навантаження та підвищити комфорт взаємодії. **Об'єктом дослідження** є процес контекстно-обізнаної адаптації користувацького візуального інтерфейсу в реальному часі залежно від відстані до користувача, а саме – типографічних параметрів користувацького інтерфейсу. **Предметом дослідження** є підходи до визначення відстані до користувача на основі даних відеопотоку. **Метою дослідження** є розробка методу контекстно-обізнаної адаптації користувацького інтерфейсу на основі визначення відстані до користувача, що забезпечує динамічне масштабування типографічної системи для зниження когнітивного навантаження та підвищення комфорту взаємодії в реальному часі. Для досягнення поставленої мети було вирішено такі **завдання**: проаналізовано існуючі підходи до визначення відстані до об'єкта, розроблено метод динамічної зміни масштабу користувацького візуального інтерфейсу, що ґрунтується на пропорційній залежності між відстанню інтерфейсу до користувача та розмірами базової типографічної одиниці (кеглю). Для експериментального підтвердження високої точності визначення відстані від візуального інтерфейсу до користувача було досліджено вплив варіативних факторів (нахил голови у фронтальній та сагітальній площинах, поворот в горизонтальній площині), а також обраного підходу до визначення відстані до об'єкта на точність вимірювання відстані. **Розроблений метод** забезпечує адаптацію не лише розміру шрифту, а й пов'язаних із ним типографічних параметрів (міжрядковий та міжлітерний інтервали, абзацні відступи). **Встановлено**, що підхід на основі глибинної карти з апроксимацією забезпечує найбільш збалансовані результати за точністю та практичною придатністю, демонструючи похибку 1.83–3.67% у діапазоні 30–90 см. Також показано, що орієнтація голови користувача є критичним фактором впливу на точність вимірювання, а максимальна похибка на близькій дистанції сягає 201.3%.

Ключові слова: інтерфейс; відстань; комп'ютерний зір; карта глибини; модель; метод.

Бібліографічні описи / Bibliographic descriptions

Барковська О. Ю. Метод контекстно-обізнаної адаптації візуального користувацького інтерфейсу на основі оцінки відстані до користувача. *Сучасний стан наукових досліджень та технологій в промисловості*. 2026. № 2 (36). С. 5–14. DOI: <https://doi.org/10.30837/2522-9818.2026.2.005>

Barkovska, O. (2026), "Method for context-aware adaptation of a visual user interface based on estimating the distance to the user", *Innovative Technologies and Scientific Solutions for Industries*, No. 2 (36), P. 5–14. DOI: <https://doi.org/10.30837/2522-9818.2026.2.005>

Yuliia Husieva, Igor Chumachenko, Ivan Nekrasov, Illia Khudiakov, Nataliia Dotsenko

A CONCEPTUAL MODEL FOR MONITORING PORTFOLIO MANAGEMENT IN PROJECT OFFICES BASED ON BLOCKCHAIN

The subject of this study is the processes of ensuring transparency, accountability, and data integrity in project portfolio management systems based on distributed ledger technologies. **The objective** of this work is to develop a conceptual model, the Blockchain Portfolio Governance Model (BPGM), to enhance the transparency, integrity, and manageability of strategic portfolio management processes. **Objectives:** to develop a multi-level model architecture that combines traditional management cycles with cryptographic event logging mechanisms; to formalize management decisions as distributed ledger objects using asymmetric cryptography; to propose a comprehensive management quality assessment metric that accounts for both technical integrity and procedural discipline; to validate the model through simulation modeling of business processes. **Research methods:** systems analysis, methods of mathematical and simulation modeling in the Bizagi Modeler environment, asymmetric encryption, and hashing algorithms to ensure data integrity in distributed networks. **Results.** This paper proposes and justifies the architecture of the Blockchain Portfolio Governance Model, comprising five levels: governance, data aggregation, decision formalization, cryptographic integrity, and audit. A mathematical framework for event logging has been developed, where each decision is signed using the ECDSA digital signature algorithm. A new comprehensive metric has been introduced – the Portfolio Governance Compliance Index – which enables the detection of "shadow" management actions by comparing the number of requests initiated in external systems with the number of validated transactions on the blockchain. A series of simulation experiments demonstrated that implementing Proof-of-Authority consensus algorithms in a corporate network introduces negligible time delays (less than 1% of the total cycle), while the majority of the process time is spent on expert analysis. **Conclusions:** The application of the BPGM model enables transforming subjective portfolio management into a transparent, algorithmic process. The proposed solution ensures the creation of a «single source of truth» for all stakeholders, significantly simplifies audit procedures, and enhances the organization's institutional reliability without compromising its operational efficiency.

Keywords: blockchain; project portfolio management (PPM); management transparency; distributed ledger technology (DLT); business process simulation; portfolio management compliance index (PGCI); data accountability.

Introduction

The rapid increase in the complexity of project activities, the accelerating pace of change in the business environment, and the need for more transparent strategic decision-making are driving the search for new approaches to project portfolio management. Under these conditions, traditional management mechanisms do not always ensure a sufficient level of traceability, consistency, and auditability of management decisions, especially when it comes to multi-level approvals, shifting priorities, and resource reallocation. This creates a need for models that simultaneously combine process formalization, data integrity control, and the ability to quantitatively analyze management effectiveness.

The results of the Project Management Institute (PMI) study "Step Up: Redefining the Path to Project Success with M.O.R.E.: Ready to unlock transformative value for your teams, your organization, and beyond?" in 2025 confirm that project success depends to a large extent on the availability of systematic monitoring,

regular reassessment of results, and alignment of projects with the organization's strategic goals. This justifies the need to establish transparent mechanisms for monitoring the execution of project portfolios, which ensure the timely identification of deviations, support management decision-making processes, and enhance trust among stakeholders. According to the results of a PMI study, implementing such approaches can increase the project success rate based on stakeholders' perception of value (NPSS) from the current average of 47% to 94% [1].

Traditional centralized management systems, including ERP solutions and specialized project management tools, provide advanced planning and monitoring capabilities; however, they remain vulnerable to data manipulation and do not guarantee the immutability of decisions made. This is particularly critical at the portfolio management level, where strategic decisions regarding budget reallocation, changes in priorities, and timelines are often made without formalized validation mechanisms or a transparent audit trail. As a result, conflicts arise among stakeholders, trust in management outcomes declines,

and reconstructing an objective picture of the decisions made becomes practically impossible.

In this context, blockchain technology becomes particularly relevant as a tool for ensuring the immutability, transparency, and traceability of management events. Unlike traditional approaches, a distributed ledger ensures that every recorded decision cannot be altered or deleted without the knowledge of system participants, creating fundamentally new opportunities for auditing and verifying the integrity of portfolio management. The corporate sector's growing interest in permissioned blockchain platforms, particularly Hyperledger Fabric, indicates organizations' readiness to implement such solutions in real-world management practice [2]. This creates a need for models that combine the formalization of management decisions with cryptographic mechanisms for recording them, data integrity control, and the ability to quantitatively analyze management effectiveness.

Literature Review and Problem Definition

As the complexity of distributed project environments grows, ensuring transparency, accountability, and trust in project and project portfolio management becomes a critically important task. Traditional centralized management systems, including ERP solutions and tools such as Microsoft Project and Jira, provide advanced planning and monitoring capabilities; however, they remain vulnerable to data manipulation, do not ensure non-repudiation, and do not guarantee a single source of truth. These limitations are particularly evident at the project portfolio level, where strategic decisions regarding budget reallocation, schedule changes, and priority adjustments are often made without formalized validation mechanisms or a transparent audit trail, leading to conflicts among stakeholders and reducing management effectiveness.

Blockchain technology is viewed as a promising tool for addressing these issues due to its properties of immutability, decentralization, and traceability. At the same time, recent research highlights an important conceptual distinction between the governance of blockchain systems and governance through blockchain, with the latter area remaining under-explored in the context of project portfolio management.

An analysis of empirical case studies [3, 4] shows that existing research primarily covers the operational level of individual projects: management of information, supply chains, and contractual obligations dominates.

The PMBOK Guide's [4, 5] classification by knowledge areas confirms this trend – the highest representation is observed for communications and stakeholder management, while functions critical to portfolio management (integration, risks, scope of work) are minimally represented. The strategic level of portfolio management is virtually absent in these studies, indicating that blockchain's potential for ensuring the integrity of decisions at the portfolio level remains untapped.

An analysis of 95 scientific publications from 2015–2022 [4] noted a lack of research dedicated to project portfolio management (PPM) via blockchain. A systematic review [6] confirms that current approaches focus primarily on the technical aspects of blockchain platform operation, while issues of stakeholder accountability, the distribution of decision-making rights, and organizational management remain overlooked. A similar pattern is observed in the construction industry, where a study [7] covers a wide range of areas – from smart contracts and BIM to supply chain management and energy efficiency – but is also limited to the level of individual projects and does not extend to portfolio management. The absence of research on project portfolio governance in several independent reviews indicates that this is effectively an unexplored niche, which this study aims to address.

Individual empirical studies [8–10] confirm that blockchain integration improves data quality and traceability at the individual project level, but do not propose mechanisms for scaling these benefits to the portfolio level. Works [11–13], dedicated to smart contracts in infrastructure projects, demonstrate the effectiveness of automation at the operational level but create "islands of automation" not linked by a unified logic of portfolio decisions. The BPGM architecture proposed in this paper solves this problem through data aggregation and governance levels that ensure systematic control and synchronization across portfolio projects.

In the domestic literature [14], blockchain in project management is considered primarily in the context of Ukraine's reconstruction using maturity models (PMMM, OPM3, CMMI), but without formalized mechanisms for verifying management decisions through an immutable digital ledger. The BPGM model can serve as such a mechanism, complementing existing approaches to maturity assessment.

Thus, most current research focuses on the application of blockchain at the level of individual projects or the technical aspects of blockchain platforms,

while formalized portfolio management, ensuring the integrity of management decisions, and their quantitative assessment remain unexplored. In parallel, the field of AI-based project management is developing [15, 16], confirming the general trend toward the adoption of digital technologies in project management, particularly in turbulent environments. At the same time, AI- and blockchain-based approaches are complementary: while AI provides adaptability and forecasting, blockchain ensures the immutability and auditability of management decisions.

In this regard, this paper, as a logical continuation of the authors' previous research [1, 17–20], proposes a blockchain-based project portfolio management model (BPGM) aimed at formalizing management events, ensuring their immutability, and developing quantitative metrics for assessing portfolio integrity.

The proposed model is particularly relevant in the context of Ukraine's large-scale reconstruction, where, as of early 2026, over 8,000 reconstruction projects have been registered in the DREAM system. Reconstruction projects are implemented according to the "build back better" principle with the integration of sustainable development goals, which requires a multidimensional assessment of the results of both individual projects and their portfolios [21, 22]. At the same time, donors and international partners emphasize the need for a verifiable audit trail for every management decision.

Regarding research at the portfolio management level, the PMI standard defines portfolio management as a discipline that ensures project investments align with the organization's strategic goals and involves a clear distribution of authority and responsibility for decision-making [23]. At the same time, a study of seven decades of PPM evolution [24] indicates that despite a significant accumulation of methodological knowledge, practical tools for transparency and accountability at the portfolio management level remain underdeveloped – most of the proposed methods have proven insufficient or impractical in real organizational conditions. The authors [25] distinguish between portfolio management as an operational practice and portfolio governance as a strategic, continuous oversight process, emphasizing that studies explicitly focused on portfolio governance remain rare. This distinction is fundamental to this work: the BPGM model addresses specifically the governance level – the formalization and verification of strategic decisions – rather than the operational management of individual projects, which distinguishes it from most existing approaches.

Formulation of the research objective

The objective of this research is to develop a conceptual Blockchain Portfolio Governance Model (BPGM) to enhance the transparency, integrity, and manageability of strategic portfolio management processes.

To achieve this goal, the following research tasks must be addressed:

1. Develop a multi-level model architecture that combines traditional management cycles with cryptographic mechanisms for recording events.
2. Formalize management decisions as objects in a distributed ledger using asymmetric cryptography.
3. Propose a comprehensive management quality assessment metric that accounts for both technical integrity and procedural discipline.
4. Validate the model through simulation modeling of business processes.

Thus, the study aims to combine process modeling, simulation analysis, and blockchain mechanisms to create a tool that supports transparent, controlled, and auditable portfolio management.

Materials and methods

BPGM is a conceptual multi-level project portfolio management model in which strategic decisions affecting the budget, timelines, priorities, and resource allocation are formalized as events, cryptographically recorded in a distributed ledger, and made available for independent verification and audit.

The key idea behind BPGM is the transformation of subjective management agreements into verifiable digital facts: every decision made at the governance level undergoes formalization, cryptographic protection, and immutable recording in the ledger. This provides a single source of truth for all portfolio stakeholders and prevents retroactive changes to recorded events.

A set of methods was used to develop and validate the model. System analysis methods were applied to decompose the portfolio management process at the decision-making level and to define functional requirements for each level of the BPGM architecture. BPMN modeling was used to formalize and visualize the management process – from the initiation of a change request to the audit verification of the result – taking into account the interaction between the Portfolio Manager, PMO, steering committee, and blockchain components. The ECDSA and SHA-256 cryptographic algorithms

form the basis of the mechanism for digitally signing management events and constructing a hash chain, ensuring the immutability of records and the ability to independently verify their integrity. The Proof-of-Authority (PoA) consensus algorithm was chosen as the basis for the permissioned blockchain network, as it provides the necessary throughput and organizational accountability with a limited set of authorized validators – members of the steering committee or independent auditors. Simulation modeling was performed in the Bizagi Modeler environment – a specialized tool for simulating business processes built in BPMN notation.

It should be noted that the model is not intended to replace traditional portfolio management systems, but rather to enhance their trust, traceability, and accountability through a blockchain infrastructure with permissioned access.

The goal of BPGM is to ensure:

- the immutability of strategic decisions within the portfolio;
- transparency of changes to budgets, timelines, and priorities;
- control over the reallocation of resources between projects;
- the ability to conduct independent audits of decisions;
- a reduction in information asymmetry among stakeholders;
- increased institutional trust in portfolio management.

In this context, "management in project offices" refers to a set of governance processes implemented by the PMO at the portfolio level: coordination of interdependent project initiatives with a shared resource pool, standardization and validation of decision-making procedures (changes to budget, timelines, priorities), monitoring of alignment with the organization's strategic goals, and ensuring stakeholder accountability to sponsors and the steering committee. The class of tasks for which the BPGM model is designed involves

$$E = \{PID, Type, \Delta Cost, \Delta Time, DecisionBody, Timestamp, Signature\},$$

where *PID* is the project identifier; *Type* is the decision type; $\Delta Cost$ is the budget change; $\Delta Time$ is the schedule change; *DecisionBody* is the body or entity that made the decision; *Timestamp* is the time of recording; *Signature* is the digital signature.

This layer ensures the standardization of management events and establishes a unified logic for all portfolio decisions.

situations where multiple projects are executed simultaneously amid competition for resources, where every management decision potentially impacts the entire portfolio and requires permanent documentation for future audit. It is precisely under such conditions that the absence of a single verified source of management events most often leads to systemic failures: duplication of decisions, data inconsistencies among stakeholders, and the inability to reconstruct the chronology of strategic changes.

Results

BPGM comprises five levels that ensure an end-to-end process from data collection to audit:

1. Governance Layer

This layer is responsible for the governance context of decision-making. It is here that strategic actions are formulated regarding budget changes; schedule changes; priority changes; resource reallocation; and the approval or rejection of portfolio initiatives. This layer includes the portfolio manager; the governance committee; the PMO; sponsors; and financial and operational stakeholders.

2. Data Aggregation Layer

To minimize the risk of manipulation during the change initiation phase, this layer provides automated integration with traditional project management information systems (Microsoft Project, Jira) and corporate ERP solutions. Instead of manual entry, which does not guarantee objectivity, the model's multi-level architecture allows for the extraction of up-to-date metrics regarding budgets, resources, and the current status of tasks via API. This aggregated data serves as an objective foundation for the operation of the Decision Formalization Layer.

3. Decision Formalization Layer

Any management decision is translated into a formalized event record:

4. Cryptographic Integrity Layer

At this layer, the event is converted into a cryptographically secured record:

$$h = H(E),$$

where *H* is a hash function, such as SHA-256.

The technological foundation for recording these entries is a multi-tier smart contract architecture deployed in a permissioned environment (specifically, using

frameworks such as Hyperledger Fabric or private Ethereum networks). To ensure the necessary throughput and organizational accountability of transactions, the Proof-of-Authority (PoA) consensus algorithm is used. In this configuration, the right to validate blocks is granted exclusively to authorized nodes (e.g., members of the steering committee or independent auditors), which guarantees the irreversibility of management actions.

The choice of a permissioned blockchain (specifically Hyperledger Fabric) for the BPGM model is justified by the specifics of corporate project portfolio management (PPM), where the priority is on confidentiality, access control, and efficiency, rather than full decentralization.

A permissioned blockchain restricts access to authorized participants only (PMO, Portfolio Manager, committee), ensuring role-based permissions and identity verification, unlike permissionless blockchains (Bitcoin/Ethereum), where anyone can validate. This is critical for PPM, where decisions contain confidential information about budgets and strategies, preventing leaks via a public ledger.

The subsequent chain of events unfolds as follows:

$$B_n = E \{ e_n, h_n, h_{n-1} \},$$

where each new entry contains the event itself; its hash; and the hash of the previous entry.

This creates a sequential chain of events that makes it impossible to secretly alter decisions that have already been recorded.

5. Audit Layer

The audit layer ensures data integrity and consistency between the document and the registry entry:

$$V(D, h) = \begin{cases} \text{true}, & \text{if } H(D) = h \\ \text{false}, & \text{otherwise,} \end{cases}$$

where D is a document or dataset; h is a previously recorded hash.

This mechanism allows an independent auditor to verify whether the document has been altered since it was recorded.

Let the project portfolio be defined as:

$$P = \{P_1, P_2, \dots, P_n\},$$

where P_i is an individual project within the portfolio.

Then the set of all portfolio events is:

$$E = \bigcup_{i=1}^n E_i,$$

where E_i is the set of project events.

This means that the portfolio is viewed as a collection of related events, rather than merely as a set of projects.

The system is considered consistent if the following holds for every record:

$$h_n = H(E_n \| h_{n-1}),$$

where $\|$ means data concatenation.

Then the integrity of the entire system is defined as:

$$\text{IntegrityIndex} = \begin{cases} 1, & \text{if all event sare valid} \\ 0, & \text{otherwise.} \end{cases}$$

In practice, this means that any disruption in the sequence of events is immediately detected during verification.

A new metric is proposed to assess the quality and reliability of the portfolio registry:

$$PEII = \frac{\text{ValidEvents}}{\text{TotalEvents}},$$

where ValidEvents is the number of events that remained unchanged and passed verification; TotalEvents is the total number of recorded strategic events.

The proposed scale for interpreting $PEII$ values is presented in Table 1.

Table 1. Interpretation of $PEII$ values

$PEII$ values	Interpretation
0,90–1,00	High level of compliance
0,70–0,89	Acceptable level
< 0,70	Critical level – audit required

To improve responsiveness to business-critical decisions, a weighted version of the index is proposed:

$$PEII_w = \frac{\sum_{e \in \text{ValidEvents}} w_e}{\sum_{e \in \text{TotalEvents}} w_e} \times 100\%,$$

where w_e is the event weight (e.g., 3 for BudgetChange, 2 for ScheduleChange, 1 for other types).

To assess the quality and transparency of portfolio management within the BPGM model, a comprehensive metric is introduced – the Portfolio Governance Compliance Index ($PGCI$). The $PGCI$ allows for a quantitative assessment of an organization's management discipline:

$$PGCI = \frac{N_{val}}{N_{init}} \times PEII,$$

where N_{val} is the number of management transactions that have been successfully validated by consensus algorithms and recorded on the blockchain; N_{init} is the

total number of initiated requests for changes to the portfolio (including those that were rejected or deleted outside the system).

This approach allows for the identification not only of technical data integrity errors but also of "gray areas" in governance – cases where decisions are made or changed in circumvention of formalized BPGM procedures. A $PGCI$ value close to 1 indicates high process maturity and full transparency in portfolio management.

The implementation of the $PGCI$ integrated indicator allows for a distinction between the concepts of data integrity and governance compliance. While $PEII$ focuses on the immutability of existing records, the $PGCI = N_{val}/N_{init}$ identifies "off-system" management actions that have not been formalized in the distributed ledger. Thus, the $PGCI$ serves as a tool for detecting "shadow" project portfolio management.

To demonstrate the advantages of using the $PGCI$ over traditional technical integrity metrics, let's consider two contrasting scenarios for the operation of a portfolio management system. These scenarios allow us to distinguish between the technical robustness of a distributed ledger and actual management discipline within an organization.

Scenario A: *High technological reliability with low managerial discipline*

In this case, the organization uses a perfectly configured blockchain infrastructure that ensures absolute data immutability ($PEII = 1.0$). However, at the operational management level, protocols are being ignored: managers enter only 50% of the decisions actually made into the ledger, leaving the rest of the changes in the "shadow" zone (outside the BPGM).

Under these conditions, $PGCI = 0.5 \times 1.0 = 0.5$. A low index value signals management risks: despite the system's technical perfection, it does not fulfill the function of transparent control, since a significant portion of processes remains informal.

Scenario B: *High management discipline despite technical compromises*

In this scenario, the organization's management demonstrates a strong commitment to established procedures, recording every decision made on the blockchain ($N_{val}/N_{init} = 1.0$). However, due to insufficient decentralization (for example, a low number of active nodes in the private network), a technical error

or a targeted attack occurred, resulting in the corruption of 20% of historical records.

In this case, $PGCI = 1.0 \times 0.8 = 0.8$. Here, the index indicates a breach of technical trust. Although the organizational culture is strong, the registry data can no longer be considered an absolutely reliable source for auditing.

A comparative analysis confirms that the $PGCI$ is an effective tool for diagnosing the "health" of a governance system. It allows stakeholders to clearly identify the nature of the problems: whether they lie in the realm of human error and regulatory violations (Scenario A) or in the realm of technical infrastructure vulnerabilities (Scenario B). This enables the adoption of targeted corrective actions to improve the overall transparency of portfolio management.

Thus, the model can be presented as a sequence of layers: Governance Layer → Data Aggregation Layer → Decision Formalization Layer → Cryptographic Integrity Layer → Permissioned Blockchain Ledger → Audit Layer.

That is, a management decision arises in the portfolio environment, the decision is formalized as an event, the event is hashed and included in the chain of records, the record is stored in a permissioned blockchain, and an independent auditor performs an integrity check.

Thus, we propose the concept of portfolio management as a system of cryptographically verified events, formalize the mechanism for translating strategic decisions into events suitable for immutable storage, and introduce the $PEII$ and $PGCI$ metrics, which allow for the quantitative assessment of the level of integrity and trust in the portfolio registry.

BPMN modeling was used to formalize and visualize the process of management decisions in a portfolio context. The modeling results are presented in Figure 1.

The process encompasses the interaction between management, financial, control, and blockchain components. This structure corresponds to the logic of BPMN swimlanes, where each lane represents a separate responsibility or process participant.

The Portfolio Manager initiates a change request: budget, timeline, priority, or resource allocation. The reason for the change, the expected impact on the portfolio, and the justification for the committee are described. The request is forwarded to the PMO or committee for preliminary analysis. This stage creates an event source, which will subsequently serve as the basis for formalization and verification.

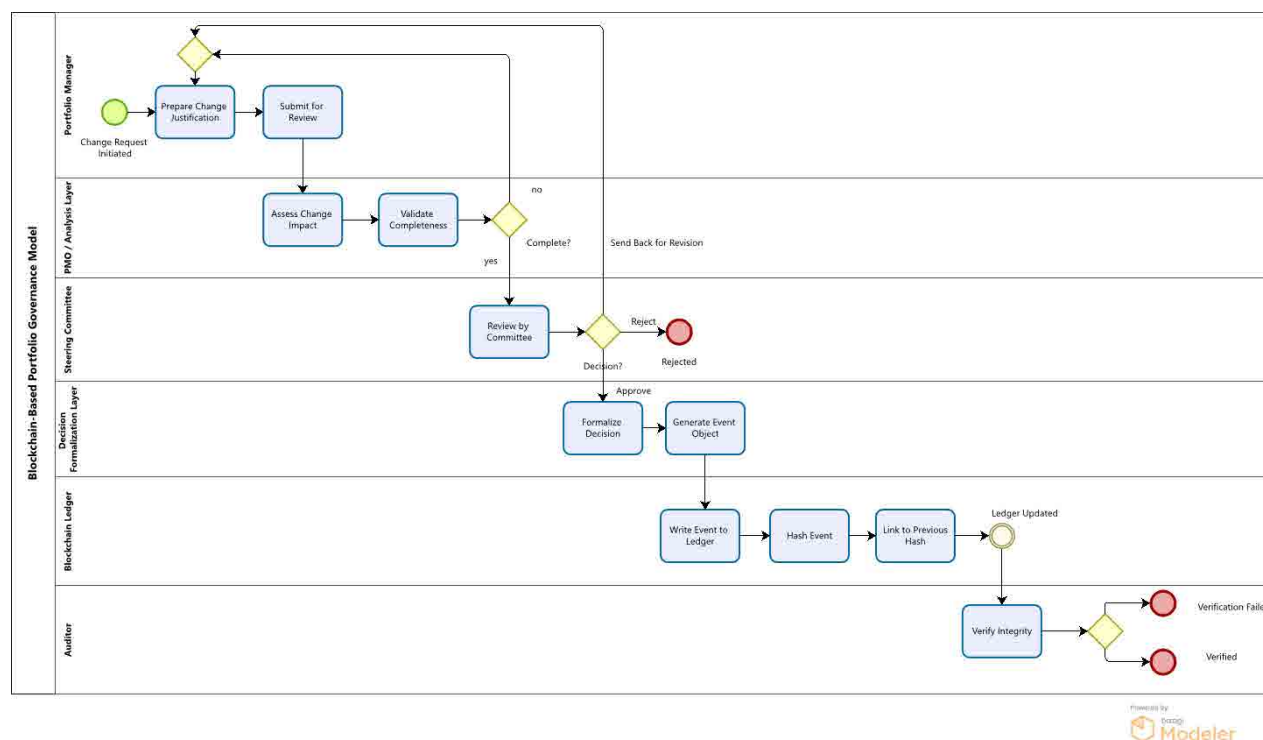


Fig. 1. BPMN-based process model of strategic change handling in the Blockchain-Based Portfolio Governance Model (BPGM)

The PMO analyzes the proposal's impact on the budget, timelines, risks, interdependencies between projects, and portfolio KPIs. It is verified whether there is sufficient data for review: the presence of calculations, signatures, financial estimates, and explanations. If the data is insufficient, the process is returned to the Portfolio Manager for revision. If the package is complete, it is forwarded to the committee for review.

Next, the committee reviews the request in terms of strategic feasibility, priority, and resource constraints. At this stage, three options are possible: Approve, Reject, Send Back for Revision. The corresponding gateway is key to the model, as this is where the management decision is formed, which then becomes an event at the blockchain level.

After the decision is approved, a structured record is generated:

- project ID;
- decision type;
- budget change;
- timeline change;
- decision-making body;
- timestamp;
- digital signature.

The event is transformed into a formal object for subsequent recording in the registry. It is here that the management decision is transformed into a blockchain-

ready event. The formalized event is recorded in a permissioned blockchain. The event's hash is calculated. The event is appended to the chain of previous records. Once the record is saved, the process proceeds to integrity verification. This ensures the immutability of strategic decisions and a transparent audit trail. The auditor verifies whether the recorded hash matches the current document or event. If everything is correct, the event is deemed valid. If not, a discrepancy is recorded. The audit result concludes the process. This stage links the blockchain mechanism with management accountability, which is one of the model's key advantages.

In the first stage, the BPMN model was modeled using only time parameters and branching probabilities, without explicit resource constraints. This approach allowed us to estimate the process duration, rework cycles, gateway behavior, and the resulting Portfolio Event Integrity Index (*PEII*) under various decision-making scenarios.

Tasks performed by humans were modeled using triangular distributions defined by minimum, most likely, and maximum durations, while technical blockchain operations were assigned a fixed duration due to their low variability (Table 2). The time parameters for the blockchain are conservative (intentionally inflated for simulation purposes), since in real Private Ethereum or

Hyperledger networks, this time typically does not exceed a few seconds (this demonstrates that even in the "worst-case" scenario, the blockchain network does not "slow down" business). Gateway routing was determined using branching probabilities for the outcomes of confirmation, revision, rejection, and integrity.

The initial modeling parameters were derived based on the process logic, expert estimates of the duration of individual operations, and assumptions regarding the probabilities of transition between stages. The time characteristics reflect the approximate duration of management actions within the model, while the probabilistic parameters were used to formalize possible process development scenarios at control points. This approach is appropriate in the absence of a complete set of empirical data, as it allows for the reproduction of realistic system behavior, the assessment of its sensitivity to changing conditions, and the comparability of results across scenarios.

The characteristics of the scenarios are presented below.

1. Optimistic. The process proceeds with almost no returns:

- approve: 70%.
- revision: 20%.
- reject: 10%.
- integrity pass: 99%.

2. Base case. The most realistic:

- approve: 60%.
- revision: 25%.
- reject: 15%.
- integrity pass: 98%.

3. Constrained. More revisions and delays:

- approve: 45%.
- revision: 35%.
- reject: 20%.
- integrity pass: 95%.

Table 2. Time characteristics of the process

Task	Distribution	Min	Mode	Max
Prepare Change Justification	Triangular	1	2	4
Submit for Review	Fixed	0,1	0,1	0,1
Assess Change Impact	Triangular	1	3	5
Validate Completeness	Triangular	0,5	1	2
Review by Committee	Triangular	2	4	7
Resubmit Change Request	Fixed	0,1	0,1	0,1
Formalize Decision	Triangular	0,5	1	2
Generate Event Object	Fixed	0,05	0,05	0,05
Write Event to Ledger	Fixed	0,05	0,05	0,05
Hash Event	Fixed	0,01	0,01	0,01
Link to Previous Hash	Fixed	0,01	0,01	0,01
Verify Integrity	Triangular	0,1	0,3	0,6

The simulation was performed using Bizagi Modeler. The simulation results for the base case scenario are shown in Figure 2.

The simulation results for the base case scenario showed that out of 451 initiated requests, 441 were completed within the model, and the average process duration was 17.32 days. The most time was spent on preparing the justification for changes, impact assessment, and committee review, while the stages of recording the decision in the ledger were completed without delays, confirming their operational ease.

The simulation also showed that some requests are filtered out during the completeness check or rejected after review, reflecting the presence of control points in the process. At the same time, 351 completed

cycles passed integrity verification, and only 6 cases ended unsuccessfully, indicating the high reliability of the data recording and verification mechanism in this scenario.

Simulation modeling confirmed that the proposed model ensures controlled movement of the request through all stages of portfolio approval, with the main time losses concentrated on management analysis rather than on blockchain operations. This means that blockchain integration does not complicate the process but enhances its transparency and control over the integrity of management decisions.

Within the scope of the study, the BPGM model was developed and formalized as a comprehensive architecture for supporting portfolio governance based

on a combination of process logic and blockchain recording mechanisms. The model reflects the sequence of management actions from the initiation of a change request to its verification, approval, or rejection, and also

provides for the immutable recording of critical events in a distributed ledger. This structure ensures the traceability of a decision at all stages of its processing and creates a basis for subsequent auditing.

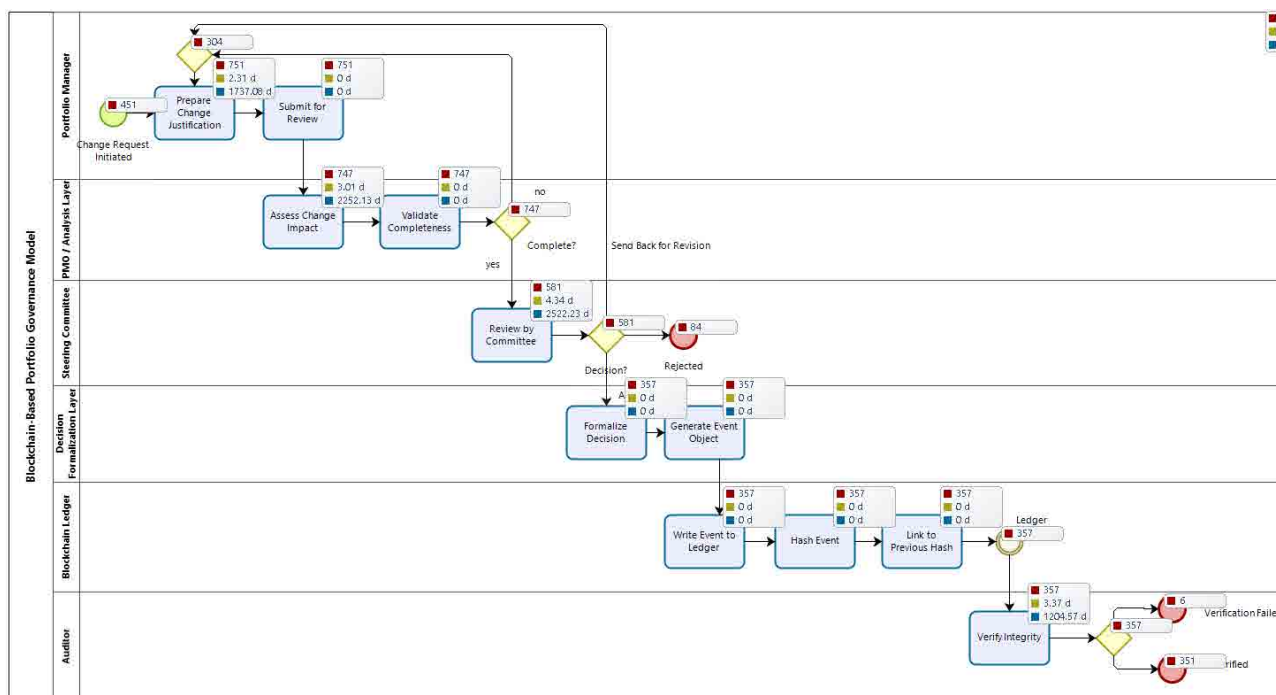


Fig. 2. Modeling results (Base case scenario)

One of the key outcomes is that the BPGM model allows for the separation of the actual managerial logic of decision-making from the technical aspects of verifying and documenting those decisions. As a result, every significant action within the process becomes a controllable event that can be verified based on temporal, procedural, and informational criteria. This increases the transparency of portfolio management and reduces the risks of loss, distortion, or duplication of critical data.

Simulation modeling of the BPGM BPMN process was performed for three scenarios: optimistic, base case, and constrained. In all scenarios, the time characteristics of tasks, the probabilities of transitions through gates, and the total duration of the portfolio decision process from request initiation to verification completion were analyzed. The shortest lead time was obtained in the optimistic scenario, while the constrained scenario was characterized by the longest duration due to a higher proportion of re-approvals and revisions.

In the base case scenario, the Review by Committee and Revise Request stages contributed the most to the total process time, indicating their role as the main time

bottlenecks. The Decision gate provided the main branching point of the process, with the majority of trajectories ending in approval, while a smaller portion proceeded to revision or rejection. The blockchain-recording stages – Write Event to Ledger, Hash Event, and Link to Previous Hash – had virtually no impact on the total duration due to their technical nature and low variability in execution time.

The *PEII* metric remained high in all scenarios, indicating that the integrity and controllability of portfolio events were maintained within the model's scenarios. The greatest decrease in *PEII* was observed in the constrained scenario, where the proportion of repeat cycles increased and the probability of a request being returned for revision rose.

Discussion of Results

BPGM is a conceptual model that combines portfolio management, event formalization, and blockchain integrity mechanisms. Its key advantage lies in the fact that strategic decisions cease to be merely managerial agreements and are transformed into verified digital facts.

This creates a foundation for transparent, controllable, and audit-ready management of project portfolios.

The model can be applied to the management of IT portfolios; construction and infrastructure programs; government digital transformation portfolios; cross-organizational projects; and audits of large programs where the traceability of decisions is critical. Its practical value lies in creating a single source of truth for strategic portfolio decisions.

Scaling the model depending on the complexity of the tasks being solved can be approached in several ways.

1. *Number of projects in the portfolio*

An increase in the number of projects increases the volume of transactions at the blockchain level, but thanks to PoA consensus, this load scales linearly and does not affect latency (confirmed by simulation – <1% of a cycle). However, the load on the governance level increases: more projects mean more cross-project dependencies, which requires expanding the pool of validators and refining the coordination protocols.

2. *Number of hierarchy levels and participants*

For simple portfolios (one PMO, two to three stakeholders), a minimal validator configuration is sufficient. For complex portfolios – corporate or government – multi-level approvals emerge (PMO → sponsor → regulator), and BPGM adapts by adding validator nodes without changing the architecture. This is a structural advantage of a permissioned blockchain.

3. *Frequency and diversity of management decisions*

The higher the rate of change and frequency of decisions (typical for IT portfolios in an Agile environment), the greater the value of a consistent audit trail – and the higher the sensitivity of *PGCI* to deviations. For slow-moving portfolios (infrastructure programs), the model remains valid, but the threshold values for *PGCI* and *PEII* may require calibration for a specific domain.

4. *Degree of project interdependence*

With weak interdependence, BPGM operates as a parallel set of independent decision chains. With strong interdependence, a data aggregation level is activated that tracks cross-project events. The complexity of the task here determines which part of the architecture is involved, but does not break the model.

The results indicate that the value of the BPGM model lies not only in the use of blockchain as a technological component, but primarily in building a managed portfolio governance model where the decision-making process becomes transparent and

verifiable. Unlike traditional approaches, in which management actions are often scattered among different participants and information systems, BPGM forms a single logical control loop. This allows for increased confidence in decision-making outcomes and ensures better portfolio manageability.

From a methodological standpoint, BPGM should be viewed as a platform for the further development of digital portfolio management. Its architecture creates the conditions for integration with analytical tools, decision support systems, and portfolio performance monitoring mechanisms. In this context, the model has practical value as a foundation for building more robust, transparent, and formalized management processes.

The results confirm that the proposed BPGM model is sensitive to changes in the logic of approval and the quality of the initial submission of changes. In practical terms, this means that even with a stable blockchain infrastructure, the overall effectiveness of governance is determined primarily by the organization of the process, rather than the technical layer of data storage. Simulation modeling confirmed that the implementation of blockchain tools does not create new technological bottlenecks. The main delay remains in the cognitive sphere – the time required for stakeholders to analyze the governance situation.

Thus, blockchain in a governance model should not replace the management process, but rather enhance its transparency and auditability. At the same time, the simulation itself shows that technical integrity does not compensate for poor-quality management input if requests are frequently returned for revision.

The implementation of BPGM in the real economy is accompanied by a number of limitations that must be taken into account:

1. *The Oracle Problem:* blockchain guarantees the technical immutability of data after it is recorded in the ledger, but it cannot automatically verify the accuracy of information at the stage of its initial entry. If a project manager intentionally enters false data about task status into the aggregation system, the BPGM model will cryptographically record this fact but will not detect the misinformation without additional external audit mechanisms.

2. *Complexity of integration with existing systems:* The effectiveness of the Data Aggregation Layer directly depends on the availability of open APIs in corporate software (ERP, CRM). In organizations with outdated IT infrastructure, automating data collection may be

technically challenging, requiring a temporary transition to semi-automated input methods.

3. *Cost-effectiveness and scalability*: Although simulation modeling confirmed low time costs, deploying and maintaining a private blockchain network (validation nodes) requires additional expenses for infrastructure and energy consumption. For small businesses with a limited number of projects, the administrative costs of maintaining the BPGM may outweigh the benefits of transparency.

4. *Digital Asset Management (Key Management)*: The model relies on asymmetric cryptography, which requires stakeholders to possess a high level of digital literacy in handling private keys. The loss or compromise of a key could lead to a blockage in the approval process or legal disputes regarding the authenticity of signatures.

The identified limitations outline avenues for future research, particularly regarding the development of algorithms for automatic verification of input data using artificial intelligence and the optimization of key storage mechanisms in the corporate sector.

Conclusions

This article proposes and justifies the BPGM model as a tool for enhancing the transparency, traceability, and integrity of portfolio management. Its architecture combines the process logic of management decisions with mechanisms for immutable event logging, creating a reliable foundation for controlling critical stages of portfolio management.

It is shown that the BPGM model allows formalizing the flow of a management request through all key stages – from initiation to result verification – and thereby reduces the risks of data loss, distortion,

or inconsistency. This is particularly important for environments where decisions are strategic in nature and require a high level of accountability.

The results obtained confirm the feasibility of further developing the BPGM as a foundation for the digitalization of portfolio management processes. In future research, it is advisable to expand the model by integrating analytical monitoring tools, performance evaluation mechanisms, and support for adaptive portfolio management.

Conflict of interest

The authors declare that they have no conflict of interest, including financial, personal, authorial, or any other nature, that could influence the research or the results published in this article.

Funding

The study was funded by the Ministry of Education and Science of Ukraine as part of research project 0125U001544 on the topic "Methodology for ensuring the processes of monitoring and controlling the implementation of project and program portfolios for project offices in the context of Ukraine's reconstruction".

Data availability

The manuscript has no associated data.

Use of artificial intelligence

The authors confirm that they did not use artificial intelligence technologies to write this paper.

References

1. Project Management Institute (2025), "Step Up: Redefining the Path to Project Success with M.O.R.E.: Ready to Unlock Transformative Value for Your Teams, Your Organization, and Beyond?", available at: <https://www.pmi.org/learning/thought-leadership/path-to-project-success> (accessed 18 April 2026).
2. Saari, A., Sinclair, S., Leshinsky, R. and Junnila, S. (2025), "Best practices for blockchain-driven digital transformation in cross-industry settings", *Digital Business*, 5(2), Article 100127. DOI: <https://doi.org/10.1016/j.digbus.2025.100127>
3. Husieva, Y., Chumachenko, I., Skachkov, O., Kadykova, I. and Dotsenko, N. (2025), "Using blockchain technology for monitoring projects", in *Integrated Computer Technologies in Mechanical Engineering – 2024 (ICTM 2024)*, Lecture Notes in Networks and Systems, Vol. 1473, Cham: Springer, pp. 461–472. DOI: https://doi.org/10.1007/978-3-031-94845-9_38
4. Sönmez, R., Sönmez, F.Ö. and Ahmadisheykhsarmast, S. (2023), "Blockchain in project management: a systematic review of use cases and a design decision framework", *Journal of Ambient Intelligence and Humanized Computing*, 14, pp. 8433–8447. DOI: <https://doi.org/10.1007/s12652-021-03610-1>

5. Project Management Institute (2021), *PMBOK® Guide: A Guide to the Project Management Body of Knowledge*. 6th ed., Newtown Square, PA: Project Management Institute. Available at: <https://www.pmi.org/standards/pmbok> (accessed 18 April 2026).
6. Liu, Y., Lu, Q., Zhu, L., Paik, H.Y. and Staples, M. (2023), "A systematic literature review on blockchain governance", *Journal of Systems and Software*, 197, Article 111576. DOI: <https://doi.org/10.1016/j.jss.2022.111576>
7. Shishehgarhaneh, M.B., Moehler, R.C. and Moradinia, S.F. (2023), "Blockchain in the construction industry between 2016 and 2022: a review, bibliometric, and network analysis", *Smart Cities*, 6(2), pp. 819–845. DOI: <https://doi.org/10.3390/smartcities6020040>
8. Alkhudary, R. and Gardiner, P. (2024), "Utilizing blockchain to enhance project management information systems: insights into project portfolio success, knowledge management and learning capabilities", *International Journal of Managing Projects in Business*, 17(4/5), pp. 731–754. DOI: <https://doi.org/10.1108/IJMPB-01-2024-0021>
9. Perera, R., Wickramarachchi, R. and Rajapakse, C. (2023), "Applications of blockchain technology in project management – a systematic literature review", in *2023 2nd International Conference for Innovation in Technology (INOCON)*, Bangalore, India, pp. 1–7. DOI: <https://doi.org/10.1109/INOCON57975.2023.10101235>
10. Polcumpally, A.T., Pandey, K.K., Kumar, A. and Samadhiya, A. (2024), "Blockchain governance and trust: a multi-sector thematic systematic review and exploration of future research directions", *Heliyon*, 10(12), Article e32975. DOI: <https://doi.org/10.1016/j.heliyon.2024.e32975>
11. Lytvynenko, D., Malyeyeva, O. and Yelizieva, A. (2021), "Blockchain technologies in communication management of infrastructure projects", *Radioelectronic and Computer Systems*, 3, pp. 169–181. DOI: <https://doi.org/10.32620/reks.2021.3.14>
12. Vorobiov, V., Cherechin, O., Romaniv, I., Bereza, V., Cherepaniak, V. and Cherepaniak, A. (2023), "Digitalization of the IT project management process", *Academic Visions*, 26. Available at: <https://www.academy-vision.org/index.php/av/article/view/1206> (accessed 18 April 2026).
13. Topazly, R. (2025), "Research of modern technologies of construction management, digitalization of processes and their adaptation to Ukrainian realities", *Public Administration and Customs Administration*, 1(44), pp. 43–49. DOI: <https://doi.org/10.32782/2310-9653-2025-1.8>
14. Petkun, S. and Pomazanskyi, Yu. (2025), "Regarding priority of project management and blockchain technologies", *Public Administration and Digital Practices*, 1(4), pp. 60–72. DOI: <https://doi.org/10.31673/2786-7412.2025.013691>
15. Bushuiev, S., Bushuieva, N., Lobok, Ye. and Murovanskyi, H. (2025), "Management of innovative projects based on artificial intelligence applications in a turbulent environment", *Innovative Technologies and Scientific Solutions for Industries*, 2(32), pp. 168–176. DOI: <https://doi.org/10.30837/2522-9818.2025.2.168>
16. Kobylkin, D., Zachko, O., Mytsko, R., Zakharchyshyn, S. and Elmas, C. (2025), "Management of critical parameters of logistics and infrastructure projects by means of computer experiment (on the example of the security and defense sector in war conditions)", *Innovative Technologies and Scientific Solutions for Industries*, 3(33), pp. 33–44. DOI: <https://doi.org/10.30837/2522-9818.2025.3.033>
17. Husieva, Y., Chumachenko, I., Dotsenko, N., Kadykova, I. and Skachkov, O. (2026), "Information support for process transformation: tools of the grey method of earned requirements", in *Smart Technologies in Urban Engineering (STUE 2024)*, Lecture Notes in Networks and Systems, Vol. 1659, Cham: Springer, pp. 37–45. DOI: https://doi.org/10.1007/978-3-032-06832-3_4
18. Dotsenko, N., Chumachenko, D., Nekrasov, I., Lutsiv, Y. and Husieva, Y. (2025), "Development of processes for monitoring program implementation in stakeholder-oriented resource management", *Advanced Information Systems*, 9(3), pp. 42–49. DOI: <https://doi.org/10.20998/2522-9052.2025.3.05>
19. Husieva, Yu., Nekrasov, I. and Chumachenko, I. (2025), "Intelligent models for value monitoring in programs and project portfolios", *Municipal Economy of Cities. Series: Information Technology and Engineering*, 4(192), pp. 16–21. DOI: <https://doi.org/10.33042/3083-6727-2025-4-192-16-21>
20. Husieva, Y., Chumachenko, I., Dotsenko, N., Biletskyi, I., Nekrasov, I. and Khudiakov, I. (2025), "Intelligent framework for monitoring sustainable development goals within project portfolios and programs", *CEUR Workshop Proceedings*, 4164, pp. 323–333. Available at: <https://ceur-ws.org/Vol-4164/paper21.pdf> (accessed 18 April 2026).
21. Kyryllova, O., Shakhov, V., Rossomakha, O., Lozytskyy, O., Zakrevskiy, S. and Pitera, V. (2023), "Evaluation of the quality of research projects of universities of transport engineering based on the portfolio management information system", in *2023 IEEE 18th International Conference on Computer Science and Information Technologies (CSIT)*, pp. 1–4. DOI: <https://doi.org/10.1109/CSIT61576.2023.10324064>
22. Kolesnikova, K., Olekh, T., Mukhamediyeva, A., Haidabrus, B. and Kairatuly, Z. (2025), "Multidimensional assessment of sustainable development projects", in *2025 IEEE 13th International Conference on Intelligent Data Acquisition and Advanced Computing Systems: Technology and Applications (IDAACS)*, pp. 1–6. DOI: <https://doi.org/10.1109/IDAACS68557.2025.11322315>

23. Project Management Institute (2017), *The Standard for Portfolio Management*. 4th ed., Newtown Square, PA: Project Management Institute. Available at: <https://www.pmi.org/standards/for-portfolio-management> (accessed 18 April 2026).
24. Hansen, L.K. and Svejvig, P. (2022), "Seven decades of project portfolio management research (1950–2019) and perspectives for the future", *Project Management Journal*, 53(3), pp. 277–294. DOI: <https://doi.org/10.1177/87569728221089537>
25. Tuominen, S. and Martinsuo, M. (2025), "Alternative approaches to innovation project portfolio governance", *Project Management Journal*, 56(1), pp. 5–21. DOI: <https://doi.org/10.1177/87569728241242429>

Received (Надійшла) 30.04.2025

Accepted for publication (Прийнята до друку) 21.05.2026

Publication date (Дата публікації) 27.06.2026

Відомості про авторів / About the Authors

Гусєва Юлія Юріївна – доктор технічних наук, доцент, Харківський національний університет міського господарства імені О.М. Бекетова, професор кафедри управління проєктами в міському господарстві і будівництві, Харків, Україна;

Yuliia Husieva – Doctor of Technical Sciences, Associate Professor, O.M. Beketov National University of Urban Economy in Kharkiv, Professor of the Project Management in Urban Economy and Construction Department, Kharkiv, Ukraine;

e-mail: yulia.guseva@kname.edu.ua

ORCID ID: <http://orcid.org/0000-0001-6992-543X>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=57194413028>

Чумаченко Ігор Володимирович – доктор технічних наук, професор, Харківський національний університет міського господарства імені О.М. Бекетова, завідувач кафедри управління проєктами в міському господарстві і будівництві, Харків, Україна;

Igor Chumachenko – Doctor of Technical Sciences, Professor, O.M. Beketov National University of Urban Economy in Kharkiv, Head of the Project Management in Urban Economy and Construction Department, Kharkiv, Ukraine;

e-mail: igor.chumachenko@kname.edu.ua

ORCID ID: <https://orcid.org/0000-0003-2312-2011>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=57194419994>

Некрасов Іван Борисович – кандидат технічних наук, Центральний науково-дослідний інститут озброєння та військової техніки Збройних Сил України, науковий співробітник, Київ, Україна;

Ivan Nekrasov – Candidate of Technical Sciences, Central Research Institute of Armaments and Military Equipment of Armed Forces of Ukraine, Research Fellow, Kyiv, Ukraine;

e-mail: ib.nekrasov@gmail.com

ORCID ID: <http://orcid.org/0009-0002-6528-7973>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=60013729200&origin=resultslist>

Худяков Ілля Олександрович – доктор філософії, Харківський національний університет міського господарства імені О.М. Бекетова, старший викладач кафедри управління проєктами в міському господарстві і будівництві, Харків, Україна;

Illia Khudiakov – PhD, O.M. Beketov National University of Urban Economy in Kharkiv, Senior Lecturer of the Project Management in Urban Economy and Construction Department, Kharkiv, Ukraine;

e-mail: Illia.Hudyakov@kname.edu.ua

ORCID ID: <https://orcid.org/0000-0002-0049-2979>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=58020241400>

Доценко Наталія Володимирівна – доктор технічних наук, професор, Харківський національний університет міського господарства імені О.М. Бекетова, професор кафедри управління проєктами в міському господарстві і будівництві, Харків, Україна;

Nataliia Dotsenko – Doctor of Technical Sciences, Professor, O.M. Beketov National University of Urban Economy in Kharkiv, Professor of the Project Management in Urban Economy and Construction Department, Kharkiv, Ukraine;

e-mail: Nataliya.Dotsenko@kname.edu.ua

ORCID ID: <http://orcid.org/0000-0003-3570-5900>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=60022696900>

КОНЦЕПТУАЛЬНА МОДЕЛЬ МОНІТОРИНГУ ПОРТФЕЛЬНОГО УПРАВЛІННЯ В ПРОЄКТНИХ ОФІСАХ НА ОСНОВІ БЛОКЧЕЙНУ

Предметом вивчення є процеси забезпечення прозорості, підзвітності та цілісності даних у системах управління портфелями проєктів на основі технологій розподіленого реєстру. **Мета дослідження** – розроблення концептуальної моделі Blockchain Portfolio Governance Model (BPGM) для підвищення прозорості, цілісності та керованості процесів стратегічного портфельного управління. **Завдання:** розробити багаторівневу архітектуру моделі, що поєднує традиційні управлінські цикли з криптографічними механізмами фіксації подій; формалізувати управлінські рішення як об'єкти розподіленого реєстру з використанням асиметричної криптографії; запропонувати комплексний показник оцінки якості управління, що містить як технічну цілісність, так і процедурну дисципліну; здійснити валідацію моделі способом імітаційного моделювання бізнес-процесів. **Методи дослідження:** системний аналіз, методи математичного та імітаційного моделювання в середовищі Bizagi Modeler, алгоритми асиметричного шифрування й гешування для забезпечення цілісності даних у розподілених мережах. **Результати.** У роботі запропоновано й обґрунтовано архітектуру моделі Blockchain Portfolio Governance Model, яка передбачає п'ять рівнів: управління, агрегацію даних, формалізацію рішень, криптографічну цілісність і аудит. Розроблено математичний апарат фіксації подій, де кожне рішення засвідчується за допомогою алгоритму цифрового підпису ECDSA. Упроваджено новий інтегральний показник – індекс відповідності портфельного управління (Portfolio Governance Compliance Index), що дає змогу виявляти "тіньові" управлінські дії за допомогою порівняння кількості ініційованих запитів у зовнішніх системах із кількістю валідованих транзакцій у блокчейні. Проведено серію імітаційних експериментів, які довели, що впровадження алгоритмів консенсусу Proof-of-Authority в корпоративну мережу додає незначні часові затримки (менше 1% від загального циклу), тоді як основний час процесу припадає на експертний аналіз. **Висновки.** Застосування моделі BPGM допомагає трансформувати суб'єктивне управління портфелем у прозорий алгоритмічний процес. Запропоноване рішення забезпечує створення "єдиного джерела істини" для всіх стейкхолдерів, значно спрощує процедури аудиту й підвищує інституційну надійність організації без зниження її операційної ефективності.

Ключові слова: блокчейн; управління портфелями проєктів (PPM); прозорість управління; розподілений реєстр (DLT); імітаційне моделювання бізнес-процесів; індекс відповідності портфельного управління (PGCI); підзвітність даних.

Бібліографічні описи / Bibliographic descriptions

Гусєва Ю. Ю., Чумаченко І. В., Некрасов І. Б., Худяков І. О., Доценко Н. В. Концептуальна модель моніторингу портфельного управління в проєктних офісах на основі блокчейну. *Сучасний стан наукових досліджень та технологій в промисловості*. 2026. № 2 (36). С. 15–28. DOI: <https://doi.org/10.30837/2522-9818.2026.2.015>

Husieva, Y., Chumachenko, I., Nekrasov, I., Khudiakov, I., Dotsenko, N. (2026), "A conceptual model for monitoring portfolio management in project offices based on blockchain", *Innovative Technologies and Scientific Solutions for Industries*, No. 2 (36), P. 15–28. DOI: <https://doi.org/10.30837/2522-9818.2026.2.015>

Lesia Dubchak, Nazar Vivchar, Oleg Savenko, Bohdan Derysh, Oleg Zastavnyy

SOFTWARE TOOL FOR IMPLEMENTING AN INTELLIGENT METHOD FOR ANALYZING WIND TURBINE BLADE DEFECTS WITH LIMITED DATASET

The **subject** of the study is methods and software tools for intelligent diagnostics of wind turbine blade defects based on neuro-fuzzy models capable of operating effectively in conditions of limited, fuzzy, or partially defined input data. The **goal** of the study is to develop a software implementation and conduct an experimental study of a neuro-fuzzy Wang–Mendel network for classifying wind turbine blade defects, and to justify its feasibility as the basis of an intelligent system for monitoring the technical condition of energy facilities. The main **tasks** are to analyze modern methods of defect classification, develop the architecture of a neuro-fuzzy model, create a software package for training and classification, implement an algorithm for automatically forming a fuzzy rule base, optimize the parameters of membership functions, and conduct experimental studies to assess the effectiveness of the proposed approach. The study used **methods** of fuzzy logic, neuro-fuzzy systems, mathematical modeling, data classification, parameter optimization, and software implementation of intelligent systems in Python using modern data processing and numerical computing libraries. The Wang–Mendel algorithm provides automatic formation of fuzzy rules based on training data and adaptation of model parameters to increase classification accuracy. As a **result**, an intelligent software complex for classifying wind turbine blade defects was developed, which provides effective data processing and decision-making based on fuzzy rules. Experimental studies were conducted on a training sample of 160 records and a test sample of 64 records, describing the probabilities of the presence of various types of defects. **Conclusions:** according to the training results, an accuracy of 91% was achieved on the training sample and 94% on the test sample with a low level of mean square error, which confirms the high efficiency, generalization ability, and practical suitability of the proposed model for use in intelligent systems for monitoring the technical condition of wind turbines in conditions of limited data and resources.

Keywords: wind turbines; defects; neuro-fuzzy systems; Wang–Mendel; Python.

Introduction

Over the past 5–10 years, wind power has shown steady and rapid growth worldwide and in Europe. Global wind power capacity is expected to increase by 50% in the coming years [1, 2]. Europe is a leader in the development of this sector, with an average annual growth rate of 15–16% over the past two decades [3, 4]. The largest capacities are concentrated in Germany, Spain, the United Kingdom and France, which together provide more than 60% of all wind installations on the continent [5].

Wind power is expected to remain a key driver of the global energy transition between 2030 and 2050, driven by technological innovation, cost reductions and supportive policies [6, 7]. Mid-term projections also predict a significant reduction in the cost of wind power: by 2050, the cost could fall by 37–49%, reaching \$20–30/MWh for onshore and \$40–60/MWh for offshore wind farms [8, 9].

At the same time, the rapid growth of the industry places new demands on the reliability and efficiency of wind turbine operation. The economic performance of energy companies, the competitiveness of renewable

energy, and the success of the transition to a low-carbon economy directly depend on the performance of the equipment.

Wind turbine failures have a direct and significant impact on both operating costs and performance. Failures increase maintenance requirements, reduce energy production, and shorten the life of the equipment, which together lead to increased costs, reduced efficiency, and consequently increased electricity costs [10, 11]. Minor failures, such as blade surface erosion, are a major cause of unplanned repairs, the cost of which is much higher than the cost of repairs for major structural failures [12].

Thus, the development of wind energy requires increased attention to the condition of turbines and their maintenance methods. Ensuring the stability of equipment operation is becoming an important prerequisite for further growth of the industry, which emphasizes the relevance of research in the field of diagnostics of wind turbine blade defects.

The research aims to develop, implement, and investigate the effectiveness of a Wang–Mendel neuro-fuzzy network for the task of classifying defects in wind turbine blades.

To achieve the stated objective, the following research tasks must be addressed:

1. To analyze modern methods and approaches for diagnostics and classification of wind turbine blade defects, including machine learning, deep learning, and neuro-fuzzy systems, and to identify their advantages and limitations.
2. To develop the model of an intelligent defect classification system based on the Wang–Mendel neuro-fuzzy network, capable of automatically generating a fuzzy rule base from training data.
3. To implement a software system for wind turbine blade defect classification using the Python programming language, incorporating data preprocessing, fuzzification, fuzzy inference, and adaptive parameter optimization.
4. To train the neuro-fuzzy model using experimental data and determine the optimal parameters of membership functions and the structure of the fuzzy rule base.
5. To conduct experimental studies to evaluate the performance of the developed model using training and testing datasets, and to assess its classification accuracy, error rate, and generalization capability.
6. To perform a comparative analysis of the proposed approach with existing defect classification methods and to justify the feasibility of applying the Wang–Mendel neuro-fuzzy network in intelligent monitoring systems for wind turbine condition assessment.

Analysis of last achievements and publications

The development of wind energy is accompanied by increasing requirements for the reliability of wind turbines, in particular for their key structural elements – blades. Traditional diagnostic methods based on visual inspection or classical signal analysis algorithms are often laborious, expensive and limited in accuracy.

In this regard, machine learning is increasingly being used to improve the efficiency of blade defect detection. Current research applies various machine and deep learning methods using visual and sensory data for more accurate identification and classification of damage. Convolutional neural networks, in particular Xception, ResNet-50, AlexNet, VGG-19, as well as specialized architectures, have demonstrated high efficiency in defect classification tasks: the use of the Xception model with transfer learning provided accuracy up to 99.92% on small samples and up to 100% on large ones [13].

The development of real-time object detection architectures has led to the widespread use of advanced YOLO models, such as YOLOv5s, YOLOX, and YOLOv8n. Their combination with attention mechanisms, feature fusion technologies, and efficient loss functions has significantly increased the average accuracy and speed of detection [14]. Signal analysis methods are actively used to detect internal defects, including cracks and delamination: guided waves, wavelet transforms, feature extraction algorithms in combination with k-nearest neighbors' classifiers, support vector methods, and random forest. Such approaches provide classification accuracy of up to 97% [15]. Additionally, adaptive neuro-fuzzy inference systems (ANFIS) and Gaussian processes are used, which reduce computational costs and allow for early damage detection with an accuracy of up to 91% [16].

Key advantages of these methods include high accuracy and speed of data processing, the ability to detect a wide range of damage, and the possibility of integration with unmanned aerial vehicles for remote monitoring [15, 16]. The use of weakly supervised learning and drones for image collection can significantly reduce the need for manual annotation and reduce the risk of human error. Integration with signal processing systems and the use of optimized models make it possible to perform large-scale automated maintenance with minimal resource costs [17].

At the same time, these approaches have certain limitations. Deep neural networks require large volumes of well-annotated data, which complicates their application in practice. Classical algorithms remain vulnerable to the detection of small or hidden defects, although modern improvements partially solve this problem. The computational requirements are also significant, but the use of lightweight models and effective feature selection methods allows for reducing the load [18]. A current research direction is to increase the resistance of models to a decrease in the quality of input data (low resolution, overlapping objects) and to develop more effective methods of merging and supplementing data.

Despite the high performance of modern machine learning methods mentioned earlier, they have a number of limitations, including significant data requirements, the need for careful annotation, and high computational load. In practical wind energy applications, only limited or partial data of small volume are often available, which reduces the effectiveness of classical neural networks

or deep learning methods. In this context, approaches that can work with incomplete or fuzzy input data without losing the accuracy of defect classification are of particular importance.

Given the successful application of fuzzy logic in medicine, robotics and other areas [19–21], its use in wind turbine damage diagnostics is relevant. Neural-fuzzy networks, which combine the capabilities of fuzzy inference systems and artificial neural networks, demonstrate high efficiency in cases where the data are limited or contain uncertainty. Of particular interest is the Wang-Mendel neuro-fuzzy network, which provides flexible formation of inference rules based on input data and allows achieving high accuracy even on small samples [19]. The functioning of such systems is based on fuzzification – the transformation of clear input values into fuzzy sets using membership functions and the reverse process of defuzzification, which allows obtaining one clear value at the output. This provides an interpreted representation of knowledge in the form of a system of fuzzy rules, bringing the decision-making process closer to the logic of human experts. Adding a neural component allows for the implementation of a learning mechanism through which the parameters of the membership functions can adapt and the fuzzy inference rules can change to increase the accuracy of solving the problem.

Modification of the Wang–Mendel neuro-fuzzy network

Fuzzy inference systems and neural networks have similar principles of operation. As inference systems, they require training, the formulation of a certain representation of the knowledge obtained during training, and an algorithm for processing new data using certain rules. The similarity is as follows:

1. Neural networks are trained on input data, while fuzzy systems mostly rely on an expert who sets the inference rules.
2. The "knowledge" of a neural network is stored in the format of weight coefficients, and of a fuzzy system, in the form of rules represented by membership functions.
3. Both systems process input data based on acquired knowledge and a predefined algorithm.

Fuzzy systems and neural networks operate as dynamic systems capable of parallel and distributed data processing. Both approaches are excellent at approximating functions without relying on predefined

mathematical models, instead drawing knowledge directly from sample data. These systems are not only similar in their adaptive nature but also provide unique advantages when applied to real-world scenarios [6].

A hybrid model that combines the concepts of fuzzy logic and artificial neural networks is called a neuro-fuzzy network. Due to the peculiarity of this combination, there is a system that is based on the strengths of both methods and allows you to work effectively with inaccurate, partial or limited information about complex systems. A great advantage of such an architecture is the combination of the interpretability of a fuzzy inference system with the ability to use training methods familiar to "classical" neural networks. Due to this, the accuracy of the hybrid model is better than that of a simple fuzzy inference system, and such a system is also better able to handle classification tasks based on large volumes of data.

One type of neural fuzzy network is the Wang–Mendel neural fuzzy network, proposed by Jerry Mendel and Li-Xin Wang in 1992 [8]. The architecture of such a neural network consists of four layers:

- fuzzification layer;
- aggregation layer;
- linear layer;
- defuzzification layer.

The fuzzification layer, which also acts as an input layer, takes as input a data vector $x_1, x_2, \dots, x_N$ of length N and fuzzified each variable separately to formulate a fuzzy set A calculating the value L of each of the membership functions $\mu_l(x_n)$ ($l = 1, 2, \dots, L$; $n = 1, 2, \dots, N$).

The generalized Gaussian function is used as the membership function:

$$\mu_A(x_n) = \frac{1}{1 + \left(\frac{x_n - c_n}{\sigma_n} \right)^2} \quad (1)$$

To initialize the fuzzification layer, it is necessary to determine the number L of membership functions, and the parameters c_n and σ_n ($n = 1, 2, \dots, N$), which will be refined during the training process to obtain a better refinement of the final prediction of the network. In some cases, the initial values of the parameters c_n and σ_n and the number of membership functions can be either calculated automatically or statically specified before the network training begins. We will consider the case with predefined parameters of the first layer. At the output of the fuzzification layer, L vectors of

dimension N are obtained, where each element is calculated by the formula

$$\mu_{lr}(x_r) = \frac{1}{1 + \left(\frac{x_r - c_{lr}}{\sigma_{lr}} \right)^2} \quad (l = 1, 2, \dots, L; \quad r = 1, 2, \dots, N), \quad (2)$$

where x_r – the r -th element of the input feature vector, c_{lr} – the center of the Gaussian membership function, σ_{lr} – the width (variance) of the function.

The aggregation layer performs the composition of the membership function values L calculated by the previous layer, formulating the inference rules. The calculation is carried out according to the formula:

$$\mu_l(x) = \prod_j \mu_{lj}(x_j). \quad (3)$$

The aggregation layer does not contain parameters that could change during training.

The linear layer is parametric and contains L synaptic weights $w_1, w_2, \dots, w_L$ whose values will adapt during training. The initial value of these weights is chosen randomly from the interval $[-1; 1]$. This layer is used to aggregate inference rules according to the formulas $f_1 = \sum_l w_l \cdot \mu_l(x)$, $f_2 = \sum_l \mu_l(x)$.

The fourth layer is responsible for data defuzzification and obtaining the system output signal according to the formula

$$y(x) = f_1 / f_2. \quad (4)$$

The value $y(x)$ obtained at the end will be the prediction of the Wang–Mendel neural fuzzy network. Thus, the task of the network is that for any pair of data $\langle x, d \rangle$, the prediction of the neural network $y(x)$ corresponds to the expected value d .

Several common algorithms are used to train a Wang–Mendel neural fuzzy network, the most common of which is based on the principle of minimizing the objective function. This method is similar to the method of training conventional neural networks and consists in selecting such values of weights $w_1, w_2, \dots, w_L$ and parameters c_n and σ_n to obtain the most accurate result. For K training samples, the objective function can be specified as

$$E = \frac{1}{2} \sum_{k=1}^K (y(x^k) - d^k), \quad (5)$$

where x^k – input training vector, $y(x^k)$ – prediction of the fuzzy neural network for vector x^k , d^k – expected value of the network output for vector x^k .

Training a Wang–Mendel fuzzy neural network using the objective function minimization algorithm occurs through the alternation of two stages:

1. For the fixed parameters c_n and σ_n the first layer of the network, the values $w_1, w_2, \dots, w_L$ of the synaptic weights of the third layer are calculated.
2. For the fixed values of the synaptic weights of the third layer, the values of the parameters c_n and σ_n the first layer are calculated.

At the first stage of training a neural network, training K samples $\langle x^k, d^k \rangle$ are selected. Based on these samples, a vector $D = \{d^1, d^2, \dots, d^K\}$, ($k = 1, 2, \dots, K$) of expected network responses is formed. Based on the obtained values, a system of K linear equations is formulated:

$$PV \cdot W = D, \quad (6)$$

where W – the desired vector of values $w_1, w_2, \dots, w_N$ of synaptic weights of the third layer, PV – the matrix of coefficients, where each value of the element of each row corresponds to the coefficient calculated for the corresponding value of the input vector of the corresponding sample, which can be expressed as

$$PV_r^k = \prod_j \mu_{rj}(x_j^k) / \sum_{l=1}^L \prod_j \mu_{lj}(x_j^k). \quad (7)$$

In this case, a much larger number of training samples K is selected than the dimension of the input vector N , accordingly, the number of rows of the coefficient matrix PV will be much larger than the number of its rows.

From equation (6) we can find the value of the vector W and represent it as an equation

$$W = PV^+ \cdot D, \quad (8)$$

where PV^+ – pseudo-inverse matrix to the PV matrix. The pseudo-inverse matrix can be searched for by any of the corresponding algorithms, but this does not affect the training of the network. Therefore, by solving the system of linear equations, the vector W of synaptic weights of the third layer will be obtained. After this, the stage of synaptic weight refinement is considered complete.

At the second stage, to calculate the values of the synaptic weights of the third layer of the neural network,

the parameters c_n and σ_n (i.e., the coefficients of the Gaussian membership function) of the first layer are refined. For this, the gradient descent method is used, for which the values and on the training cycle can be expressed as

$$\begin{cases} c_{ij}(k+1) = c_{ij}(k) - \eta_c \cdot \frac{\partial E(k)}{\partial c_{ij}} \\ \sigma_{ij}(k+1) = \sigma_{ij}(k) - \eta_\sigma \cdot \frac{\partial E(k)}{\partial \sigma_{ij}}, \end{cases} \quad (9)$$

where η_c and η_σ are coefficients that determine the speed of parameter adaptation c_n and σ_n , respectively, $\partial E(k)/\partial c_{ij}$ and $\partial E(k)/\partial \sigma_{ij}$ are partial derivatives of network errors.

After refinement c_n and σ_n , the learning process is restarted from the first stage and repeated until either

the desired accuracy is achieved or all parameters (of the first and third layers are stabilized). When implementing this learning algorithm in practice, the first stage (refinement of the weights of the third layer by solving a system of linear equations) is of the greatest importance, which is performed in one step, and the second stage is secondary. The second stage of learning is repeated a significant number of times with each new iteration of learning.

Software implementation

The algorithmic logic of the method implementation in the software system can be formalized in the form of generalized pseudocode, which reflects the complete system workflow from data acquisition to obtaining the classification result, as shown in Figure 1.

```

Input: Training dataset D_tr, Testing dataset D_te
Output: Predicted class labels for D_te

1: Read D_tr, D_te from CSV files
2: Preprocess data (numeric conversion, NaN removal, shuffling)
3: Normalize all input variables to [0,1]
4: Define fuzzy partitions for each input variable
5: Initialize empty rule base RB
6: for each sample (x, y) in D_tr do
7:   for each input variable x_i do
8:     determine fuzzy term A_i with max membership  $\mu(A_i, x_i)$ 
9:   end for
10:  compute rule weight  $w = \prod \mu(A_i, x_i)$ 
11:  form rule R: IF  $x_1$  is  $A_1$  AND ... AND  $x_n$  is  $A_n$  THEN class y
12:  if conflict with existing rule in RB then
13:    keep rule with larger weight
14:  else
15:    add R to RB
16:  end if
17: end for
18: for each sample x in D_te do
19:  compute activation of all rules in RB
20:  aggregate activations per class
21:  assign class with maximum aggregated activation
22: end for
  
```

Fig. 1. Pseudocode of the fuzzy rule processing method implementation

The program code implementing the modified Wang–Mendel neuro-fuzzy network is structured as separate functional blocks responsible for data input, normalization, dataset diagnostics, fuzzy model construction, and implementation of the Wang–Mendel algorithm. In particular, the code includes functions for loading and analyzing training and testing datasets, normalizing numerical features, generating fuzzy partitions, and auxiliary procedures for displaying statistical information about the data, which

allows monitoring the correctness of preprocessing stages (Fig. 2).

From a data flow perspective, the software system forms a closed processing pipeline in which the results of each stage are directly used in the subsequent stage, and the output decisions are transferred to the evaluation subsystem for generating the confusion matrix and quantitative performance metrics. The software architecture has a modular structure and includes subsystems for data input and storage,

preprocessing and normalization, fuzzy model definition, Wang–Mendel rule generation, fuzzy inference mechanism, and result evaluation and visualization. This separation of functional responsibilities ensures

the extensibility of the software system and provides the foundation for its further integration with other intelligent methods or hardware-based data acquisition systems.



Fig. 2. Main modules of the Wang–Mendel network software implementation

Two data sets were used for training: training (160 records) and test (64 records). Each record was a vector of five elements: the first four corresponded to the probabilities of the presence of a certain defect (from 0 to 1), and the fifth element contained the actual class, one of four possible states of the blade. Three types of defects were taken into account: erosion, corrosion

and crack, as well as a class corresponding to the absence of damage. Accordingly, each record clearly belonged to one of the four classes that determined the real state of the wind turbine. Before feeding into the model, the data were normalized. An example of the initial (before normalization) records for each of the classes is given in Table 1.

Table 1. Example of input data sample for training the proposed neuro-fuzzy network

Probability of erosion	Corrosion probability	Probability of a crack	Probability of normal condition
0.889	0.0134	0.0972	0.0005
0.0694	0.4155	0.2806	0.2345
0.1482	0.2308	0.4903	0.1308
0.0009	0.0568	0.055	0.8872
0.889	0.0134	0.0972	0.0005
0.0694	0.4155	0.2806	0.2345

The Wang–Mendel fuzzy neural network was trained using the objective function minimization algorithm by alternately optimizing the parameters of the membership functions of the first layer and the synaptic weights of the third layer.

The implementation of the Wang–Mendel algorithm was carried out using the interpreted object-oriented programming language Python in the PyCharm development environment. The entire architecture of the neural network was created without the use of third-party libraries. To perform individual mathematical operations, analyze data and display them, auxiliary libraries were used: math (simple mathematical operations), time (measurement of program execution time), pandas (data manipulation and analysis), numpy (performance of matrix operations and high-level mathematical

operations), matplotlib (display of data in a graphical representation), tqdm (display of learning progress).

The optimal number of membership functions was determined empirically: increasing their number generally improved accuracy, but after a certain threshold, the efficiency began to decrease, while the training time continued to increase. The best results were obtained with 30 membership functions. The initial centers of the Gaussian functions were evenly distributed in the interval [1, 4], and the widths were fixed at 0.2. Since these parameters were optimized during the training process, their initial choice did not have a critical impact on the final result and was based on expert judgment.

The network was trained for 10 epochs with an early stop mechanism: if the classification accuracy and

the standard deviation remained unchanged for three consecutive epochs, training was stopped. As a result, the model achieved an accuracy of 91% with a standard deviation of 0.0975. The training process was completed after four full epochs according to the early stop criterion. The total training time was 20 minutes 39 seconds.

Experimental studies

The program available at the GitHub repository link [22] implements the Wang–Mendel method, which combines the principles of fuzzy logic and heuristic analysis algorithms to create a fuzzy control system, particularly for data classification and solving optimization problems. A detailed description of its main components is provided below.

The program is implemented in Python and its logic involves sequential data processing, starting with loading the necessary libraries and data sets, followed by their preliminary preparation, which includes cleaning from gaps, normalization and standardization of features to ensure correct operation of the algorithm and increase the stability of the learning process. The central element of the program is the implementation of the Wang–Mendel algorithm, which forms a fuzzy classification model based on training data, automatically generating a database of fuzzy rules that reflect the patterns between the input features and the corresponding classes. To increase the efficiency of the model, a parameter optimization mechanism is provided, which allows you to adapt the membership functions and other system parameters in order to achieve the optimal ratio between classification accuracy and computational costs. The main functions of the software tool provide a full cycle of model operation, including the training procedure based on the training sample and the formation of predictions for new data, which allows using the developed system as a universal classification tool.

The program is organized in an interactive Jupyter Notebook environment, which provides the ability to step-by-step execute the algorithm, monitor intermediate results, and analyze the model's effectiveness in real time. After executing the corresponding code blocks, the user receives numerical indicators of the classification quality, in particular, the accuracy and error values, which allows you to assess the effectiveness of the model and its generalization ability. In addition to numerical results, the program generates graphical representations of the classification results, including error matrices

and scatter plots, which provides a visual representation of the correspondence between the true and predicted classes and allows you to assess the nature of the error distribution. In particular, Fig. 3 shows a scatter plot of the true and predicted classes for the test sample, which demonstrates a high level of correspondence between the classification results and the actual values, which confirms the effectiveness of the proposed approach.

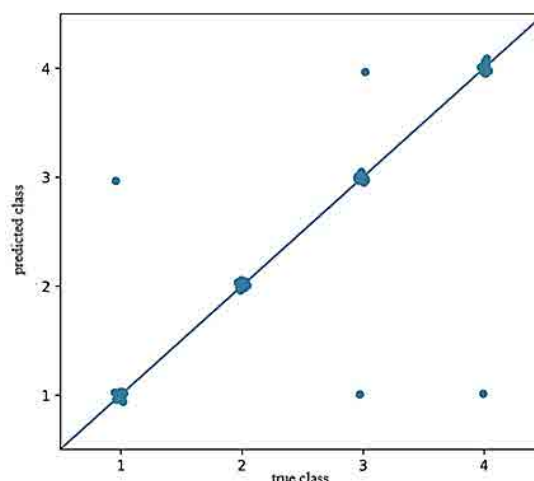


Fig. 3. Scatter plot of true and predicted classes for the test sample

The principle of operation of the software tool is based on the use of fuzzy logic and heuristic learning, which allows you to work effectively in conditions of uncertainty and incomplete data. The fuzzy model forms a system of rules of the "if-then" type, which allows you to perform classification based on the degree of belonging of the input data to the corresponding fuzzy sets, ensuring a smooth transition between classes and increased resistance to noise. The heuristic nature of the algorithm allows you to automatically form the optimal structure of the model based on the analysis of training data without the need for complex iterative training procedures, which reduces computational costs and increases the speed of the system. Thanks to this, the user has the opportunity to adapt the model parameters in accordance with the specifics of a particular task, which ensures the versatility of the software tool and the possibility of its use for different types of data.

Experimental studies were conducted on a test sample of 64 records, each record containing four features corresponding to the estimated probabilities of erosion, corrosion, crack, or no defect, and belonging to one of four classes. The results obtained demonstrate the high efficiency of the developed software tool: the classification accuracy on the test sample reached 94% accuracy at $RMSE = 0.0686$. This indicates a good

generalization ability of the model and its efficiency when working with new data, and also confirms the feasibility of using the Wang–Mendel method as the basis of an intelligent defect classification system in conditions of limited amounts of training data and the need to ensure high accuracy and interpretability of the results.

Comparison with analogues

A comparative analysis of modern approaches to defect classification shows that their effectiveness is largely determined by the balance between accuracy, training data volume requirements, computational resources, and interpretability of the results obtained. As shown in Fig. 4, the maximum accuracy values are demonstrated by deep convolutional neural networks, in particular Xception TL, which provide up to 99.92% of correct classifications. The main advantage of such models is the ability to automatically generate informative features without the need for manual design. At the same time, their practical application is significantly limited by high requirements for the training sample volume, significant computational costs, and complexity of implementation, which complicates their use in embedded or resource-limited monitoring systems.

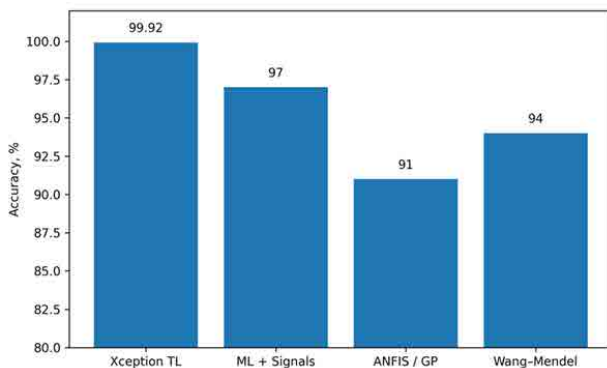


Fig. 4. Comparison of accuracy with analogues

Signal analysis approaches using classical machine learning methods demonstrate somewhat lower but consistently high accuracy rates, reaching 97%. Their characteristic feature is lower computational costs compared to deep neural networks and the possibility of effective use with limited resources. However, the effectiveness of such methods significantly depends on the quality of sensor data and the correctness of feature vector formation, which reduces their versatility and complicates their application in variable or noisy environments.

Methods based on neuro-fuzzy systems, in particular ANFIS and Gaussian processes, demonstrate increased efficiency when working with small samples and provide accuracy at the level of about 91%. Their key advantage is the ability to work under uncertainty and take into account the fuzzy nature of the input data. However, the complexity of parameter tuning, the increase in computational complexity with increasing dimensionality of the problem, and limited scalability may limit their use in real monitoring systems with large data streams.

The proposed approach based on the Wang–Mendel method demonstrates classification accuracy at the level of 94%, which exceeds the indicators of classical neuro-fuzzy systems and is comparable to more complex machine learning methods, while maintaining significantly lower requirements for resources and the amount of training data. The key feature of this method is the formation of an explicit base of fuzzy rules without the need for complex iterative training procedures, which ensures high computational efficiency and stability of results even with a limited training sample. Unlike deep neural networks, the proposed approach does not require large arrays of labeled data and can function effectively in conditions of limited information, which is critically important for monitoring tasks of technical objects, where obtaining large samples is a difficult or expensive process.

An additional advantage of the Wang–Mendel method is its interpretability, since the classification results are formed on the basis of a clearly defined base of fuzzy rules, which allows for the analysis of the decision-making logic and ensures the transparency of the system. This is a fundamental difference from deep neural networks, which function as a "black box" and do not allow for a direct explanation of the process of forming the result.

Thus, the proposed method provides an optimal compromise between classification accuracy, computational efficiency, ability to work with small samples, and interpretability of results. This confirms the feasibility of its use as the basis of an intelligent defect classification system in resource-limited and real-world operating conditions.

Conclusions

This paper analyzes the current state and trends in wind energy development, demonstrating the need to

improve the reliability and efficiency of wind turbines, in particular their blades. It is shown that defects in structural elements significantly affect the productivity, economic performance and duration of equipment operation, and therefore the competitiveness of the entire industry.

A review of existing diagnostic methods revealed that traditional approaches are inferior to modern machine and deep learning algorithms in terms of accuracy and speed. Still, the latter require large volumes of qualitatively annotated data and have high computational costs. In this regard, a promising direction is the use of neuro-fuzzy systems that can operate effectively under data incompleteness or uncertainty.

Within the framework of the study, a neuro-fuzzy network based on the Wang–Mendel method was implemented for the classification of defects in wind turbine blades. The developed model combines the interpretability of fuzzy inference systems with the adaptability of artificial neural networks, which made it possible to achieve high classification accuracy with a relatively small amount of training data. According to the results of the experiments, the training accuracy was 91%, and on the test set – 94% with a standard deviation of 0.0686, which confirms the effectiveness of the proposed architecture for defect diagnosis tasks.

The obtained results indicate the feasibility of using Wang–Mendel neuro-fuzzy networks in wind turbine technical condition monitoring systems, especially in cases of limited or partially fuzzy data. Further research should be directed towards expanding the sample size, optimizing membership functions, and combining neuro-

fuzzy systems with computer vision and deep learning methods to increase the model's robustness to noise and input variability.

Conflict of interest

The authors declare that they have no conflicts of interest, including financial, personal, authorship, or any other conflicts that could influence the research or the results published in this article.

Funding

This research was supported by the Ministry of Education and Science of Ukraine and funded by the European Union's external aid instrument as part of Ukraine's commitments under the EU Framework Program for Research and Innovation "Horizon 2020". The work was carried out as part of the research project "Intelligent Defect Detection System for Green Energy Facilities Using UAVs", state registration number 0124U004665 (2023–2026).

Data availability

Data will be provided upon reasonable request

Use of artificial intelligence

The authors confirm that they did not use artificial intelligence technology in the creation of this paper.

References

1. Arshad, M., O'Kelly, B. (2019), "Global status of wind power generation: theory, practice, and challenges", *International Journal of Green Energy*, Vol. 16, pp. 1073–1090. DOI: <https://doi.org/10.1080/15435075.2019.1597369>
2. Bukovsky, M., Barthelmie, R., Leung, L., Pryor, S. and Sakaguchi, K. (2020), "Climate change impacts on wind power generation". *Nature Reviews Earth & Environment*, Vol. 1, pp. 627–643. DOI: <https://doi.org/10.1038/s43017-020-0101-7>
3. Hatziaegyriou, N., Zervos, A. (2001), "Wind power development in Europe", *Proceedings of the IEEE*, Vol. 89, pp. 1765–1782. DOI: <https://doi.org/10.1109/5.975906>
4. Lukas, M., Friedrich, K. (2017), "Wind energy statistics in Europe: onshore and offshore", *Wind Energy Management*, pp. 339–347. DOI: https://doi.org/10.1007/978-3-319-45659-1_36
5. Dubis, B., Dunn, J., Beldycka-Bórawska, A., Jankowski, K., Bórawski, P. (2020), "Development of wind energy market in the European Union", *Renewable Energy*, 161, pp. 691–700. DOI: <https://doi.org/10.1016/j.renene.2020.07.081>
6. Viktor, P., Ali, B., Alzoubi, H., Jaszczur, M., Salman, H., Al-Musawi, T., Hassan, Q., Al-Jiboory, A., Sameen, A., Algburi, S. (2024), "The renewable energy role in the global energy transition", *Renewable Energy Focus*, 100545 p. DOI: <https://doi.org/10.1016/j.ref.2024.100545>
7. Alphan, H., Bilgili, M. (2022), "Global growth in offshore wind turbine technology", *Clean Technologies and Environmental Policy*, 24, pp. 2215–2227. DOI: <https://doi.org/10.1007/s10098-022-02314-0>
8. Lantz, E., Stehly, T., Shields, M., Cooperman, A., Kitzing, L., Beiter, P., Kikuchi, Y., Wiser, R., Berkhout, V., Telsnig, T. (2021), "Wind power costs driven by innovation and experience with further reductions on the horizon", *Wiley Interdisciplinary Reviews: Energy and Environment*, Vol. 10, Issue 5. DOI: <https://doi.org/10.1002/wene.398>

9. Baker, E., Beiter, P., Rand, J., Gilman, P., Seel, J., Wiser, R., Lantz, E. (2021), "Expert elicitation survey predicts 37% to 49% declines in wind energy costs by 2050", *Nature Energy*, 6, pp. 555–565. DOI: <https://doi.org/10.1038/s41560-021-00810-z>
10. Fan, Y., Peng, H., Zhang, H., Li, S., Shangguan, L. (2023), "Analysis of wind turbine equipment failure and intelligent operation and maintenance research", *Sustainability*, 15(10), 8333p. DOI: <https://doi.org/10.3390/su15108333>
11. Zhou, S., Wu, H., Du, Y., Peng, Y., Jing, X., Kwok, N. (2020), "Damage detection techniques for wind turbine blades: a review", *Mechanical Systems and Signal Processing*, Vol. 141, 106445 p. DOI: <https://doi.org/10.1016/j.ymssp.2019.106445>
12. Thomsen, K., Mishnaevsky, L. (2020), "Costs of repair of wind turbine blades: influence of technological aspects", *Wind Energy*, pp. 2247–2255. DOI: <https://doi.org/10.1002/we.2552>
13. Altice, B., Nazario, E., Davis, M., Shekaramiz, M., Moon, T., Masoum, M. (2024), "Anomaly detection on small wind turbine blades using deep learning algorithms", *Energies*, 17(5), 982 p. DOI: <https://doi.org/10.3390/en17050982>
14. Liu, Y., Zheng, Y., Shao, Z., Wei, T., Cui, T., Xu, R. (2024), "Defect detection of the surface of wind turbine blades combining attention mechanism", *Advanced Engineering Informatics*, 59, 102292 p. DOI: <https://doi.org/10.1016/j.aei.2023.102292>
15. Ogaili, A., Jaber, A., Hamzah, M. (2023), "A methodological approach for detecting multiple faults in wind turbine blades based on vibration signals and machine learning", *Curved and Layered Structures*, Vol. 10. DOI: <https://doi.org/10.1515/cls-2022-0214>
16. Dubchak, L., Sachenko, A., Bodyanskiy, Y., Wolff, C., Vasylykiv, N., Brukhanskyi, R., Kochan, V. (2024), "Adaptive neuro-fuzzy system for detection of wind turbine blade defects", *Energies*, 17(24), 6456 p. DOI: <https://doi.org/10.3390/en17246456>
17. Dubchak, L., Rusyn, B., Wolff, C., Ciszewski, T., Sachenko, A., Bodyanskiy, Y. (2026), "Hypersector-based method for real-time classification of wind turbine blade defects", *Energies*, 19, 442 p. DOI: <https://doi.org/10.3390/en19020442>
18. He, Y., Niu, X., Hao, C., Li, Y., Kang, L., Wang, Y. (2024), "An adaptive detection approach for multi-scale defects on wind turbine blade surface", *Mechanical Systems and Signal Processing*, 200, 111592 p. DOI: <https://doi.org/10.1016/j.ymssp.2024.111592>
19. Casal-Guisande, M., Cerqueiro-Pequeño, J., Bouza-Rodríguez, JB, Comesaña-Campos, A. (2023), "Integration of the Wang & Mendel algorithm into the application of fuzzy expert systems to intelligent clinical decision support systems", *Mathematics*, 11(1), 2469 p. DOI: <https://doi.org/10.3390/math11112469>
20. Chen, L., Su, W., Wu, M., Pedrycz, W., Hirota, K. (2020), "A fuzzy deep neural network with sparse autoencoder for emotional intention understanding in human–robot interaction", *IEEE Transactions on Fuzzy Systems*, 28(7), pp. 1252–1264. DOI: <https://doi.org/10.1109/TFUZZ.2020.2966167>
21. Talpur, N., Abdulkadir, SJ, Alhussian, H., Hasan, MH, Aziz, N., Bamhdi, A. (2022), "A comprehensive review of deep neuro-fuzzy system architectures and their optimization methods", *Neural Computing and Applications*, 34(1), pp. 1837–1875. DOI: <https://doi.org/10.1007/s00521-021-06807-9>
22. Vivchar, N. "Wang-Mendel", available at: <https://github.com/NazarVivchar/wang-mendel/blob/59e6d3e9aec4533feced179cda53c7711a4dcdcf/main.ipynb>

Received (Надійшла) 24.02.2026

Accepted for publication (Прийнята до друку) 21.05.2026

Publication date (Дата публікації) 27.06.2026

Відомості про авторів / About the Authors

Дубчак Леся Орестівна – кандидат технічних наук, доцент, Західноукраїнський національний університет, завідувачка кафедри комп'ютерної інженерії, Тернопіль, Україна;

Lesia Dubchak – Candidate of Technical Sciences, Associate Professor, West Ukrainian National University, Head of the Computer Engineering Department, Ternopil, Ukraine;

e-mail: dlo@wunu.edu.ua

ORCID ID: <https://orcid.org/0000-0003-3743-2432>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=56008186500>

Вівчар Назар Тарасович – Західноукраїнський національний університет, здобувач вищої освіти ступеня доктора філософії кафедри комп'ютерної інженерії, Тернопіль, Україна;

Nazar Vivchar – West Ukrainian National University, Postgraduate Student of the Computer Engineering Department, Ternopil, Ukraine;

e-mail: nvivcharn@gmail.com

ORCID ID: <https://orcid.org/0009-0003-8196-4978>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=59953891200>

Савенко Олег Станіславович – доктор технічних наук, професор, Хмельницький національний університет, професор кафедри комп'ютерної інженерії та інформаційних систем, Хмельницький, Україна;

Oleg Savenko – Doctor of Technical Sciences, Professor, Khmelnytskyi National University, Professor of the Computer Engineering and Information Systems Department, Khmelnytskyi, Ukraine;

e-mail: savenko_oleg_st@ukr.net

ORCID ID: <https://orcid.org/0000-0002-4104-745X>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=54421023400>

Дериш Богдан Богданович – Західноукраїнський національний університет, викладач кафедри комп'ютерної інженерії, Тернопіль, Україна;

Bohdan Derysh – West Ukrainian National University, Lecturer of the Computer Engineering Department, Ternopil, Ukraine;

e-mail: dbb@wunu.edu.ua

ORCID ID: <https://orcid.org/0000-0002-7215-9032>

Заставний Олег Михайлович – кандидат технічних наук, Західноукраїнський національний університет, старший викладач кафедри спеціалізованих комп'ютерних систем, Тернопіль, Україна;

Oleg Zastavnyy – Candidate of Technical Sciences, West Ukrainian National University, Senior Lecturer of the Specialized Computer Systems Department, Ternopil, Ukraine;

e-mail: o.zastavnyi@wunu.edu.ua

ORCID ID: <https://orcid.org/0000-0001-8630-8791>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=8366871500>

ПРОГРАМНИЙ ЗАСІБ РЕАЛІЗАЦІЇ ІНТЕЛЕКТУАЛЬНОГО МЕТОДУ АНАЛІЗУ ДЕФЕКТІВ ЛОПАТЕЙ ВІТРОВИХ ТУРБІН З ОБМЕЖЕНОЮ ВИБІРКОЮ ДАНИХ

Предметом дослідження є методи й програмні засоби інтелектуальної діагностики дефектів лопатей вітрових турбін на основі нейронечітких моделей, здатних ефективно функціонувати в умовах обмежених, нечітких або частково визначених вхідних даних. **Мета** – розроблення програмної реалізації та експериментальне дослідження нейронечіткої мережі Ванга–Менделя для класифікації дефектів лопатей вітрових турбін, а також обґрунтування доцільності її використання як основи інтелектуальної системи моніторингу технічного стану енергетичних об'єктів. **Завдання дослідження:** аналіз сучасних методів класифікації дефектів, розроблення архітектури нейронечіткої моделі, створення програмного комплексу для навчання й класифікації, реалізація алгоритму автоматичного формування бази нечітких правил, оптимізація параметрів функцій належності та проведення експериментальних досліджень для оцінювання ефективності запропонованого підходу. У роботі використано **методи** нечіткої логіки, нейронечітких систем, математичного моделювання, класифікації даних, оптимізації параметрів і програмної реалізації інтелектуальних систем мовою Python із застосуванням сучасних бібліотек оброблення даних і чисельних обчислень. Алгоритм Ванга–Менделя забезпечує автоматичне формування нечітких правил на основі навчальних даних і адаптацію параметрів моделі з метою підвищення точності класифікації. **Результати.** Розроблено інтелектуальний програмний комплекс для класифікації дефектів лопатей вітрових турбін, який забезпечує ефективне оброблення даних і підтримку прийняття рішень на основі нечітких правил. Експериментальні дослідження проведено на навчальній вибірці обсягом 160 записів і тестовій вибірці обсягом 64 записи, що описують імовірності наявності різних типів дефектів. **Висновки:** за результатами навчання досягнуто точності 91% на навчальній вибірці та 94% на тестовій вибірці за низького рівня середньоквадратичної похибки, що підтверджує високу ефективність, здатність до узагальнення та практичну придатність запропонованої моделі для використання в інтелектуальних системах моніторингу технічного стану вітрових турбін в умовах обмежених даних і ресурсів.

Ключові слова: вітрові турбіни; дефекти; нейронечіткі системи; алгоритм Ванга–Менделя; Python.

Бібліографічні описи / Bibliographic descriptions

Дубчак Л. О., Вівчар Н. Т., Савенко О. С., Дериш Б. Б., Заставний О. М. Програмний засіб реалізації інтелектуального методу аналізу дефектів лопатей вітрових турбін з обмеженою вибіркою даних. *Сучасний стан наукових досліджень та технологій в промисловості*. 2026. № 2 (36). С. 29–39. DOI: <https://doi.org/10.30837/2522-9818.2026.2.029>

Dubchak, L., Vivchar, N., Savenko, O., Derysh, B., Zastavnyy, O. (2026), "Software tool for implementing an intelligent method for analyzing wind turbine blade defects with limited dataset", *Innovative Technologies and Scientific Solutions for Industries*, No. 2 (36), P. 29–39. DOI: <https://doi.org/10.30837/2522-9818.2026.2.029>

Oleh Zolotukhin, Yevgeniy Bodyanskiy, Valentin Filatov, Maryna Kudryavtseva, Oleksandr Vasylets

STACKING HYBRID SYSTEM OF COMPUTATIONAL INTELLIGENCE BASED ON KERNEL ACTIVATION-MEMBERSHIP FUNCTIONS AND ITS ONLINE LEARNING IN PATTERN RECOGNITION TASK

Subject of research. The subject of the research is a stacking hybrid computational intelligence system based on a kernel activation-membership function, which combines the advantages of neural networks and fuzzy models for solving pattern recognition tasks under conditions of uncertainty, noise, and limited dimensionality of data sets. **The goal of the work.** The goal of the work is to develop a stacking hybrid system with a kernel activation-membership function and online learning to improve the accuracy, stability, and speed of pattern recognition in non-stationary environments in real-time. **Task.** Develop the architecture of a stacking hybrid system; propose a kernel activation-membership function and an algorithm for its parameterization; construct an online learning algorithm with step-by-step parameter updates; investigate the properties of convergence and computational complexity; conduct an experimental comparison with basic machine learning models and classical neuro-fuzzy approaches. **Methods.** Machine learning methods, ensemble and stacking approaches, kernel methods, adaptive optimization in online mode, statistical analysis of classification quality, modelling and computational experiments on synthetic and real datasets applied to pattern recognition tasks. **Results.** The proposed model provides higher classification accuracy and better generalization ability compared to individual baseline models, demonstrates resistance to noise and changes in data distributions, as well as rapid adaptation in streaming data conditions thanks to online learning. It is shown to reduce the error and ensure stable operation of the system in non-stationary environments. **Conclusions.** The developed stacking hybrid system with a kernel activation-membership function is an effective tool for real-time pattern recognition tasks. The combination of the stacking approach and online learning increases the accuracy, stability and adaptability of the system, which makes the proposed approach promising for practical application in intelligent information systems and cyber-physical environments.

Keywords: computational intelligence; data stream mining; kernel activation-membership function; machine learning; neuro-neo-fuzzy system; stacking hybrid system.

Introduction

Computational intelligence (CI) systems have become popular today for solving various information processing problems [1–3] and, above all, classification problems – pattern recognition, mainly images, which is explained from a formal point of view by their universal approximating capabilities. Here, the most widespread are deep convolutional neural networks (DCNN) [4–6], which allow restoring nonlinear separating multidimensional hypersurfaces of any complex shape between an arbitrary number of possible classes. However, some problems may arise in the case of overlapping classes when observations can simultaneously belong to several classes with different levels of membership. In this case, hybrid systems of computational intelligence (HSCI) and, above all, neuro-fuzzy systems are the most effective.

Analysis of modern publications

At the same time, deep systems are quite cumbersome with a fixed multilayer architecture, containing a very large number of tuned parameters and requiring too large volumes of training samples for their learning, which are not always available when solving many practical tasks. In addition, the stationarity of the characteristics of the objects forming this sample is assumed a priori. One way or another, traditional DCNNs are quite complex and slow systems.

Essentially faster in terms of the learning process are systems (primarily neural networks) with kernel activation functions [7], the most common of which are Radial Basis Function Networks (RBFNs) [8–11] which are also characterized by universal approximating properties, but are able to adjust their synaptic weights online. This is explained by the fact that the output signal of these networks linearly depends on the adjusted

parameters that allows the use of speed-optimized learning algorithms, including the Gaussian Newtonian ones.

In [12], a convolutional RBFN was proposed, which demonstrated high speed and quality of learning compared to DCNNs based on elementary Rosenblatt perceptrons.

At the same time, traditional RBFNs suffer from so-called "curse of dimensionality" effect when the number of kernel activation functions (R -neurons) grows rapidly with the increase in the dimensionality of the input signal-image. Finding the required number of R -neurons in the network could be done using the methods of evolutionary computational intelligence systems [13–15], however the gradual addition of kernel functions to the network [16] slows down the learning process.

In our opinion, some of the noted problems can be overcome by using the stacking approach [17], when the system as a whole is formed from separate subsystems, each of which is either a neural network or a neuro-fuzzy or neo-fuzzy system [18–19] trained independently of the other subsystems using different algorithms. This approach will increase the speed of training and does not require too large volumes of training samples.

The goal of the work. The goal of the work is to improve the efficiency of a stacking hybrid system with a kernel activation-membership function and online learning to improve the accuracy, stability, and speed of pattern recognition in non-stationary environments in real-time.

Methods and Materials

Architecture of a stacking hybrid computational intelligence system based on kernel activation-membership functions

The feedforward architecture of the proposed SHSCI is shown in Fig. 1 and contains two stacks and an output layer formed by a softmax activation function. The zero receptive layer receives a vector input signal-image $x(k) = (x_1(k), \dots, x_i(k), \dots, x_n(k)) \in R^n$,

where $k = 1, 2, \dots, N, N+1, \dots$ either the number of a specific observation in the training sample, or the index of the discrete current time in the case when the data for processing are received online.

In this case, it is assumed a priori that all data belong to (or are pre-coded into) some limited interval $x_{i \min} \leq x_i(k) \leq x_{i \max}, \forall i = 1, 2, \dots, n$ first stack is

essentially a simplified RBFN and is formed by $h > n$ so-called R -neurons, i.e. multidimensional kernel activation functions $h > n$. In this layer the dimensionality of the input space R^n increases which is formed in space R^h by the R -neuron system. This approach, on which all radial-basis neural networks are based according to the T.M. Cover theorem [20–21], allows transforming the initial linearly inseparable problem of pattern recognition in space R^n into linearly-separable one in a space of increased dimension R^h .

Multidimensional Gaussians are usually used as activation functions in RBFNs due, first of all, to the simplicity of their derivatives, which, in turn, simplifies the learning process

$$\varphi_l(x(k)) = u_l(k) = \varphi_l(\|x(k) - c_l\|_2^2, \sigma_l^2) = e^{-\frac{\|x(k) - c_l\|_2^2}{2\sigma_l^2}}.$$

Here the vector c_l specifies the coordinates of the center of the kernel function φ_l , and a scalar parameter σ^2 is its "width". Although theoretically any kernel (most often bell-shaped) function can be used as an activation function, it is easy to construct an RBFN with matrix inputs using Gaussians [22, 23]. This means that the input of such a network can be supplied with not only a vector signal but also a matrix signal, i.e. an image without its prior transformation using the convolution operation. In this case, the radial basis activation function can be written in the form

$$\varphi_l(x(k)) = u_l(k) = e^{\frac{(Sp(x(k) - c_l)(x(k) - c_l))^T}{2\sigma^2}},$$

where $Sp(\cdot)$ is the matrix trace symbol, is the square of the Frobenius matrix norm, which is consistent with the Euclidean vector norm and specifies the distance between the matrix input signal $x(k)$ and the matrix center of the Gaussian c_l

$$(Sp(x(k) - c_l)(x(k) - c_l))^T.$$

At the output of this stack, a vector signal of increased dimension is formed which is fed to the second stack. This vector signal is fed to a generalized neo-fuzzy neuron (GNFN) [24], which, in turn, is a multidimensional modification of the standard neo-fuzzy neuron [18, 19, 25].

$$\varphi(k) = u(k) = \varphi_1(x(k)) \dots, \varphi_l(x(k)) \dots, \varphi_h(x(k)) \dots = (u_1(k) \dots, u_l(k) \dots, u_h(k)) \in R^h.$$

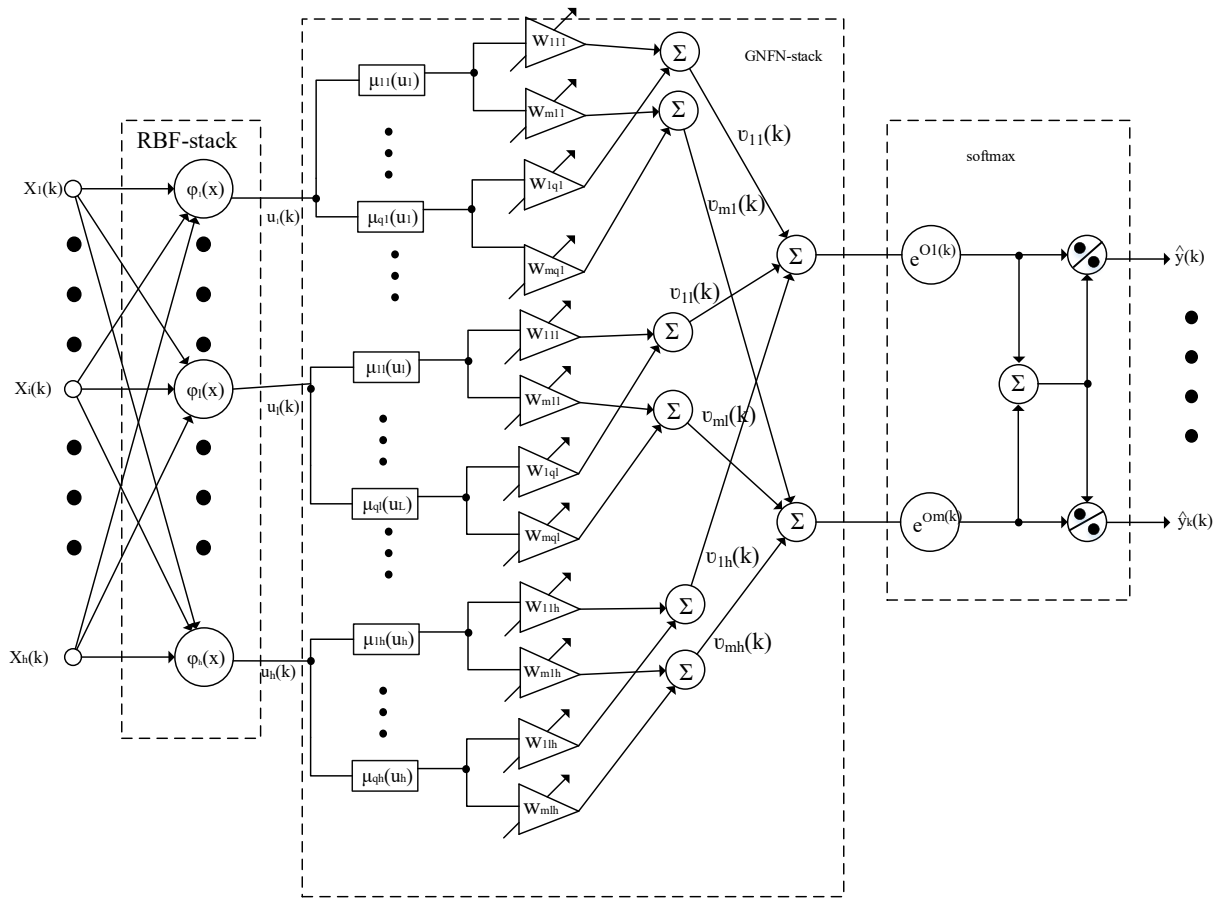


Fig. 1. Architecture of stacking hybrid system of computational intelligence based on kernel activation-membership functions

The GNFN in SHSCI contains qh kernel one-dimensional membership functions $\mu_{pl}(u_l)$, $p = 1, 2, \dots, q$; $l = 1, 2, \dots, h$; m tuned synaptic weights w_{jpl} , $j = 1, 2, \dots, m$ and $mh + m$ elementary adders, at the outputs of which signals $o_1(k), \dots, o_j(k), \dots, o_m(k)$ are formed, where m – the number of classes in the learning set

These signals are fed to the softmax activation functions at the output of the system, which calculate the level of fuzzy membership of each classified observation to a specific class $j = 1, 2, \dots, m$ in the form

$$\hat{y}_j(k) = \frac{e^{o_j(k)}}{\sum_{j=1}^m e^{o_j(k)}}.$$

In principle, instead of GNFN, a system of m ordinary neo-fuzzy neurons could be used, but in this case the number of membership functions would increase by m times, which would make the system as a whole more cumbersome.

As membership functions in NFN and GNFN, triangular functions that satisfy the Ruspini unity partitioning conditions are usually used. If these conditions are not met, then it is necessary to perform a defuzzification operation [26], which complicates the calculation process. Using B -splines instead of triangular structures improves the approximation qualities of the system, but also complicates the numerical implementation. Therefore, it was proposed [27] to use parabolic membership functions based on Epanechnikov kernels [28] instead of triangles, which provides better results compared to piecewise linear approximation when using triangles. In this case, the defuzzification operation is not required, since this procedure is implemented by the softmax activation function at the output of the system.

Membership functions of SHSCI are shown in Fig. 2.

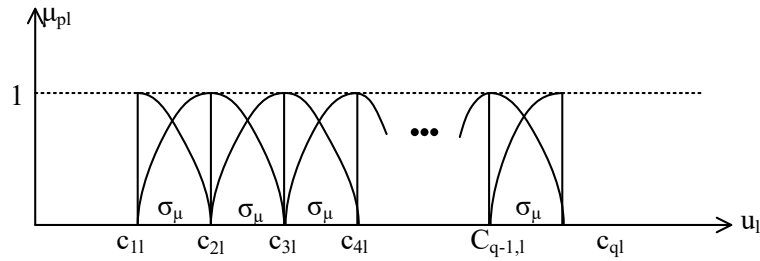


Fig. 2. Membership functions of SHSCI

The general form of the SHSCI membership functions is as follows:

$$\mu_{pl}(u_l) = \left[1 - \frac{(u_l - c_{pl})^2}{\sigma_\mu^2} \right]_+$$

where $[\cdot]_+ = \max\{0\}$, which is also shown in Fig. 2.

The use of such membership functions leads to the fact that at each training cycle of the system not mqh synaptic weights are adjusted, but only $2mh$ weights, which naturally simplifies the training of the system as a whole.

Each of the membership functions $\mu_{pl}(u_l)$ is associated with m (the number of system outputs) synaptic weights w_{jpl} , which are to be adjusted during

the training process. These synaptic weights are associated with two layers of elementary summation blocks. The first layer contains mh adders and the second contains m adders. These adders calculate the values

$$v_{jl}(k) = \sum_{p=1}^q w_{jpl} \mu_{pl}(u_l(k)),$$

and

$$o_j(k) = \sum_{l=1}^h v_{jl}(k) = \sum_{l=1}^h \sum_{p=1}^q w_{jpl} \mu_{pl}(u_l(k)).$$

The signals from the output of the GNFN stack are fed to the SHSCI output layer, formed by the softmax activation function, resulting in the formation of a vector of outputs

$$\hat{y}(k) = (y_1(k), \dots, y_j(k), \dots, y_m(k))^T = \left(\frac{e^{o_1(k)}}{\sum_{j=1}^m e^{o_j(k)}}, \dots, \frac{e^{o_j(k)}}{\sum_{j=1}^m e^{o_j(k)}}, \dots, \frac{e^{o_m(k)}}{\sum_{j=1}^m e^{o_j(k)}} \right)^T,$$

each component of which specifies the level of fuzzy membership of each observation $x(k)$ to each of the classes $j = 1, 2, \dots, m$.

SHSCI training

Training of the proposed system can be conditionally divided into three relatively independent tasks: tuning of synaptic weights w_{jpl} , $j = 1, 2, \dots, m$; $p = 1, 2, \dots, q$; $l = 1, 2, \dots, h$ based on the concept of controlled learning with teacher; tuning of centers of kernel activation functions c_l , $l = 1, 2, \dots, h$; based on self-learning; and tuning of centers of membership

functions c_{pl} , $p = 1, 2, \dots, q$; $l = 1, 2, \dots, h$ also based on self-learning without a teacher.

The main task here is to tune the synaptic weights, which is related to optimizing the cross-entropy learning criterion.

$$E = - \sum_{k=1}^N \sum_{j=1}^m y_j^*(k) \ln y_j(k) = - \sum_{k=1}^N \sum_{j=1}^m y_j^*(k) \ln \frac{e^{o_j(k)}}{\sum_{j=1}^m e^{o_j(k)}},$$

where the components of the external reference signal that can take only two values: "1" or "0".

$$y_j^*(k), j = 1, 2, \dots, m$$

Considering the $(q_h \times 1)$ vector of membership function values

$$\mu(k) = (\mu_{11}(u_1(k)), \dots, \mu_{1l}(u_1(k)), \mu_{12}(u_2(k)), \dots, \mu_{pl}(u_l(k)), \dots, \mu_{qh}(u_h(k)))^T$$

and the vector of synaptic weights associated with the j -th output of SHSCI, we can write the signal value at the j -th output at the k -th tuning cycle in the form

$$w_j = (w_{j11}, \dots, w_{jq1}, w_{j12}, \dots, w_{jpl}, \dots, w_{jqh})^T,$$

$$o_j(k) = w_j^T (k-1) \mu(k),$$

$$\hat{y}_j(k) = \frac{e^{o_j(k)}}{\sum_{j=1}^m e^{o_j(k)}} = \frac{e^{w_j^T (k-1) \mu(k)}}{\sum_{j=1}^m e^{w_j^T (k-1) \mu(k)}}.$$

Optimization of the cross-entropy learning criterion using a gradient procedure leads to an algorithm for tuning synaptic weights of the form:

$$w_j(k) = w_j(k-1) + \eta_j(k)(y_j^*(k) - y_j(k))\mu(k),$$

where η is the learning rate parameter $\eta_j(k)$.

$$w_j(k) = w_j(k-1) + \frac{y_j^*(k) - y_j(k)}{\|\mu(k)\|^2} \mu(k) = w_j(k-1) + (y_j^*(k) - y_j(k))\mu^{T+}(k)$$

where $(\cdot)^+$ is the pseudoinversion symbol.

This algorithm can be given additional filtering (following) properties if it is modified in the following way [27]:

$$\begin{cases} w_j(k) = w_j(k-1) + \beta^{-1}(k)(y_j^*(k) - y_j(k))\mu(k), \\ \beta(k) = \alpha\beta(k-1) + \|\mu(k)\|^2, \end{cases}$$

where $0 \leq \alpha \leq 1$ is the forgetting (smoothing) factor. It is easy to see that $\alpha=0$ provides optimal speed, and $\alpha=1$ turns the algorithm into a stochastic approximation procedure.

The installation and adjustment of the centers of the kernel activation functions of the RBFN stack can be ensured by the combined use of the principle of "Neurons of data points" [31–32] and self-learning according to T. Kohonen [33–34]. The center of the first activation function c_1 is installed at a point with coordinates that coincide with the components of the first observation of the training sample $x(1)$, i.e. $c_1 = x(1)$.

Next, a certain threshold of indistinguishability of two neighboring centers is introduced and after the observation $x(2)$ arrives (it does not matter whether it is a vector or a matrix), the condition is checked. If this condition is true, then the first center is adjusted so that $\|x(2) - c_1\| < r$.

$$c_1(2) = \frac{c_1 + x(2)}{2}.$$

If the inequality is not satisfied, then this center is adjusted according to the self-learning rule "Winner Takes More": $r < \|x(2) - c_1(1)\| \leq 2r$

$$c_1(2) = c_1(1) + \eta_c(2)\psi_1(2)(x(2) - c_1(1))$$

with neighborhood function

$$\psi_1(x) = e^{-\|x(2) - c_1(1)\|^2 / 2r}.$$

In turn, this algorithm can be optimized for the rate of convergence by appropriate choice of the step. So, if we choose $\eta_j(k)$

$$\eta_j(k) = \|\mu(k)\|^{-2},$$

we arrive at the optimal adaptive algorithm in terms of speed by S. Kaczmarz, B. Widrow, M. Hoff [29, 30]:

If the inequality $\|x(2) - c_1(1)\| > 2r$ is not satisfied, then a second RBF activation function is formed

$$\varphi_2(x) = e^{-\|x - c_2\|^2 / 2\delta^2} = e^{-\|x - x(2)\|^2 / 2\delta^2}.$$

The following iterations regarding the formation of activation functions proceed in a similar way. Let us suppose that the system has already received k observations and formed $l < h$ activation functions, the last of which has the center $c_l(k)$. Let us suppose also that an observation of the image $x(k+1)$ has been received. Then, among all the already formed R -neurons, there is a winning neuron whose distance $\|x(k+1) - c_l^*(k)\|$ is minimal among all l formed R -neurons.

The correction of the centers c_l is subsequently carried out as follows:

$$\text{If } \|x(k+1) - c_l^*(k)\| < r,$$

$$\text{then } c_l(k+1) = \frac{c_l^*(k) + x(k+1)}{2},$$

$$\text{if } r < \|x(k+1) - c_l^*(k)\| \leq 2r,$$

then

$$\begin{cases} c_l(k+1) = c_l^*(k) + \eta_c(k+1)\psi_l(k+1)(x(k+1) - c_l^*(k)) \\ \psi_l(x) = e^{-\|x - c_l^*(k)\|^2 / 2r}, \end{cases}$$

if $\|x(k+1) - c_l^*(k)\| > 2r$, the next radial basis function is formed (a new R -neuron is added with center $c_{l+1}(k+1) = x(k+1)$).

This process continues until the maximum possible number h of kernel activation functions is formed in the first stack. It is easy to see that this approach is based on the ideas of evolving systems [13–15] and self-learning according to T. Kohonen [32].

As for the membership function $\mu_{pl}(u_l)$ in the CNFN stack, the adjustment of their centers c_{pl} is completely similar to the adjustment of the centers of the radial basis functions c_l . The only difference

is the choice of the indistinguishable thresholds p_μ .

Fig. 3 shows the appearance of the membership functions after k -cycles of their adjustment.

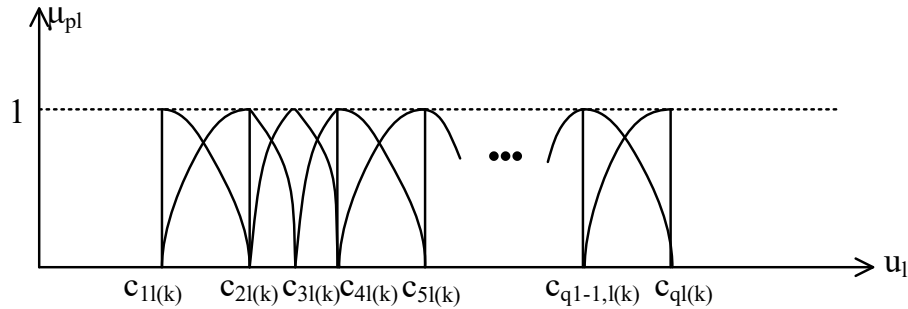


Fig. 3. Membership functions of CHSCI after k -cycles of their tuning

It is easy to see that these functions become asymmetric, and in analytical form they take the form

$$\begin{cases} \mu_{plL}(u_l) = \left[1 - \frac{(u_l - c_{pl}(k))^2}{(c_{pl-1}(k) - c_{pl}(k))^2} \right]_+, \\ \mu_{plR}(u_l) = \left[1 - \frac{(u_l - c_{pl}(k))^2}{(c_{pl+1}(k) - c_{pl}(k))^2} \right]_+, \end{cases}$$

where the indices L , R mean "left" and "right" relative to the center $c_{pl}(k)$.

It should also be noted that this method of self-learning membership function centers is computationally much simpler than setting up these same centers based on learning with a teacher [27], and these advantages are manifested in multidimensional systems, when the number of these membership functions can be quite large.

Results of the computational experiment

The objective of the computational experiment is to evaluate the performance of the proposed stacking hybrid system of computational intelligence in a pattern recognition task with online learning, focusing on classification accuracy, learning speed, and adaptability to data streams using short training samples.

The system was tested on benchmark datasets commonly used in pattern recognition:

– Fashion (Modified National Institute of Standards and Technology database): 28×28 grayscale test images of Zalando products (clothing, footwear, accessories that are presented on the Zalando platform), with 60000 training and 10000 test samples.

– CIFAR-10: 32×32 color images in 10 classes, with 50000 training and 10000 test samples.

Data samples were streamed to the system one-by-one or in small batches of 10–50 samples simulating real-time input.

The experimental settings are as follows:

1. Input format: matrix input was used without preliminary convolution, enabling direct use of image matrices in the RBF kernel layer.
2. System configuration:
 - maximum number of R -neurons in the first stack is $h = 100$;
 - number of membership functions per input channel is $q = 5$;
 - learning rate is $\eta = 0.01$;
 - forgetting factor is $\alpha = 0.1$;
 - indistinguishability thresholds for RBF and membership functions are $p_c = 0.1$, $p_\mu = 0.15$;
3. Evaluation metrics:
 - classification accuracy (overall and per class);
 - number of epochs or time to convergence;
 - evolution of the number of R -neurons and membership functions during online learning;
 - memory capacity (total parameters adjusted during learning).

The procedure of the computational experiment is as follows.

1. Initialization:
 - first R -neuron center c_1 was initialized using the first training sample;
 - membership function centers were initialized uniformly over the input range.

2. Online training:

- samples were fed to SHSCI sequentially;
- RBF centers and membership function centers were updated using T. Kohonen-style self-learning;
- synaptic weights w_{jpl} were updated using the S. Kaczmarz, B. Widrow, M. Hoff algorithm with pseudo-inverse and forgetting factor for filtering.

3. Testing:

- After each group of 500 samples, accuracy was evaluated on a fixed test set.
- The evolution of the network structure (number of R -neurons, active membership functions) was recorded.

The results of the computational experiment using Fashion MNIST and CIFAR-10 data samples are given in Table 1.

Table 1. The results of the computational experiment

Metric	Fashion MNIST	CIFAR-10
Accuracy (Final)	96.8%	84.3%
Accuracy after first 1000 samples	91.5%	74.0%
Number of R -neurons formed	87	93
Total membership functions active	340	365
Time to convergence	~4.8s	~6.1s
Memory capacity (parameters)	~16K	~18K

Additional results:

1. Training dynamics: the system quickly adapted to data, with most R -neurons and membership functions formed in the first 20% of data. Accuracy steadily improved with each batch.

2. Comparison with CNN:

- comparable accuracy to shallow CNNs but with faster training and significantly lower memory usage.
- better adaptation to streaming data and superior performance with small sample sizes.

Visualization of the computational experiment is presented below:

- evolution of membership functions (Fig. 4).
- online formation of R -neuron centers visualized with PCA (Principal Component Analysis) projection of input data space (Fig. 5).
- accuracy vs. number of processed samples plotted to compare learning speed with standard CNN (Fig. 6).

This figure illustrates how the membership functions in the SHSCI system become asymmetric over multiple online learning cycles. Initially symmetric, their shapes adapt based on data distribution, improving approximation accuracy without requiring defuzzification.

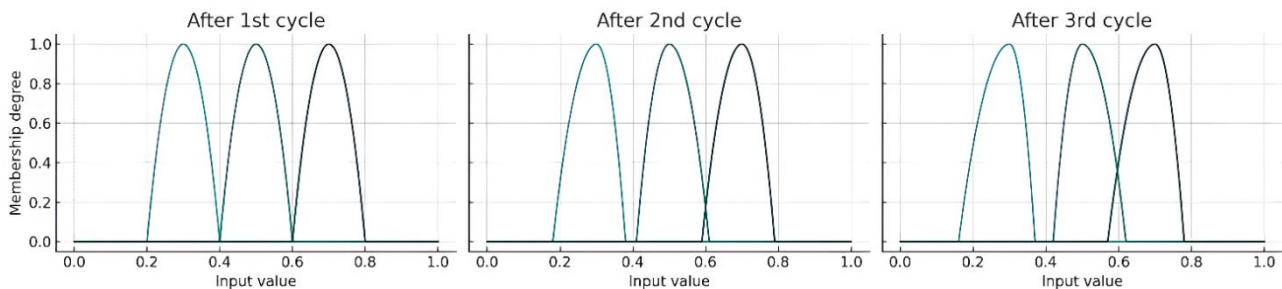


Fig. 4. Evolution of membership functions after k -cycles of tuning

This figure shows a PCA projection of the input data from a digit recognition dataset, for example, Fashion MNIST, with red cruciform markers representing the R -neuron centers dynamically formed during the online learning process of SHSCI. These centers act as kernels for the radial basis functions in the first stack, helping to separate overlapping class regions in the higher-dimensional space.

This plot shows that SHSCI achieves faster convergence in accuracy with fewer training samples compared to a standard CNN. This highlights SHSCI's efficiency in online learning scenarios and with limited data, making it well-suited for real-time tasks and stream-based pattern recognition tasks.

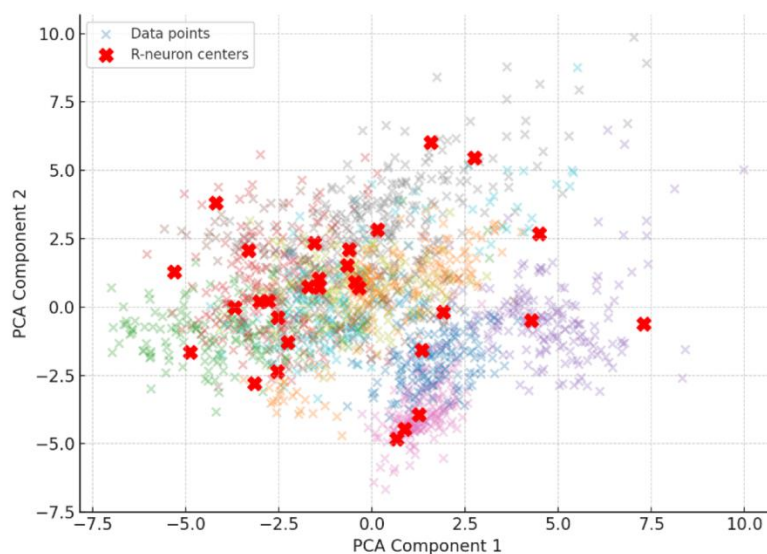


Fig. 5. PCA projection of data and R -neuron centers

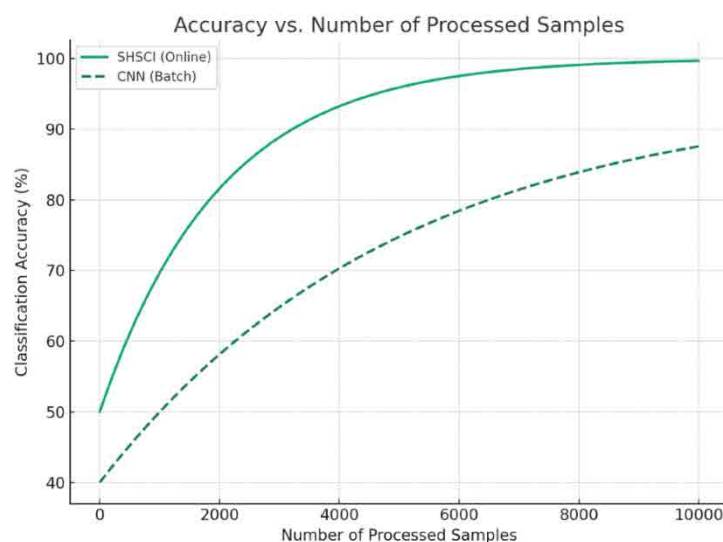


Fig. 6. Accuracy vs. number of processed samples plotted to compare learning speed with standard CNN

Conclusion

A stacking hybrid system of computational intelligence, the stacks of which are formed by a modified radial-basis-function neural network, which provides an increase in the dimensionality of the input space of signals-images and a generalized neo-fuzzy-neuron, which provides high approximating properties when restoring separating hypersurfaces between classes is proposed. The SHSCI is configured based on combined learning and self-learning, at the same time, not only synaptic weights are configured, but also the centers of kernel activation-membership functions. The process of configuring all parameters is quite simple in numerical implementation and occurs in the online mode as information is received for processing. The proposed

SHSCI is intended for solving problems of classification and recognition of images (pictures), at the same time, these images can be presented to the system input in both traditional vector and matrix forms. The advantages of the proposed system are manifested under conditions of short training samples, when traditional deep convolutional networks become ineffective, or when solving data stream mining problems, when training must be performed in real time.

In future work, the proposed stacking hybrid computational intelligence system can also be applied to the important practical problem of dynamic resource management in data networks as a pattern recognition task. In this context, pattern recognition is the transformation of an input signal (image/feature vector) into a class/label or output vector [35–36].

Conflict of interest

The authors declare that they have no conflict of interest, in particular financial, personal, authorial, or any other nature, that could influence the research, as well as the results published in this article.

Horizon 2020 MSCA RISE programme under grant agreement No. 10100828.

Data availability

The manuscript has no linked data.

Funding

This paper is part of the DIOR project that has received funding from the European Union's

Use of artificial intelligence

The authors confirm that they did not use artificial intelligence technologies to write this work.

References

1. Mumford, C.L. and Jain, L.C. (2009), *Computational Intelligence: Collaboration, Fusion and Emergence*, Springer, Berlin Heidelberg, 732 p. DOI: <https://doi.org/10.1007/978-3-642-01799-5>
2. Kruse, R., Borgelt, C., Braune, C., Mostaghim, S. and Steinbrecher, M. (2016), *Computational Intelligence: A Methodological Introduction*, Springer, London, 564 p. DOI: <https://doi.org/10.1007/978-1-4471-7296-3>
3. Kacprzyk, J. and Pedrycz, W. (eds.) (2015), *Springer Handbook of Computational Intelligence*, Springer, Berlin Heidelberg, 1634 p. DOI: <https://doi.org/10.1007/978-3-662-43505-2>
4. Goodfellow, I., Bengio, Y. and Courville, A. (2016), *Deep Learning*, MIT Press, Cambridge, MA, 800 p. ISBN 978-0-262-03561-3.
5. Schmidhuber, J. (2015), "Deep Learning in Neural Networks: An Overview", *Neural Networks*, Vol. 61, pp. 85–117. DOI: <https://doi.org/10.1016/j.neunet.2014.09.003>
6. Kung, S.-Y. (2014), *Kernel Methods and Machine Learning*, Cambridge University Press, Cambridge, 591 p. DOI: <https://doi.org/10.1017/CBO9781139176224>
7. Karamichailidou, D., Gerolymatos, G., Patrinos, P., Sarimveis, H. and Alexandridis, A. (2024), "Radial Basis Function Neural Network Training Using Variable Projection and Fuzzy Means", *Neural Computing and Applications*, Vol. 36, No. 33, pp. 21137–21151. DOI: <https://doi.org/10.1007/s00521-024-10274-3>
8. Ismayilova, A. and Ismayilov, V.E. (2024), "On the Universal Approximation Property of Radial Basis Function Neural Networks", *Annals of Mathematics and Artificial Intelligence*, Vol. 92, No. 3, pp. 174–185. DOI: <https://doi.org/10.1007/s10472-023-09901-x>
9. Marrero, I. (2025), "Relaxed Conditions for Universal Approximation by Radial Basis Function Neural Networks of Hankel Translates", *AIMS Mathematics*, Vol. 10, No. 5, pp. 10852–10865. DOI: <https://doi.org/10.3934/math.2025493>
10. Jensen, V., Bianchi, F.M. and Anfinsen, S.N. (2024), "Ensemble Conformalized Quantile Regression for Probabilistic Time Series Forecasting", *IEEE Transactions on Neural Networks and Learning Systems*, Vol. 35, No. 7, pp. 9014–9025. DOI: <https://doi.org/10.1109/TNNLS.2022.3217694>
11. Zolotukhin, O.V., Kudryavtseva, M.S., Bodyanskiy, Y.V., Filatov, V.O., Antilikatorov, A.V. and Kalinin, D.V. (2026), "Physical-Informed Neural Network in Signal Processing and Network Traffic Communications", *System Research and Information Technologies*, No. 1, pp. 155–169. DOI: <https://doi.org/10.20535/SRIT.2308-8893.2026.1.11>
12. Angelov, P., Filev, D. and Kasabov, N. (eds.) (2010), *Evolving Intelligent Systems: Methodology and Applications*, Wiley/IEEE Press, Hoboken, NJ, 444 p. DOI: <https://doi.org/10.1002/9780470569962>
13. Angelov, P., Zhou, X. and Lughofer, E. (2008), "Evolving Fuzzy Classifiers Using Different Model Architectures", *Fuzzy Sets and Systems*, Vol. 159, No. 23, pp. 3160–3182. DOI: <https://doi.org/10.1016/j.fss.2008.06.019>
14. Baruah, R.D., Angelov, P. and Baruah, D. (2020), "Evolving Intelligent Systems", In: *Wiley Encyclopedia of Electrical and Electronics Engineering*, pp. 1–17. DOI: <https://doi.org/10.1002/047134608X.W8405>
15. Filatov, V., Zolotukhin, O. and Kudryavtseva, M. (2025), "Intellectual Data Analysis in Relational Information and Analytical Systems", *Innovative Technologies and Scientific Solutions for Industries*, Vol. 4, No. 34, pp. 101–111. DOI: <https://doi.org/10.30837/2522-9818.2025.4.101>
16. Jiang, W., Chen, Z., Xiang, Y., Shao, D., Ma, L. and Zhang, J. (2019), "SSEM: A Novel Self-Adaptive Stacking Ensemble Model for Classification", *IEEE Access*, Vol. 7, pp. 120337–120349. DOI: <https://doi.org/10.1109/ACCESS.2019.2933262>
17. Uchino, E. and Yamakawa, T. (1994), "Neo-Fuzzy-Neuron Based New Approach to System Modeling, with Application to Actual System", In: *Proceedings of the Sixth International Conference on Tools with Artificial Intelligence (TAI'94)*, pp. 701–706. DOI: <https://doi.org/10.1109/TAI.1994.346442>

18. Uchino, E. and Yamakawa, T. (1997), "Soft Computing Based Signal Prediction, Restoration and Filtering", In: *Intelligent Hybrid Systems: Fuzzy Logic, Neural Networks and Genetic Algorithms*, Vol. 413, pp. 331–351. DOI: https://doi.org/10.1007/978-1-4615-6191-0_14
19. Brereton, R.G. and Lloyd, G.R. (2010), "Support Vector Machines for Classification and Regression", *Analyst*, Vol. 135, pp. 230–267. DOI: <https://doi.org/10.1039/B918972F>
20. Aggarwal, C.C. (2018), *Neural Networks and Deep Learning: A Textbook*, Springer, Cham, 520 p. DOI: <https://doi.org/10.1007/978-3-319-94463-0>
21. Filatov, V., Semenets, V. and Zolotukhin, O. (2019), "Synthesis of Semantic Model of Subject Area at Integration of Relational Databases", In: *Proceedings of the International Conference on Advanced Optoelectronics and Lasers (CAOL)*, pp. 598–601. DOI: <https://doi.org/10.1109/CAOL46282.2019.9019532>
22. Bodyanskiy, Y., Zolotukhin, O., Yerokhin, A., Kudryavtseva, M. and Yerokhin, M. (2025), "Fast Stacking Neuro-Neo-Fuzzy System for Inverse Modeling in Online Mode", *International Journal of Computing*, Vol. 24, No. 4, pp. 661–667. DOI: <https://doi.org/10.47839/ijc.24.4.4330>
23. Lemos, A., Caminhas, W.M. and Gomide, F. (2014), "A Fast Learning Algorithm for Evolving Neo-Fuzzy Neuron", *Applied Soft Computing*, Vol. 14, Part B, pp. 194–209. DOI: <https://doi.org/10.1016/j.asoc.2013.03.022>
24. Zurita, D., Delgado, M., Carino, J.A., Ortega, J.A. and Clerc, G. (2016), "Industrial Time Series Modelling by Means of the Neo-Fuzzy Neuron", *IEEE Access*, Vol. 4, pp. 6151–6160. DOI: <https://doi.org/10.1109/ACCESS.2016.2611649>
25. Chang, W.J., Chang, C.-H. and Ku, C.C. (2010), "Fuzzy Controller Design for Takagi–Sugeno Fuzzy Models with Multiplicative Noises via Relaxed Non-Quadratic Stability Analysis", *Proceedings of the Institution of Mechanical Engineers, Part I: Journal of Systems and Control Engineering*, Vol. 224, No. 8, pp. 918–931. DOI: <https://doi.org/10.1243/09596518JSCE1035>
26. Mokodompit, M., Nasib, S.K., Djakaria, I., Yahya, N.I. and Hasan, I.K. (2025), "Implementation of Fuzzy Time Series Markov Chain Method Using Kernel Smoothing in Forecasting the Stock Price of PT. Elnusa Tbk.", *Indonesian Journal of Computational and Applied Mathematics*, Vol. 1, No. 1, pp. 18–28. DOI: <https://doi.org/10.64182/indocam.v1i1.9>
27. Chen, Y.-C., Genovese, C.R. and Wasserman, L. (2016), "A Comprehensive Approach to Mode Clustering", *Electronic Journal of Statistics*, Vol. 10, No. 1, pp. 210–241. <https://doi.org/10.1214/16-EJS1115>
28. Needell, D. and Tropp, J.A. (2014), "Paved with Good Intentions: Analysis of a Randomized Block Kaczmarz Method", *Linear Algebra and its Applications*, Vol. 441, pp. 199–221. DOI: <https://doi.org/10.1016/j.laa.2013.01.022>
29. Sayed, A.H. (2014), "Adaptive Networks", *Proceedings of the IEEE*, Vol. 102, No. 4, pp. 460–497. DOI: <https://doi.org/10.1109/JPROC.2014.2306253>
30. Buhmann, M.D. (2003), *Radial Basis Functions: Theory and Implementations*, Cambridge University Press, Cambridge, 272 p. DOI: <https://doi.org/10.1017/CBO9780511543241>
31. Kohonen, T. (1995), *Self-Organizing Maps*, Springer, Berlin Heidelberg, 362 p. DOI: <https://doi.org/10.1007/978-3-642-97610-0>
32. Zolotukhin, O., Filatov, V., Yerokhin, A., Kudryavtseva, M. and Semenets, V. (2021), "An Approach to the Selection of Behavior Patterns of Autonomous Intelligent Mobile Systems", In: *Proceedings of the IEEE 8th International Conference on Problems of Infocommunications, Science and Technology (PIC S&T)*, pp. 349–352. DOI: <https://doi.org/10.1109/PICST54195.2021.9772110>
33. Zolotukhin, O., Filatov, V., Yerokhin, A., Kudryavtseva, M. and Semenets, V. (2021), "The Methods for the Prediction of Climate Control Indicators in the Internet of Things Systems", *CEUR Workshop Proceedings*, Vol. 3013, pp. 391–400. DOI: <https://doi.org/10.5281/zenodo.14526027>
34. Dashenkov, D. and Smelyakov, K. (2025), "Extending the ImageNET Dataset for Multimodal Text and Image Learning", *Innovative Technologies and Scientific Solutions for Industries*, Vol. 1, No. 31, pp. 20–31. DOI: <https://doi.org/10.30837/2522-9818.2025.1.020>
35. Danylenko, S. and Smelyakov, K. (2025), "Content-Based Image Retrieval Method in a Multidimensional Data Model at Big Data Scale", *Innovative Technologies and Scientific Solutions for Industries*, Vol. 4, No. 34, pp. 18–31. DOI: <https://doi.org/10.30837/2522-9818.2025.4.018>
36. Amirian, M. and Schwenker, F. (2020), "Radial Basis Function Networks for Convolutional Neural Networks to Learn Similarity Distance Metric and Improve Interpretability", *IEEE Access*, Vol. 8, pp. 123087–123097. DOI: <https://doi.org/10.1109/ACCESS.2020.3007337>

Received (Надійшла) 14.02.2026

Accepted for publication (Прийнята до друку) 21.05.2026

Publication date (Дата публікації) 27.06.2026

Відомості про авторів / About the Authors

Золотухін Олег Вікторович – кандидат технічних наук, доцент, Харківський національний університет радіоелектроніки, декан факультету комп'ютерних наук, Харків, Україна;

Oleh Zolotukhin – Candidate of Technical Sciences, Associate Professor, Kharkiv National University of Radio Electronics, Dean of the Computer Science Faculty, Kharkiv, Ukraine;

e-mail: oleg.zolotukhin@nure.ua

ORCID ID: <https://orcid.org/0000-0002-0152-7600>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?origin=resultslist&authorId=57207774022>

Бодянський Євгеній Володимирович – доктор технічних наук, професор, Харківський національний університет радіоелектроніки, професор кафедри штучного інтелекту, Харків, Україна;

Yevgeniy Bodyansky – Doctor of Technical Science, Professor, Kharkiv National University of Radio Electronics, Professor of the Artificial Intelligence Department, Kharkiv, Ukraine;

e-mail: yevgeniy.bodyanskiy@nure.ua

ORCID ID: <https://orcid.org/0000-0001-5418-2143>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=13105377000>

Філатов Валентин Олександрович – доктор технічних наук, професор, Харківський національний університет радіоелектроніки, професор кафедри штучного інтелекту, Харків, Україна;

Valentin Filatov – Doctor of Technical Science, Professor, Kharkiv National University of Radio Electronics, Professor of the Artificial Intelligence Department, Kharkiv, Ukraine;

e-mail: valentin.filatov@nure.ua

ORCID ID: <https://orcid.org/0000-0002-3718-2077>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=56911938100>

Кудрявцева Марина Сергіївна – кандидат технічних наук, доцент, Харківський національний університет радіоелектроніки, професор кафедри штучного інтелекту, Харків, Україна;

Maryna Kudryavtseva – Candidate of Technical Sciences, Associate Professor, Kharkiv National University of Radio Electronics, Professor of the Artificial Intelligence Department, Kharkiv, Ukraine;

e-mail: maryna.kudryavtseva@nure.ua

ORCID ID: <https://orcid.org/0000-0003-0524-5528>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?origin=resultslist&authorId=57207765829>

Василюк Олександр Олексійович – Харківський національний університет радіоелектроніки, здобувач вищої освіти ступеня доктора філософії кафедри штучного інтелекту, Харків, Україна;

Oleksandr Vasylets – Kharkiv National University of Radio Electronics, Postgraduate Student of the Artificial Intelligence Department, Kharkiv, Ukraine;

e-mail: oleksandr.vasylets@nure.ua

ORCID ID: <https://orcid.org/0009-0009-1261-1689>

СТЕКОВА ГІБРИДНА СИСТЕМА ОБЧИСЛЮВАЛЬНОГО ІНТЕЛЕКТУ НА ОСНОВІ ЯДЕРНОЇ АКТИВАЦІЙНОЇ ФУНКЦІЇ ТА ЇЇ ОНЛАЙН-НАВЧАННЯ В ЗАДАЧІ РОЗПІЗНАВАННЯ ОБРАЗІВ

Предмет дослідження. Предметом дослідження є стекова гібридна система обчислювального інтелекту на основі ядерної активаційної функції, що поєднує переваги нейронних мереж і нечітких моделей для розв'язання задач розпізнавання образів в умовах невизначеності, зашумленості та обмеженої розмірності наборів даних. **Мета дослідження.** Мета роботи полягає у розробці гібридної стекової системи з ядерною функцією активації та онлайн-навчанням для забезпечення кращої точності, стабільності й швидкості розпізнавання шаблонів у нестационарних середовищах у режимі реального часу. **Завдання.** Розробити архітектуру стекової гібридної системи; запропонувати ядерну активаційну функцію та алгоритм її параметризації; побудувати алгоритм онлайн-навчання з покроковим оновленням параметрів; дослідити властивості

збіжності та обчислювальної складності; провести експериментальне порівняння з базовими моделями машинного навчання та класичними нейро-нечіткими підходами. **Методи.** Застосовано методи машинного навчання, ансамблеві та стекові підходи, ядерні методи, адаптивну оптимізацію в онлайн-режимі, статистичний аналіз якості класифікації, моделювання та обчислювальні експерименти на синтетичних і реальних наборах даних для задач розпізнавання образів. **Результати.** Запропонована модель забезпечує вищу точність класифікації та кращу узагальнювальну здатність порівняно з окремими базовими моделями, демонструє стійкість до шумів і змін розподілів даних, а також швидку адаптацію в умовах потокових даних завдяки онлайн-навчанню. Показано зменшення похибки та стабільну роботу системи в нестационарних середовищах. **Висновки.** Розроблена стекова гібридна система з ядерною активаційною функцією є ефективним інструментом для задач розпізнавання образів у режимі реального часу. Поєднання стекового підходу та онлайн-навчання підвищує точність, стійкість і адаптивність системи, що робить запропонований підхід перспективним для практичного застосування в інтелектуальних інформаційних системах і кіберфізичних середовищах.

Ключові слова: інтелектуальний аналіз потокових даних; машинне навчання; нейро-нео-фаззі система; обчислювальний інтелект; стекова гібридна система; ядерна активаційна функція.

Бібліографічні описи / Bibliographic descriptions

Золотухін О. В., Бодянський Є. В., Філатов В. О., Кудрявцева М. С., Василюк О. О. Стекова гібридна система обчислювального інтелекту на основі ядерної активаційної функції та її онлайн-навчання в задачі розпізнавання образів. *Сучасний стан наукових досліджень та технологій в промисловості*. 2026. № 2 (36). С. 40–51. DOI: <https://doi.org/10.30837/2522-9818.2026.2.040>

Zolotukhin, O., Bodyanskiy, Y., Filatov, V., Kudryavtseva, M., Vasylets, O. (2026), "Stacking Hybrid System of Computational Intelligence Based on Kernel Activation-Membership Functions and its Online Learning in Pattern Recognition Task", *Innovative Technologies and Scientific Solutions for Industries*, No. 2 (36), P. 40–51. DOI: <https://doi.org/10.30837/2522-9818.2026.2.040>

Heorhii Kuchuk, Mykyta Matvieiev, Nataliia Kosenko, Dmytro Lysytsia, Andrii Levchenko

EVALUATING THE PERFORMANCE OF A WEBAR APPLICATION'S SCENE LOADING ON A MOBILE DEVICE

Subject of the study: the patterns of CPU resource allocation and changes in browser engine timing metrics during the execution of WebAR scene initialization stages. **Object of study:** the process of initializing and loading 3D content in a client-side WebAR application on a mobile device. **The aim of the study** is to develop an approach for comprehensive performance evaluation and to identify patterns in the distribution of computational load during the initialization of WebAR applications using deep browser profiling tools. **Objectives.** Develop a mathematical model of WebAR application operation during the scene initialization phase and decompose the loading process into key stages with an assessment of their time costs. Conduct empirical profiling, identify peak CPU loads and blocking tasks, and formulate well-founded directions for further performance optimization. **Methods.** Performance was investigated through empirical performance testing of the WebAR application during its most resource-intensive stage – the initialization and loading of a 3D scene. Data collection regarding CPU thread load within the context of a web page was performed using the Chrome DevTools developer tools. A Google Pixel 8 mobile device served as the hardware environment for the experiments, and the statistical reliability of the results was ensured by a 10-fold iteration of each test scenario. **Research results.** The developed mathematical model formalizes the process of loading a WebAR application scene and allows for the assessment of CPU load at each stage. Based on the analysis of the time profile obtained using the improved mathematical model and Chrome DevTools, four loading stages were identified: initialization of an empty scene, loading a 3D object from the server, model parsing, and model rendering. A nonlinear nature of the CPU load was established. The parsing stage proved to be the most resource-intensive, during which peak CPU load and the longest blocking task were recorded, caused by the synchronous deserialization of binary data. **Conclusions.** An approach is proposed to improve the performance of loading a WebAR application scene on a mobile device by balancing the CPU load. The study confirmed that the critical phase of initialization is the parsing of the 3D model. Given the simulation results and empirical studies conducted, a promising direction is the implementation of incremental loading and processing of 3D content.

Keywords: web; augmented reality; WebAR application; performance evaluation; parsing; rendering; Angular.

1. Introduction

Digital technologies have become an integral part of modern society, transforming the ways we communicate, learn, work, and consume information [1]. One of the most dynamic and promising technologies is augmented reality (AR), which combines virtual digital content with the physical environment in real time [2]. The integration of virtual objects into real space provides a new level of interactivity and immersion, significantly expanding the possibilities for user interaction with information compared to traditional multimedia tools [3].

WebAR technology has seen particular development; it implements AR functionality without the need to install separate applications, providing access to augmented reality directly through a web browser [4]. This approach lowers barriers to entry for users and simplifies the implementation of AR solutions for businesses [5]. Today, WebAR is actively used in e-commerce, medicine, education, and the entertainment industry, contributing to increased user engagement

and the effectiveness of digital services [6]. The growing number of scientific studies and practical implementations in this field indicates the emergence of a trend in which web resources with integrated AR capabilities may become the standard for user interaction with digital content [7].

WebAR has a number of significant advantages, but at the same time is characterized by increased computational load on the web application and the client device [8]. The fundamental contradiction lies in the fact that basic web technologies were created primarily for working with 2D graphics. In contrast, WebAR requires continuous execution of resource-intensive operations: 3D graphics processing, rendering, and real-time spatial tracking [9]. Implementing these processes through standard browser mechanisms leads to reduced performance, manifested in a drop in frame rate, increased latency, blocking of the main execution thread, and overall inefficient use of hardware resources [10]. According to research, such technical and interface barriers account for 64% of all issues preventing users

from performing their intended actions. As a result, system responsiveness and the smoothness of interaction with AR content deteriorate, negatively impacting the user experience. The implementation of optimization solutions can significantly reduce these shortcomings; however, developing effective optimization strategies is a complex and multifaceted task [11]. A key prerequisite for solving this problem is an objective and comprehensive performance evaluation that takes into account the characteristics of the browser engine, the limitations of mobile devices, and the specifics of how 3D graphics function in a web environment [12].

Website performance evaluation is based on the analysis of a set of factors that determine its operational efficiency, which are commonly summarized under the concept of web metrics [13]. These include page load time, interface response speed, content display stability, and other indicators of user experience quality. One of the key characteristics is page load time, which encompasses the processes of fetching, processing, and rendering HTML, JavaScript, and CSS resources, as well as graphical elements. Although external factors – such as network latency and connection bandwidth – influence the overall loading time, internal computational processes are of decisive importance for WebAR applications [14]. These include the duration of JavaScript code parsing and execution, 3D scene initialization, shader compilation, loading and processing of 3D models, as well as the efficiency of CPU resource utilization. The most resource-intensive stage is scene initialization, during which the browser simultaneously processes significant amounts of geometric and texture data, configures the graphics pipeline, and synchronizes the video stream from the camera. It is at this stage that there is an increased risk of the main thread blocking and the interface "freezing", which critically affects the user's perception of the service [15]. In this regard, systematic and objective performance evaluation becomes particularly relevant as a necessary condition for improving the efficiency and stability of WebAR applications.

2 Analysis of Recent Studies and Publications

A significant portion of current research focuses primarily on the network and time-based performance characteristics of web resources. For example, in [16], a framework based on machine learning methods is proposed for predicting the performance of web pages. Within this model, quality is assessed based on a set

of 16 metrics, dominated by temporal and network indicators: server response time, load duration, page size, and the number of HTTP requests. Despite taking into account indicators such as *Start Render Time* and *Time to Interactive*, the proposed methodology does not provide for in-depth profiling of computational resource consumption, particularly CPU load during the active phase of the application [15]. A similar approach can be seen in the work of *Ghattas and Sartawi* (2020), where the evaluation of web resource quality is based primarily on loading speed and layout stability metrics, without accounting for the computational cost of executing client-side scripts and complex interactive scenarios [6]. *Hussein et al.* (2023) investigated the impact of WebAR on the academic performance of elementary school students. During the technical validation of the AR.js and MyWebAR platforms, the authors used the GTmetrix toolkit, which is primarily focused on evaluating page load speed and network parameters. Thus, in this study, performance was interpreted primarily as a measure of content delivery, while an aspect critical to mobile learning – the CPU load during the execution of an AR scene – was overlooked.

A separate group consists of studies dedicated to functional accuracy. For instance, in [17], an evaluation of the WebXR Device API's performance was conducted, focusing on the accuracy of spatial tracking and the stability of virtual anchor binding through the analysis of trajectory errors and image differences. The only real-time performance metric was frame rate (FPS), which was used to confirm the smoothness of the visualization. However, the study lacks data on the impact on CPU load or device power consumption.

Attempts to quantify hardware costs are presented in [18], which analyzes CPU and RAM usage for different types of markers. Although the authors recorded specific CPU load levels, their methodology relied on the use of the *adb shell top* command. This approach provides only general aggregated data about the process, making it difficult to determine the contribution of individual functions, such as tracking, rendering, or video processing.

The practical aspect of deploying WebAR applications under conditions of limited hardware resources is examined in [19]. The authors implemented a contactless information system for a library based on the ar.js, three.js, and A-Frame libraries. To quantitatively assess performance on budget mobile devices supporting WebGL and WebRTC, four

variables were used: frame rate, animation frame request execution time, total load time, and the number of entity objects on the scene. Although this work directly considers the load time metric and analyzes rendering smoothness, the evaluation is conducted at the macro level. The lack of detailed browser profiling tools prevents an assessment of the load distribution across the processor's computational threads during scene initialization, which significantly limits the possibilities for targeted application optimization.

An analysis of the current state of research in the field of web performance evaluation reveals a significant imbalance in the selection of test objects and methods. According to systematic reviews, the vast majority of scientific works focus on the analysis of static content using basic macro-level metrics – page load speed, network latency [20]. However, for WebAR applications, this approach proves incapable of objectively reflecting the actual state of the system. The specific nature of WebAR requires continuous interaction with hardware, video stream processing, and 3D graphics rendering [21]. Existing testing methods do not provide a transparent assessment of the browser engine's internal processes: they do not allow for isolating the computational costs of JavaScript, graphics rendering, and memory management [22]. Although there is a wide range of analytical tools for evaluating web performance, the vast majority of them focus on network and macro-level metrics. Since such utilities are unable to detail the load distribution within the main thread and isolate the execution of JavaScript logic from rendering processes, the most appropriate and justified method for solving this problem is the use of instrumental browser profiling. This confirms the relevance of developing specialized approaches to profiling WebAR applications directly during the scene initialization phase.

3 Problem Statement

The goal of this work is to develop an approach for comprehensive performance evaluation and to identify patterns in the distribution of computational load during the initialization of WebAR applications using deep browser profiling tools. The study aims to quantitatively determine the CPU load on mobile devices, as well as to identify bottlenecks in the processes of loading, processing, and initial rendering of AR content. Only this level of instrumental detail allows for the acquisition of objective micro-level data, which constitutes the key

value of the study, as it forms a single reliable foundation for developing effective architectural strategies for optimizing application performance.

The scientific novelty of the study lies in the further development of a mathematical model of the WebAR application scene loading process on a mobile device. Unlike existing macro-level approaches, this work is the first to perform a deep decomposition of browser engine microprocesses. The proposed approach allows for an isolated assessment of the impact of the stages of spatial geometry parsing, direct rendering, and memory management on the blocking of the main execution thread. The level of detail achieved makes it possible to scientifically prove the nature of critical system overloads in a single-threaded environment and to formulate well-founded architectural strategies for optimizing the processing of three-dimensional data. **The object of this study** is the process of initializing and loading 3D content in a client-side WebAR application on a mobile device. **The subject of the study** is the patterns of CPU resource allocation and changes in browser engine timing metrics during the execution of WebAR scene loading stages.

4 Mathematical Model of the Scene Loading Process for a WebAR Application on a Mobile Device

4.1 Formalization of the Process

The mathematical formalization of 3D scene initialization in a WebAR application describes how the system transitions from camera and sensor data to the construction of a virtual scene in real-world coordinates.

In WebAR, several coordinate systems are typically used: the world coordinate system $CS \cdot W$; the camera coordinate system $CS \cdot C$; $CS \cdot O$; $CS \cdot S$. Transformations between them are defined by corresponding matrices.

When constructing the camera model, the position of a scene point in the world coordinate system is defined by a vector

$$P_w = \begin{bmatrix} X \\ Y \\ Z \\ 1 \end{bmatrix}. \quad (1)$$

The conversion to the camera system is performed using the following transformation [23]:

$$P_c = T_{cw} \cdot P_w, \quad (2)$$

where

$$T_{CW} = \begin{pmatrix} R & t \\ 0 & 1 \end{pmatrix}, \quad t = \begin{pmatrix} t_x \\ t_y \\ t_z \end{pmatrix}, \quad (3)$$

where t is the spatial displacement vector, and R is a 3×3 rotation matrix that satisfies the following requirements:

$$R \in SO(3), \quad \det R = 1, \quad R^T \cdot R = I = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}. \quad (4)$$

When generating a projection onto the screen plane, a matrix of the camera's internal parameters is used:

$$K = \begin{pmatrix} f_x & 0 & c_x \\ 0 & f_y & c_y \\ 0 & 0 & 1 \end{pmatrix}, \quad (5)$$

where f_x, f_y – focal lengths; c_x, c_y – image center coordinates.

In that case, the projection is calculated as follows [24]:

$$P = (u, v) = K \cdot (R|t) \cdot P_w, \quad (6)$$

where u, v – the coordinates of the corresponding pixel.

During the initialization phase, the following must be determined: the camera's position; the scene plane; the scale; and the initial anchor point – that is, a point in real space relative to which the AR system places and tracks virtual objects.

Camera pose refers to the camera's position and orientation in space, which are defined by a rotation matrix and a translation vector

$$\Theta = (R, t). \quad (7)$$

The system estimates this pose by analyzing images from the camera and tracking distinctive points or markers in the environment. To do this, it uses Simultaneous Localization and Mapping, Visual Odometry, and marker tracking methods, which help determine the camera's motion and coordinates in real time.

The model matrix describes how a 3D object is positioned in the scene: where it is located, how it is oriented, and what its dimensions are. It is formed as the product of the corresponding matrices [25]:

$$M = T \cdot R \cdot S, \quad (8)$$

where the T matrix controls the object's translation, the R matrix controls its rotation, and the S matrix controls its scaling. First, the object is scaled, then rotated, and finally moved to the desired position in

the scene. The resulting matrix is used by the graphics engine to correctly render the model in 3D space.

The final scene in 3D graphics is formed by sequentially applying several matrices to the object's vertex coordinates.

The full transformation is described by the central equation of WebAR visualization:

$$P_{clip} = P \cdot V \cdot M \cdot P_{local}, \quad (9)$$

where the M matrix transforms the vertex from the model's local coordinates to scene coordinates, the V matrix accounts for the camera's position and orientation, transforming the scene into the observer's coordinate system, and the P matrix projects the 3D scene onto a 2D screen, applying either a perspective or orthographic projection. After all these transformations, the screen coordinates of the vertex are obtained, which are used to render the final image.

Note that in markerless WebAR, the system first analyzes the surrounding space and attempts to find flat surfaces, such as a floor, table, or wall [26]. A plane is described by the equation

$$ax + by + cz + d = 0, \quad (10)$$

where the coefficients determine its position and orientation in space, and the normal to the surface

$$\vec{n} = (a, b, c) \quad (11)$$

shows the direction perpendicular to this plane.

Thanks to this normal vector, the system understands how to correctly orient the WebAR object relative to the surface. For example, a virtual chair can be automatically placed flat on the floor rather than at an angle. This method of surface detection is used in modern WebAR systems without special markers.

In marker-based WebAR, the process begins with the device's camera detecting a special marker [27]. After that, the algorithm finds the marker's key points or corners in the image. Based on the correspondence between the marker's points and their real-world geometry, the homography matrix is calculated using the following formula:

$$H = K \cdot \left(R + \frac{t \cdot n^T}{d} \right), \quad (12)$$

where d is the distance to the plane.

The resulting homography matrix describes the transformation between the marker plane and the camera image plane, allowing the system to determine the camera's pose and the marker's exact position in space. After calculating the homography, AR objects

are stably tied to the marker and move correctly with it in real time.

During the operation of AR systems, object coordinates may "jitter" slightly due to camera noise and tracking errors. To make the motion more stable, temporal filtering or data smoothing is used.

The mathematical model under development uses the exponential smoothing method [28]:

$$\hat{x}_k = \beta \cdot x_k + (1 - \beta) \cdot \hat{x}_{k-1}, \quad (13)$$

where x_k – a new position or angle measurement, and $\hat{x}_k$ – updated smoothed value. The coefficient β determines how much the system trusts the new data compared to previous values. As a result, AR objects move more smoothly and stably, without sudden jumps or flickering.

After determining the camera position, constructing the transformation matrices, and stabilizing the tracking, the system generates a complete WebAR scene. Formally, such a scene can be described as a set

$$\aleph = \{Q_i, C, L, T, \tilde{I}\}. \quad (14)$$

The component Q_i represents all 3D objects displayed in augmented reality. These can include furniture models, animations, text, or interactive elements. The component C describes the user's camera and its position in space. It is the camera that defines the viewpoint from which the scene is rendered. The element L is responsible for scene lighting and affects the realism of object rendering. For example, the system can adjust brightness or shadows according to the lighting of the real environment. The set T contains all object transformations, including translation, rotation, and scaling. These transformations are implemented through the model, view, and projection matrices used during scene rendering. The component $\tilde{I}$ describes user interactions with the WebAR environment. Such interactions include screen touches, gestures, device movement, or selecting AR objects in real time.

4.2 Performance Analysis

To analyze scene loading performance, it is advisable to examine the main stages of loading and preparing a WebAR application, as these stages determine the startup speed and smoothness of the scene. Since the maximum load on the application occurs during initialization and scene loading, the study focuses exclusively on four key stages: environment

initialization, resource loading, data parsing, and scene rendering.

The initialization stage includes launching the WebAR engine, creating the scene, configuring the camera, and connecting tracking modules. During loading, the system retrieves 3D models, textures, scripts, and other necessary resources.

During the parsing stage, the received data is processed and prepared for use by the graphics engine. This is followed by rendering, i.e., the construction of the final image and its display on the user's screen. For each of these stages, time characteristics, memory usage, and the load on the processor or graphics adapter can be determined. Thus, the performance of a WebAR application can be formalized as a set of key hardware and software parameters:

$$W_{perf} = (CPU, GPU, NET, MEM, FPS). \quad (15)$$

The *CPU* component characterizes the processor load that arises during the execution of computer vision algorithms, tracking, sensor data processing, and application logic. It is the central processing unit that performs most of the operations related to determining the camera's position, analyzing frames, finding keypoints, and calculating scene transformations. In markerless WebAR, the CPU load increases significantly due to the constant real-time execution of Simultaneous Localization and Mapping (SLAM) and Visual Odometry algorithms. High CPU load can lead to frame processing delays, overheating of the mobile device, and reduced stability of the AR scene.

The *GPU* component describes the load on the graphics processor during the rendering of 3D objects, lighting, textures, and animations. This component is responsible for constructing the final scene image and ensuring the visual quality of augmented reality. The more complex the models and effects used in the scene, the higher the load on the graphics adapter [29]. However, in most WebAR applications, the *GPU* primarily performs standard graphics operations, while the main tracking computations are performed by the *CPU* component.

The *NET* component characterizes the network overhead associated with loading models, textures, JavaScript libraries, and multimedia content. The volume of network traffic directly affects the launch speed of the WebAR application and the scene initialization time. This parameter is particularly critical for mobile devices with unstable Internet connections [30].

The *MEM* parameter describes RAM usage during application operation. Memory is used to store 3D models, textures, frame buffers, tracking results, and utility data structures. Excessive memory usage can cause performance degradation or force the WebAR scene to terminate on mobile devices.

The *FPS* (Frames Per Second) component characterizes the smoothness of scene rendering and the number of frames the system generates per second. For a comfortable augmented reality experience, it is usually necessary to maintain at least 30 FPS, while 60 FPS is considered optimal. A drop in FPS leads to image stuttering and instability of AR objects.

4.3 Assessing the CPU load of a WebAR application

Although WebAR performance depends on a combination of several parameters, the key factor in most scenarios is the *CPU* load, since the main algorithms for computer vision, camera pose estimation, environment analysis, and tracking are executed by the central processor in real time. Furthermore, unlike native AR applications, WebAR has an additional layer of browser abstraction, which increases the load specifically on the CPU. Additionally, even a slight CPU overload can cause both the accumulation of delays throughout the entire AR pipeline and disrupt the overall synchronization of the scene. Furthermore, on mobile devices, the CPU often becomes the primary source of power consumption and heat generation during AR operation [31]. Under sustained load, the system may automatically reduce the processor frequency (thermal throttling), which directly reduces FPS and tracking stability.

The CPU load of a WebAR application is determined by the volume of computations the system performs during scene analysis and augmented reality rendering. To estimate the load, the number of operations, their execution time, the processor clock frequency, the level of parallelism, and the intensity of data processing are taken into account. To calculate the load in the mathematical model, the following formula is used

$$CPU_{stage} = \frac{N_{ops}}{f_{cpu} \cdot \eta}, \quad (16)$$

where N_{ops} characterizes the number of arithmetic and logical operations performed at a specific stage of the WebAR system's operation, the parameter f_{cpu} determines the processor clock frequency and affects the speed of computation, while the coefficient η

describes the efficiency of CPU resource utilization, taking into account code optimization, parallel execution, and performance losses. As the number of operations or the complexity of computer vision algorithms increases, the processor load rises significantly [32]. Therefore, to ensure the stable operation of WebAR applications, it is important to optimize algorithms, reduce the computational load, and make effective use of multithreaded data processing.

Let's calculate the processor load step by step.

4.3.1 The initialization stage is the first and key step in preparing a 3D application for operation. At this stage, a WebGL context is created, which enables interaction with graphics hardware and sensors. Engine initialization includes loading the necessary libraries and settings, as well as preparing internal data structures. An important component is the creation of the scene graph – a structure that describes the hierarchy of scene objects and their relationships. Additionally, sensors are connected that are responsible for collecting information from the environment or the user.

The CPU load at this stage can be calculated using the following formula:

$$CPU_{init} = \frac{N_{ob} \cdot C_{ob} + N_{comp} \cdot C_{comp}}{f_{cpu} \cdot \eta_{init}}, \quad (17)$$

where N_{ob} – quantity of stage objects; C_{ob} – specific number of computational units (CU) per object; N_{comp} – number of components; C_{comp} – the specific number of CU for their initialization; coefficient η_{init} takes into account additional overhead costs associated with organizing the process.

This formula allows you to estimate how much CPU resources will be used during system startup. It is important for optimizing startup, especially in cases involving large scenes or complex components. The initialization phase is often critical in terms of startup time, so the efficiency of this step directly affects the overall performance of the application. Proper organization of this phase helps reduce delays.

4.3.2 The resource loading phase includes several important processes, such as HTTP request processing, data decompression, thread management, and caching. These operations are performed to ensure the correct loading and processing of information before it is displayed. HTTP processing involves analyzing incoming requests and responses and generating

appropriate responses. Decompression is required to unpack compressed data transmitted over the network to reduce traffic volume. Thread management allows for the efficient allocation of CPU resources among different tasks, ensuring parallelism and speed [33]. Caching helps reduce the number of repeated requests to the same resources, which improves performance.

To quantitatively estimate the CPU load during this stage, the following formula is used:

$$CPU_{load} = \frac{S_{dec} \cdot C_{data} + N_{res} \cdot C_{res}}{f_{cpu} \cdot \eta_{load}}, \quad (18)$$

where S_{dec} – the amount of data to be decompressed; C_{data} – the specific amount of decompression CU per byte; N_{res} – the number of requests; C_{res} – the specific amount of CU per request; η_{load} – CPU load efficiency.

The formula defines the total load as the sum of the resources spent on decompressing and processing each resource, taking into account their quantity and cost. This helps estimate the CPU power that will be required to perform all necessary operations during the loading phase. Balanced CPU utilization at this stage helps avoid delays and ensures fast content loading [34].

4.3.3 The parsing stage is the most CPU-intensive when processing 3D scenes. During this stage, structured data is parsed and interpreted, 3D models are decoded, shader programs are compiled, and the scene is constructed for subsequent rendering. JSON parsing involves converting text data into an internal format that the system can understand. Decoding 3D models is necessary to obtain geometry, textures, and animations. The formation of the scene graph organizes scene objects into a hierarchical structure, which facilitates the management of their positions and properties. These operations require significant computational resources, which is why the CPU load is at its peak at this point. To estimate this load, the following formula is used:

$$CPU_{parsing} = \frac{V_{sc} \cdot C_{vertex} + F_{pol} \cdot C_{pol} + S_{shader} \cdot C_{shader}}{f_{cpu} \cdot \eta_{parsing}}, \quad (19)$$

where V_{sc} – the number of vertices in the scene; C_{vertex} – the specific number of CUs required to process a single vertex; F_{pol} – the number of polygons; C_{pol} – the specific number of CUs required for a single polygon; S_{shader} – the number of shader programs; C_{shader} – the specific number of CUs required to compile a single

shader; $\eta_{parsing}$ – the processor utilization efficiency at this stage.

Parsing large models with a high number of vertices and polygons, as well as a large number of shader programs, significantly increases the demands on computational resources. This necessitates optimizing data structures and shader code to reduce processing time. Effective resource management at this stage allows you to prepare the scene for smooth and fast rendering. At the same time, it is important to monitor CPU utilization to avoid overloading the system. Proper balancing and optimization of the parsing stage are critical for WebAR performance.

It should be noted that in WebAR (augmented reality in a web browser), the time requirements for data parsing, frame processing, and script execution are extremely strict. Any delay causes visual stuttering, desynchronization of 3D objects from the real world, and can lead to motion sickness in the user. Here are the key timing requirements and constraints that the entire process is broken down into [35, 36]:

1. For AR to appear smooth, the frame rate must be 60 FPS (frames per second), so the total time per frame is 16.67 ms, or 1000 ms / 60 frames.
2. Any model parsing that takes longer than 2–3 ms will cause noticeable jank.
3. Motion-to-Photon Latency – the time from when the user physically rotates the phone to when the virtual object on the screen shifts to match the new position – should not exceed 20–30 ms. If computer vision algorithms (parsing of facial points, planes, or markers) run slower than 15 ms per frame, the virtual object will "lag" behind the real world during rapid movements.

4.3.4 The rendering stage is a key stage in the graphics process, where the CPU prepares all the information for transfer to the GPU. Several important operations are performed at this stage, such as preparing draw calls, updating transformation matrices, performing culling (removing unnecessary objects from the scene), and transferring commands directly to the GPU. Draw calls are commands that tell the GPU which objects to render and in what order. Updating matrices involves changing the transformations of objects or the camera, which determines how the scene will be rendered in the current frame. Culling saves resources by excluding objects that are not visible to the camera through visibility or collision checks. This reduces the load on the GPU by decreasing the amount of graphical

data being processed. The CPU load at this stage is determined by the following formula:

$$CPU_{render} = \frac{N_{draw} \cdot C_{draw} + N_{Renewal} \cdot C_{Renewal} + N_{cont} \cdot C_{cont}}{f_{cpu} \cdot \eta_{render}}, \quad (20)$$

where N_{draw} – number of draw calls; C_{draw} – cost of preparing one draw call; $N_{Renewal}$ – number of matrix updates; $C_{Renewal}$ – cost of renewal one matrix; N_{cont} – number of collision or visibility checks; C_{cont} – cost per check; η_{render} – processor efficiency.

The number of draw calls significantly impacts performance, as each call takes time to process. Matrix updates are essential for animation and the movement of objects in the scene. Visibility checks help avoid unnecessary rendering of invisible elements, which significantly improves efficiency [8]. This entire process must be optimized to achieve a high frame rate and smooth application performance.

4.4 CPU Load Balancing

Let the execution time for each of the considered stages of loading a WebAR application scene on a mobile device be T_{init} , T_{load} , T_{parse} , T_{render} accordingly, and after performing the calculations using equations (17) through (20), the following values for the required

$$\theta_{\xi}(\mathbf{Ind}) = (\theta_{\xi}(Ind_1), \theta_{\xi}(Ind_2), \theta_{\xi}(Ind_3), \theta_{\xi}(Ind_4)). \quad (23)$$

In that case, the problem of finding the optimal CPU load balance can be formulated as follows:

$$\theta_{\xi}^*(\mathbf{Ind}) = \arg \left(\gamma(\theta_{\xi}(\mathbf{Ind})) \right) \Big| \gamma(\theta_{\xi}(\mathbf{Ind})) > \gamma(\theta_{\xi}^*(\mathbf{Ind})) \quad \forall \theta_{\xi} \in \Theta, \theta_{\xi} \neq \theta_{\xi}^*, \quad (24)$$

where $\Theta = \{\theta_{\xi}\}$ – the set of all possible stage-override procedures that satisfy all the timing requirements and constraints 1–3 for the parsing stage.

5 Study of the scene loading process for a WebAR app on a Google Pixel 8 smartphone

5.1 Experimental conditions

This study evaluates the performance of a WebAR application developed using *Angular 19* and the *A-Frame 1.7.1* library. A Location AR scene containing a single AR object with a size of 16.4 MB was selected for testing. Since the maximum load on the application occurs during initialization and scene loading, the study focuses exclusively on this stage. CPU load in the context

CPU time were obtained: CPU_{init} , CPU_{load} , CPU_{parse} , CPU_{render} . For each stage ($i=1\dots 4$), we will calculate the load index using the formula:

$$Ind_i = \frac{CPU_i}{T_i} \cdot 100\%, \quad i = \overline{1, 4}, \quad (21)$$

where i corresponds to the sequential number of the WebAR scene loading stage in the application.

We will define the current CPU load index as a vector $\mathbf{Ind} = (Ind_1, Ind_2, Ind_3, Ind_4)$

Let's define the CPU load balance metric as

$$\gamma(\mathbf{Ind}) = \frac{\sum_{i=1}^4 |Ind_i - \sum_{i=1}^4 Ind_i / 4|}{\sum_{i=1}^4 Ind_i}. \quad (22)$$

After decomposing the stages, it becomes possible to relieve the overloaded stages by performing certain actions on the less-loaded preceding stages. As a result of this procedure θ_{ξ} , we obtain a new load index

of the web page was selected as the primary metric for evaluating performance.

Performance metrics were collected using the *Chrome DevTools* developer tools. The test mobile device was connected to the computer via a USB interface. The Performance tab was used to measure CPU parameters. It details the activity of the WebAR application, showing the time spent on Scripting, Rendering, and other categories. It is important to note that this tool does not reflect the overall CPU load of the entire device; therefore, the data obtained should be interpreted exclusively as indicators of the web page's load. Since the process of using *Chrome DevTools* itself creates an additional load, each test scenario was run 10 times. For further analysis, the run with the average values was selected.

Testing was conducted on a Google Pixel 8 smartphone (October 2023 model). The test device is equipped with a Google Tensor G3 processor, has 8 GB of RAM, and runs on the Android 16 operating system.

5.2 Experiment Description

To ensure the reproducibility of results and minimize the impact of uncontrolled background processes, a device preparation procedure was implemented. Before testing began, the mobile device's battery was charged to a level above 90%. After that, a full system reboot was performed to clear the RAM and force-close background processes. The display brightness was set to 80%. All third-party apps were disabled; only the website with the WebAR app remained active in the web browser.

To ensure consistent lighting conditions and stabilize the camera angle, the device was securely mounted in a static position using a tripod.

At the software level, developer mode and USB debugging were enabled (Fig. 1). The physical connection to the host computer was established using high-speed cables. The connection was verified using the *Chrome DevTools*, via the *chrome://inspect* page, which confirmed the successful detection of the remote device and the availability of the target browser tab for inspection (Fig. 2).

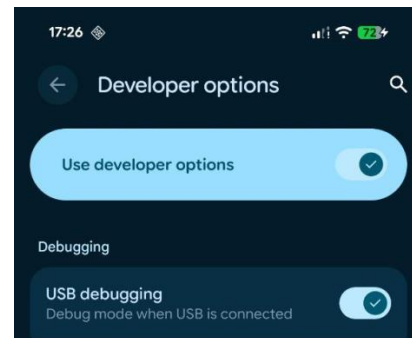


Fig. 1. Developer mode and USB debugging enabled

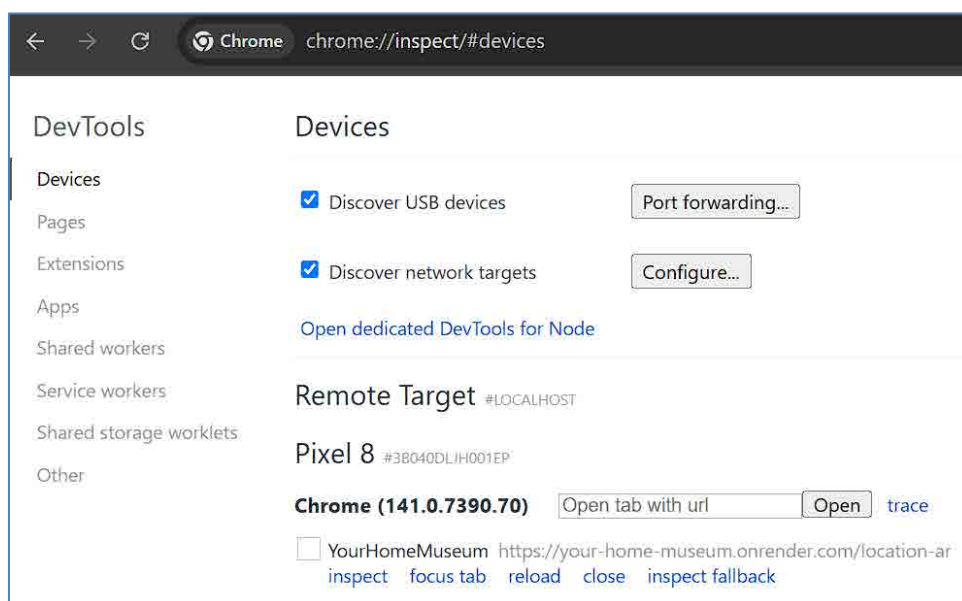


Fig. 2. Displaying the mobile tab on a computer

After completing the preparatory phase, the data collection procedure was implemented. It involved performing ten consecutive iterations of performance profiling using the Performance tab in *Chrome DevTools*. The experimental scenario covered the full cycle of loading and initializing a geolocated WebAR scene.

Upon completion of each measurement iteration, trace files containing detailed performance metrics were exported. Specifically, the time spent on Scripting, System processes, Rendering, and Painting was analyzed.

Additionally, the total script execution time (Total), average frame rate (FPS), and JavaScript heap growth were recorded.

To select the most representative sample for further detailed analysis, a comparison was made of the peak values of JS Heap memory usage and the duration of key processing stages. It is important to note that the Total metric also accounts for periods of system idle time. The resulting quantitative data has been organized and presented in Table 1.

Table 1. Test results on the Pixel 8 device

Test	JS Heap (mb)	Scripting, ms	System, ms	Rendering, ms	Painting, ms	Total, ms
1	60	2,209	558	186	145	6,743
2	60,3	2,289	573	166	127	6,401
3	41,1	1,649	438	232	153	6,181
4	42,8	2,359	543	260	161	7,238
5	45	1,705	462	214	158	5,929
6	26	2,352	580	183	140	5,898
7	42,1	2,502	568	196	156	6,262
8	42,6	2,534	845	203	180	6,898
9	42,7	2,454	574	140	135	5,904
10	60,1	2,820	668	209	165	6,765

5.3 Choosing a representative sample

To identify a representative test among the ten results obtained, the *normalized deviation method* was applied. This approach makes it possible to objectively identify the test whose scores are closest to the average performance conditions, minimizing the impact of random fluctuations and outliers.

For each test, the normalized deviation was calculated using the corresponding formula:

$$Score_i = \sum_j \frac{|X_{ij} - Me_j|}{IQR_j} \quad (25)$$

where X_{ij} – the value of the j -th metric for the i -th test; Me_j – the median of the corresponding metric; IQR_j – the interquartile range of this metric.

The interquartile range is calculated as the difference between the third and first quartiles of the sample and is determined using the formula:

$$IQR = Q_3 - Q_1 \quad (26)$$

where Q_1 – 25th percentile (lower quartile); Q_3 – 75th percentile (upper quartile).

For a sample of size $n=10$, the calculation was performed using Tukey's method: the sorted data set was divided into two equal parts of 5 elements each. In this case, Q_1 is defined as the median of the first half of the set, the 3rd element in the sorted list, and Q_3 as the median of the second half, the 8th element.

Based on Tukey's method, basic statistical indicators were determined for each metric. The results of determining the median, Q_1 , Q_3 and calculating the interquartile range IQR using equation (26) for each column are presented in Table 2.

The results of the calculation of the standardized deviation for each test using Equation (1) are presented in Table 3.

Table 2. Values of the indicators based on the Tukey method

Metric	Median	Q_1 (3-rd element)	Q_3 (8-th element)	IQR
JS Heap	42,75	42,1	60,0	17,9
Scripting	2,356	2,209	2,502	0,293
System	570,5	543	580	37,0
Rendering	199,5	183	214	31,0
Painting	154,5	140	161	21,0
Total	6,332	5,929	6,765	0,836

Table 3. Normalized test deviations on Pixel 8 device

Test	Score
1	3,18
2	3,75
3	7,39
4	4,10
5	6,40
6	2,95
7	0,87
8	10,04
9	3,79
10	6,51

The results of the calculations show that Test No. 7 has the lowest normalized value, indicating the smallest deviation of its parameters from the median characteristics of the sample. This indicates the greatest correspondence of the results obtained in this test to the typical conditions of the experiment. Thus, the specified recording was identified as representative and used for further comparative analysis between the devices under study.

5.4 Analysis of the representative recording

The representative recording obtained in the *Performance* tab of the *Chrome DevTools* toolkit is shown in Fig. 3. The analysis focuses on the application initialization period; therefore, the intervals at the beginning and end of the recording are not included. The main visual metrics are: the CPU activity graph at the top of the figure; the Main thread, which records function calls; the JS Heap memory usage graph; and the Summary section.

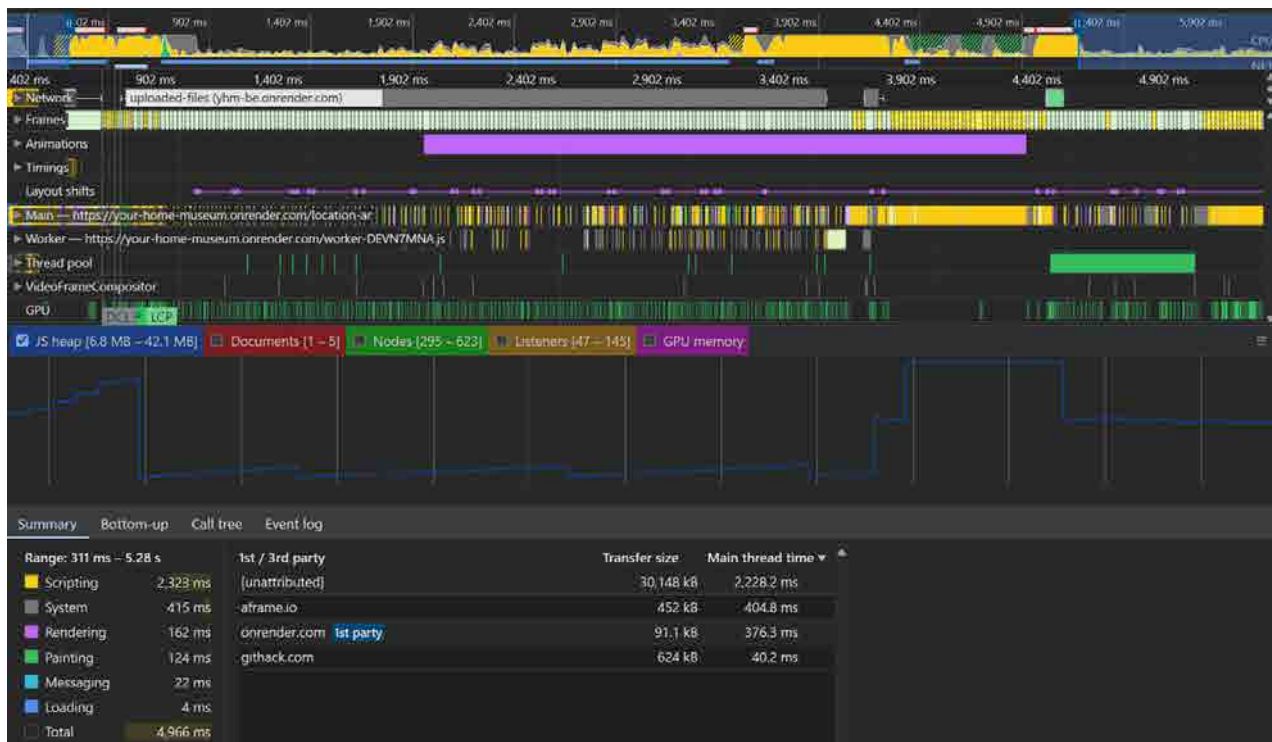


Fig. 3. Representative recording

The Summary section indicates that the total runtime of the main thread is 4.96 seconds, of which: 2.32 seconds are spent on scripts, 1.91 seconds on idle time, 0.41 seconds on system overhead, 0.16 seconds on page rendering, and 0.12 seconds on content rendering. This data is presented in a chart in Fig. 4. Inter-process communication (Messaging) and page loading are not included, as they account for less than 1% of the total time. Thus, the Scripting process accounts for the largest share (47%), followed closely by Idle (39%).

The CPU graph is divided into four intervals (Fig. 5), three of which show high CPU load, as indicated by the sections filled entirely in yellow. The unhighlighted intervals are not included because they are not directly related to the application's initialization.

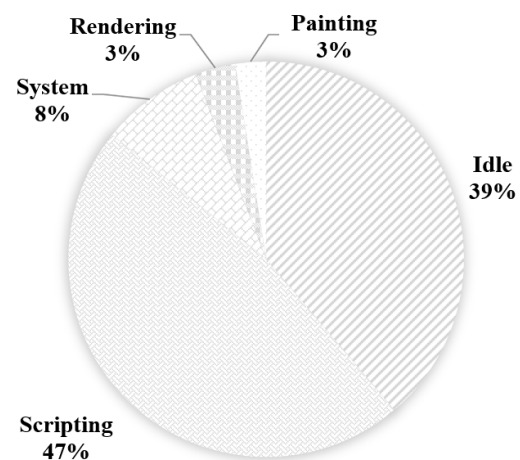


Fig. 4. Flowchart of the application loading process



Fig. 5. Division of the recording into intervals

The *first interval* of the profiling covers the time range from 0.311 to 0.795 seconds. This stage is characterized by a high computational load due to the initialization of the client application. The Scripting

process accounts for the dominant share of time, over 50% of the entire interval. This computational activity is primarily driven by two key operations. The Unattributed category had the longest duration among them, at

108 milliseconds. A detailed trace analysis revealed that this category involves the execution of native code by the JavaScript engine. Due to the specific nature of code optimization in the Angular framework, the profiler was unable to correlate these operations with specific function calls, which is typical for such architectures. The second significant component of script execution was the initialization of the A-Frame library, which took 89.1 milliseconds. An additional indicator of the active initialization phase is the rapid increase in JS Heap memory consumption from 26 to 36 megabytes. This dynamic indicates the mass creation of virtual DOM objects and the preparation of a three-dimensional scene. Overall, during this interval, the system accumulates resources for processing logic and data structures, and at the end of the stage, it performs the initial page rendering process, which takes only 28 milliseconds.

Upon completion of the initialization phase, a period of *moderate computational intensity* (0.796–3.62 s) is observed. During this period, the dominant state of the main thread is the waiting mode – 1.6 s. The application requests and waits for the transfer of a 3D model from the server, which keeps the CPU load low. Computational activity does not include blocking tasks and is primarily driven by operations such as reading the input stream, data transformation, and background processes. An important characteristic of this stage is memory optimization: instead of data accumulation, the Garbage Collector was triggered, reducing the size of the JS Heap from 37 MB to 7 MB. This indicates effective resource release after initialization and before loading the 3D model.

The third stage (3.62 s – 4.45 s) is the heavy data processing phase. After the AR object finished loading from the server, the main thread entered a state of peak load – scripting took up over 82% of the time. The longest blocking task, lasting 550 ms, was recorded during this interval. Analysis shows that this is the process of parsing the 3D model's binary data and converting it into Three.js scene structures. This is

confirmed by a sharp increase in memory usage: from 8.5 MB at the start of the interval to 42.1 MB. This spike corresponds to the decompression of the compressed model file into full-fledged objects in RAM.

The final phase (4.45 s – 5.3 s) is the 3D model rendering period. In the first half of this time interval, the browser utilizes a pool of background threads – Thread Pool Workers – which enables the parallel execution of decoding operations and the preparation of graphical resources. In the recorded trace, three workers are observed, each processing separate fragments of the model or associated texture data; this activity can be classified as the preparatory phase of rendering. While the main thread's activity is minimal in the first half of the interval, the CPU load becomes dominant in the second half. It contains calls to `WebGLRenderer.render`, `renderScene`, `renderObjects`, `renderBufferDirect`, as well as `WebGLUniforms.upload` and `safeSetTexture2D` operations. This indicates the compilation of shader programs, the loading of textures, and the direct rendering of the 3D model in the graphics scene.

The results of the analysis of a representative recording indicate four distinct stages of application loading: an intensive initialization phase of the initial empty scene lasting from 0.311 s to 0.795 s, which took 0.48 s; a moderate period of loading the 3D object from the server-side from 0.796 s to 3.62 s, which took 2.82 s; a very resource-intensive period of 3D object parsing from 3.62 s to 4.45 s, which lasted 0.83 s; and an intensive phase of 3D object rendering from 4.45 s to 5.3 s, which lasted 0.85 s.

5.5 Results

To determine the most resource-intensive stage of the WebAR application's loading process, four intervals were considered: initialization, resource loading, parsing, and rendering. For each interval, the loading index was calculated using formula (21). The results of the calculations are presented in Table 4.

Table 4. Processor performance metrics at key stages of 3D application development

№	Stage	Interval, s	T_p , ms	CPU_p , ms	Dwell time, ms	Ind_p , %
1	Initialization	0.311 – 0.795	484	311.9	172.1	64.4
2	Loading	0.796 – 3.62	2824	1318.7	1505.3	46.7
3	Parsing	3.62 – 4.45	830	732.4	97.6	88.2
4	Rendering	4.45 – 5.30	850	540.9	309.1	63.6

The most critical stage turned out to be the third *parsing* interval, during which the peak CPU load of 88.2%

was recorded (Fig. 6). Despite its relatively short duration (0.83 s), the density of computational operations during this

interval is the highest. This is due to the synchronous nature of the 3D model binary data deserialization process and the massive creation of objects in memory, which leaves the processor almost no idle time.

The *initialization* (64.4%) and *rendering* (63.6%) stages show similar average load metrics but differ

in the nature of the operations. While in the first stage the CPU is primarily occupied with JIT compilation and executing JavaScript library code, in the fourth stage interaction with the GPU and frame preparation play a significant role.

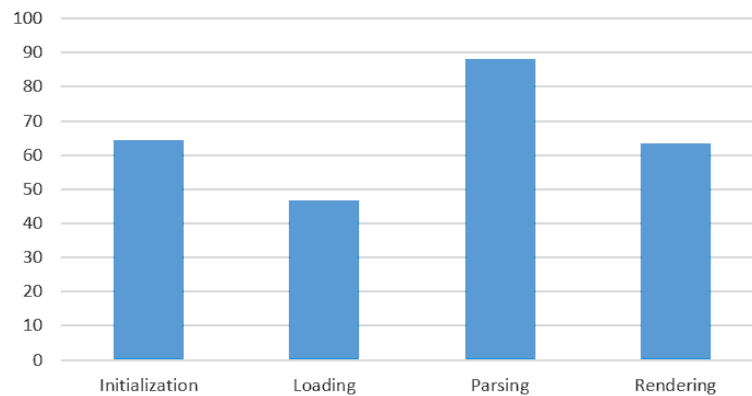


Fig. 6. Percentage load graph of intervals

The least loaded stage is *resource loading* (46.7%), which is an expected result, since network bandwidth is the limiting factor in this interval, and the main processor thread spends more than half the time waiting for data.

Thus, the highest computational load is observed during the third *Parsing* interval, while the *Loading* stage is the least loaded and has the highest idle time.

Further performance optimization should be based on resource reallocation: reducing idle time in the *Loading* stage by transferring some computational tasks from the *Parsing* stage to it. This approach will ensure the maximum increase in application performance.

6 Practical Significance

The practical significance of the research results lies in the possibility of their direct implementation by companies specializing in the development of commercial augmented reality applications. The modern e-commerce, interactive marketing, and digital education industries are actively adopting augmented reality technologies to improve customer engagement. However, the success of such solutions critically depends on the smooth operation of the software on consumers' end devices. The approach to deep profiling of browser engine microprocesses proposed in this work provides the industry with an effective tool for identifying hidden technical flaws that cannot be detected using standard network analytics tools. The refined mathematical model of the process and the empirical

data obtained form a solid scientific foundation for the transition from intuitive development to engineering-based design of client application architecture.

The key issue hindering the widespread and effective implementation of browser-based augmented reality in a business environment is the loss of control over the interface during the loading of volumetric spatial content. When a user initiates the viewing of a three-dimensional product model, there is intensive consumption of the mobile device's computational resources. The peak load on the central processing unit during geometry deserialization, as identified in the study, causes a prolonged blockage of the main execution thread. For the end user, this appears as a complete system freeze: the screen does not respond to touch, animations stop, and overall performance degrades sharply. Such a negative technical experience destroys the sense of immersion and leads to disappointment with the quality of the digital service provided.

The economic efficiency of any digital product is directly correlated with user experience metrics. In the highly competitive e-commerce market, companies think in terms of financial return and customer satisfaction levels. Technical glitches caused by inefficient allocation of processor resources have direct financial consequences. Users tend to immediately leave a webpage if its interface becomes unresponsive, even for a fraction of a second. Accordingly, eliminating architectural overloads – the nature of which was demonstrated in this study – will allow businesses to

significantly reduce user drop-off rates. Maintaining interface responsiveness ensures that a potential buyer will complete the intended action, whether it is viewing a product from all angles or placing an order. Consequently, optimizing the 3D scene initialization process becomes not merely a technical task but a strategic tool for increasing conversion rates in commercial augmented reality services.

To achieve these economic and technical benefits, developers of commercial platforms must implement new engineering standards for working with spatial data. The study's findings indicate the need to move away from monolithic synchronous loading of 3D models in the browser's main thread. Instead, the industry is advised to transition to methods of incremental geometry loading and offloading resource-intensive tasks to background computing processes. The use of modern 3D data compression formats combined with asynchronous logic processing will help balance the load on smartphone hardware. Implementing such architectural solutions will ensure a stable frame rate and high system responsiveness regardless of the cost and technical specifications of the consumer's mobile device.

Thus, this work goes beyond purely theoretical analysis and offers the industry a ready-made analytical framework. The application of the outlined principles for evaluating and optimizing performance opens the way for companies to create high-quality, stable, and cost-effective products. Meeting technical performance requirements is inextricably linked to meeting consumer needs, making the results of this study an integral part of the process of improving commercial augmented reality ecosystems.

7 Conclusions

As a result of this study, an approach was implemented for the comprehensive evaluation of performance and the identification of patterns in the distribution of computational load during the initialization of WebAR applications. To this end, an improved mathematical model of the WebAR application scene loading process on a mobile device was applied. Through the use of deep browser profiling tools, the microprocesses of the browser engine were decomposed, and the nonlinear nature of mobile device hardware resource consumption was empirically confirmed.

Analysis of isolated time profiles of a client application based on the Angular 19 component architecture and the A-Frame 1.7.1 declarative framework allowed us to identify the key phases of scene initialization. It has been scientifically proven that the critical factor in performance degradation is not the stage of direct graphics rendering, which demonstrates a stable load of 63.6 percent, but the process of parsing spatial geometry. It is during the parsing stage that the peak CPU load of 88.2 percent was recorded, along with the longest blocking task lasting 550 milliseconds, caused by the synchronous deserialization of the model's binary data. At the same time, the effectiveness of memory management in the studied stack was confirmed: the garbage collection mechanism was observed to function correctly, ensuring a rapid reduction of the peak JS Heap size from 40 megabytes to a stable operating level of 22 megabytes.

The obtained objective micro-level data form a reliable foundation for developing well-founded architectural strategies for optimizing the performance of three-dimensional web applications. Since processing a monolithic model in the main thread causes critical blocking of the main execution thread in a single-threaded environment, traditional approaches to initialization require a fundamental reevaluation. A promising and practical approach to eliminating the identified bottlenecks is the use of asynchronous and incremental geometry loading methods, as well as the mandatory delegation of resource-intensive deserialization processes to background computational threads. Implementing such engineering solutions will allow developers of commercial WebAR services to avoid critical interface delays, ensuring high user retention and improving the overall economic efficiency of digital products.

Conflict of interest

The authors declare that they have no conflict of interest with respect to this research, including financial, personal, authorship, or other conflicts that could influence the research and its results presented in this article.

Funding

The study was conducted without financial support.

Data availability

The manuscript has no associated data.

Use of artificial intelligence

The authors confirm that they did not use artificial intelligence technology in the creation of this work.

Acknowledgments

The study is being conducted within the framework of the research project 0126U003139 of the National Fund of Ukraine on the topic "Scientific Foundations for the Formation and Management of Human Capital in a Multi-Project Environment to Support the Sustainable Development of Ukraine's Recovery Programs".

References

1. Kayser, I. (2026), "Technology, Digital Transformation and Society: A Closing Editorial", *Social Sciences*, Vol. 15(4), 251 p. DOI: <https://doi.org/10.3390/socsci15040251>
2. Kalyoncu, F., Karal, H. (2025), "Design and development of augmented reality application for basic concepts of computer systems", *Education and Information Technologies*, Vol. 30(1), pp. 347-375. DOI: <https://doi.org/10.1007/s10639-024-12971-x>
3. Nadeem, M., Lal, M., Cen, J., Sharsheer, M. (2022), "AR4FSM: Mobile Augmented Reality Application in Engineering Education for Finite-State Machine Understanding", *Education Sciences*, Vol. 12(8), 555 p. DOI: <https://doi.org/10.3390/educsci12080555>
4. Butt, A., Ahmad, H., Muzaffar, A., Ali, F., Shafique, N. (2022), "WOW, the make-up AR app is impressive: a comparative study between China and South Korea", *Journal of Services Marketing*, Vol. 36, No. 1 pp. 73–88. DOI: <https://doi.org/10.1108/JSM-12-2020-0508>
5. Yin, C. Z. Y., Jung, T., tom Dieck, M. C., Lee, M. Y. (2021), "Mobile Augmented Reality Heritage Applications: Meeting the Needs of Heritage Tourists", *Sustainability*, Vol. 13(5), 2523 p. DOI: <https://doi.org/10.3390/su13052523>
6. Parekh, P., Patel, S., Patel, N., Shah M. (2020), "Systematic review and meta-analysis of augmented reality in medicine, retail, and games", *Visual Computing for Industry, Biomedicine, and Art*, Vol. 3, No. 1, 21 p. DOI: <https://doi.org/10.1186/s42492-020-00057-7>
7. Udayan, J. D., Kataria, G., Yadav, R., Kothari, S. (2020), "Augmented Reality in Brand Building and Marketing – Valves Industry", *2020 International Conference on Emerging Trends in Information Technology and Engineering (ic-ETITE)*, pp. 1–6. DOI: <https://doi.org/10.1109/ic-ETITE47903.2020.425>
8. Boutsis, A.M., Ioannidis, C., Verykokou, S. (2023), "Multi-Resolution 3D Rendering for High-Performance Web AR", *Sensors*, Vol. 23 (15), 6885 p. DOI: <https://doi.org/10.3390/s23156885>
9. Alsulami, M. H., Khayyat, M. M., Aboulola, O. I., Alsaqer, M. S. (2021), "Development of an Approach to Evaluate Website Effectiveness", *Sustainability*, Vol. 13(23), 13304 p. DOI: <https://doi.org/10.3390/su132313304>
10. Morales-Vargas, A., Pedraza-Jimenez, R., Codina, L. (2022), "Website quality in digital media: literature review on general evaluation methods and indicators and reliability attributes", *Revista Latina de Comunicacion Social*, Vol. 80, pp. 39–63. DOI: <https://doi.org/10.4185/RLCS-2022-1515>
11. Kuchuk, G., Kharchenko, V., Kovalenko, A., Ruchkov, E. (2016), "Approaches to selection of combinatorial algorithm for optimization in network traffic control of safety-critical systems", *Proceedings of 2016 IEEE East-West Design and Test Symposium*, EWDTS 2016, 7807655 p. DOI: <https://doi.org/10.1109/EWDTS.2016.7807655>
12. Matvieiev, M. (2024), "Analysis of Optimization Methods for Augmented Reality Web Applications", *Control, Navigation and Communication Systems*, No. 4(78), pp. 106-108. DOI: <https://doi.org/10.26906/SUNZ.2024.4.106>
13. Ghattas, M. M., Sartawi, P. D. B. (2020), "Performance Evaluation of Websites Using Machine Learning", *EIMJ*, Vol. 51, pp. 36–41, available at: https://www.researchgate.net/publication/345975296_Performance_Evaluation_of_Websites_Using_Machine_Learning
14. Tuah, N.M., Ahmad, W.N.W., Andrias, R.M., Dg. S. Ajor, Sura, S., Rodzuan, A.R.A. (2025), "Assessing the user experience of marker-based 3D WebAR applications using user experience questionnaire", *International Journal of Informatics and Communication Technology*, Vol. 14(1), pp. 31–41. DOI: <https://doi.org/10.11591/ijict.v14i1.pp31-41>
15. Ghattas, M., Mora, A. M., Odeh, S. (2025), "A Novel Approach for Evaluating Web Page Performance Based on Machine Learning Algorithms and Optimization Algorithms", *AI*, Vol. 6., No. 2., 19 p. DOI: <https://doi.org/10.3390/ai6020019>
16. Hussein, H. A., Ali, M. H., Al-Hashimi, M., Majeed, N. T., Hameed, Q. A., Ismael, R. D. (2023), "The Effect of Web Augmented Reality on Primary Pupils' Achievement in English", *Applied System Innovation*, Vol. 6(1), 18 p. DOI: <https://doi.org/10.3390/asi6010018>
17. Lee, D., Shim, W., Lee, M., Lee, S., Jung, K.-D., Kwon, S. (2021), "Performance Evaluation of Ground AR Anchor with WebXR Device API", *Applied Sciences*, Vol. 11(17), 7877 p. DOI: <https://doi.org/10.3390/app11177877>

18. Angeleri, C., Marini, F., Alvarez, D., Leale, G., Curras, D. (2021), "CPU and RAM Performance Assessment for Different Marker Types in Augmented Reality Applications", *Proceedings of the 16th International Joint Conference on Computer Vision, Imaging and Computer Graphics Theory and Applications*, pp. 271–277. DOI: <https://doi.org/10.5220/0010302302710277>
19. Roy, S. G., Kanjilal U. (2021), "Web-based Augmented Reality for Information Delivery Services: A Performance Study", *DESIDOC Journal of Library & Information Technology*, Vol. 41, No. 3, pp. 167–174. DOI: <https://doi.org/10.14429/djlit.41.03.16428>
20. Ghattas, M., Odeh, S., Mora, A. M. (2025), "Predicting Website Performance: A Systematic Review of Metrics, Methods, and Research Gaps (2010–2024)", *Computers*, Vol. 14, No. 10, 446 p. DOI: <https://doi.org/10.3390/computers14100446>
21. Najadat, H., Al-Badarnah, A., Alodibat, S. (2021), "A review of website evaluation using web diagnostic tools and data envelopment analysis", *Bulletin of Electrical Eng. and Informatics*, Vol. 10, pp. 258–265. DOI: <https://doi.org/10.11591/eei.v10i1.1755>
22. Niazi, M. G., Kamran, M. K. A., Ghaebi, A. (2020), "Presenting a proposed framework for evaluating university websites", *Electronic Library*, Vol. 38, No. 5–6, pp. 881–904. DOI: <https://doi.org/10.1108/EL-06-2020-0141>
23. Romanenkov, Yu., Mukhin, V., Kosenko, V., Revenko, D., Lobach, O., Kosenko, N., Yakovleva, A. (2024), "Criterion for Ranking Interval Alternatives in a Decision-Making Task", *International Journal of Modern Education and Computer Science*, Vol. 16(2), pp. 72–82. DOI: <https://doi.org/10.5815/ijmecs.2024.02.06>
24. Kovalenko, A., Kuchuk, H., Kuchuk, N., Kostolny, J. (2021), "Horizontal scaling method for a hyperconverged network", *2021 International Conference on Information and Digital Technologies (IDT)*, Zilina, Slovakia, DOI: <https://doi.org/10.1109/IDT52577.2021.9497534>
25. Kuchuk, H., Matvieiev, M. (2025), "Modeling the process of loading 3D models in a client application", *Advanced Information Systems*, Vol. 9(4), pp. 11–16. DOI: <https://doi.org/10.20998/2522-9052.2025.4.02>
26. Kumar, A., Arora, A. (2019), "A Filter-Wrapper based Feature Selection for Optimized Website Quality Prediction", *Proceedings 2019 Amity International Conference on Artificial Intelligence (AICAI)*, pp. 284–291. DOI: <https://doi.org/10.1109/aicai.2019.8701362>
27. Allison, R., Hayes, C., McNulty, C. A. M., Young, V. (2019), "A Comprehensive Framework to Evaluate Websites: Literature Review and Development of GoodWeb", *JMIR Formative Research*, Vol. 3(4). DOI: <https://doi.org/10.2196/14372>
28. Semenov, S., Mozhaiev, O., Kuchuk, N., Mozhaiev, M., Tiulieniev, S., Gnusov, Yu., Yevstrat, D., Chyrva, Y., Kuchuk, H. (2022), "Devising a procedure for defining the general criteria of abnormal behavior of a computer system based on the improved criterion of uniformity of input data samples", *Eastern-European Journal of Enterprise Technologies*, Vol. 6(4–120), pp. 40–49. DOI: <https://doi.org/10.15587/1729-4061.2022.269128>
29. Kuchuk, H., Kovalenko, A., Ibrahim, B.F. Ruban, I. (2019), "Adaptive compression method for video information", *International Journal of Advanced Trends in Computer Science and Engineering*, Vol. 8(1), pp. 66–69. DOI: <http://dx.doi.org/10.30534/ijatse/2019/1181.22019>
30. Kuchuk, H., Kalinin, Y., Dotsenko, N., Chumachenko, I., Pakhomov, Y. (2024), "A Decomposition of integrated high-density IoT data flow", *Advanced Information Systems*, Vol. 8, No. 3, pp. 77–84. DOI: <https://doi.org/10.20998/2522-9052.2024.3.09>
31. Kuchuk, H., Mozhaiev, O., Tiulieniev, S., Mozhaiev, M., Kuchuk, N., Lubentsov, A., Onishchenko, Yu., Gnusov, Yu., Brendel, O., Roh, V. (2025), "Devising a method for energy-efficient control over a data transmission process across the mobile high-density Internet of Things", *Eastern European Journal of Enterprise Technologies*, Vol. 4(4(136)), pp. 46–57, DOI: <https://doi.org/10.15587/1729-4061.2025.336111>
32. Kuchuk, N., Kashkevich, S., Radchenko, V., Andrusenko, Y. Kuchuk, H. (2024), "Applying edge computing in the execution IoT operative transactions", *Advanced Information Systems*, Vol. 8, No. 4, pp. 49–59. DOI: <https://doi.org/10.20998/2522-9052.2024.4.07>
33. Kuchuk, H., Kosenko, N., Kuchuk, N., Gopejenko, V., Kosenko, V. (2026), "Adaptive Resource Allocation Method for the Mobile Fog Layer of High-Density Industrial Internet of Things in Industry 5.0 Networks", *Innovative technologies and scientific solutions for industries*, Vol. 1(35), pp. 65–78. DOI: <https://doi.org/10.30837/2522-9818.2026.1.065>
34. Kuchuk, H., Husieva, Y., Novoselov, S., Lysytsia, D., Krykhovetskyi, H. (2025), "Load Balancing of the layers Iot Fog-Cloud support network", *Advanced Information Systems*, Vol. 9, No. 1, pp. 91–98. DOI: <https://doi.org/10.20998/2522-9052.2025.1.11>
35. Muff, F., Borcard, D., Fill, HG. (2026), "MM-AR: a web-based open-source metamodeling platform for spatial conceptual modelling", *Software and Systems Modeling*. DOI: <https://doi.org/10.1007/s10270-025-01351-9>
36. Boutsis, A.-M., Ioannidis, C., Verykokou, S. (2023), "Multi-Resolution 3D Rendering for High-Performance Web AR", *Sensors*, Vol. 23(15), 6885 p. DOI: <https://doi.org/10.3390/s23156885>

Received (Надійшла) 24.02.2026

Accepted for publication (Прийнята до друку) 21.05.2026

Publication date (Дата публікації) 27.06.2026

Відомості про авторів / About the Authors

Кучук Георгій Анатолійович – доктор технічних наук, професор, Національний технічний університет "Харківський політехнічний інститут", професор кафедри комп'ютерної інженерії та програмування, Харків, Україна;

Heorhii Kuchuk – Doctor of Technical Sciences, Professor, National Technical University "Kharkiv Polytechnic Institute", Professor of the Computer Engineering and Programming Department, Kharkiv, Ukraine;

e-mail: kuchuk56@ukr.net

ORCID ID: <http://orcid.org/0000-0002-2862-438X>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=57057781300>

Матвєєв Микита Іванович – Національний технічний університет "Харківський політехнічний інститут", здобувач вищої освіти ступеня доктора філософії кафедри комп'ютерної інженерії та програмування, Харків, Україна;

Mykyta Matvieiev – National Technical University "Kharkiv Polytechnic Institute", Postgraduate Student of the Computer Engineering and Programming Department, Kharkiv, Ukraine;

e-mail: Mykyta.Matvieiev@cit.khpi.edu.ua

ORCID ID: <https://orcid.org/0009-0000-8773-6640>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=60141774400>

Косенко Наталія Вікторівна – кандидат технічних наук, доцент, Харківський національний університет міського господарства ім. О.М. Бекетова, доцент кафедри управління проектами у міському господарстві і будівництві, Харків, Україна;

Nataliia Kosenko – Candidate of Technical Sciences, Associate Professor, Beketov National University of Urban Economy in Kharkiv, Associate Professor of the Project Management in Urban Economy and Construction Department, Kharkiv, Ukraine;

e-mail: kosnatalja@gmail.com

ORCID ID: <https://orcid.org/0000-0002-5942-3150>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=57196219605>

Лисиця Дмитро Олександрович – кандидат технічних наук, доцент, Національний технічний університет "Харківський політехнічний інститут", доцент кафедри комп'ютерної інженерії та програмування, Харків, Україна;

Dmytro Lysytsia – Candidate of Technical Sciences, Associate Professor, National Technical University "Kharkiv Polytechnic Institute", Associate Professor of the Computer Engineering and Programming Department, Kharkiv, Ukraine;

e-mail: Dmytro.Lysytsia@khpi.edu.ua

ORCID ID: <https://orcid.org/0000-0003-1778-4676>

Scopus ID: <https://www.scopus.com/authid/detail.uri?authorId=57220049627>

Левченко Андрій Олександрович – кандидат технічних наук, доцент, Військова академія, професор кафедри застосування підрозділів військової розвідки та Сил спеціальних операцій, Одеса, Україна;

Andrii Levchenko – Candidate of Technical Sciences, Associate Professor, Odessa Military Academy, Professor of the Training of Intelligence Units and Special Operations Forces Department, Odessa, Ukraine;

e-mail: katyaandreylev@gmail.com

ORCID ID: <https://orcid.org/0000-0001-5550-0027>

Scopus ID: <https://www.scopus.com/authid/detail.uri?authorId=57220805603>

ОЦІНЮВАННЯ ПРОДУКТИВНОСТІ ЗАВАНТАЖЕННЯ СЦЕНИ WEBAR ЗАСТОСУНКУ НА МОБІЛЬНОМУ ПРИСТРОЇ

Предметом дослідження є закономірності розподілу обчислювальних ресурсів центрального процесора та зміни часових метрик браузерного рушія під час виконання етапів ініціалізації WebAR-сцени. **Об'єктом даного дослідження** є процес ініціалізації та завантаження 3D-контенту у клієнтському WebAR-застосунку на мобільному пристрої. **Метою роботи** є розробка підходу до комплексного оцінювання продуктивності та встановлення закономірностей розподілу обчислювального навантаження під час ініціалізації WebAR-застосунків з використанням інструментів глибокого браузерного профайлінгу. **Завдання.** Розробити математичну модель функціонування WebAR-застосунку на етапі

ініціалізації сцени та декомпонувати процес завантаження на ключові етапи з оцінкою їхніх часових витрат. Провести емпіричне профілювання, ідентифікувати пікові навантаження на CPU та блокуючі задачі, а також сформулювати обґрунтовані напрями подальшої оптимізації продуктивності. **Методи.** Дослідження продуктивності проведено шляхом емпіричного тестування продуктивності WebAR-застосунку під час найбільш ресурсоемного етапу – ініціалізації та завантаження тривимірної сцени. Збір даних щодо навантаження на потоки CPU в контексті веб-сторінки здійснювався за допомогою інструментів розробника Chrome DevTools. Апаратним середовищем для експериментів слугував мобільний пристрій Google Pixel 8, а статистична достовірність результатів забезпечувалася 10-кратною ітерацією кожного тестового сценарію. **Результат дослідження.** Розроблена математична модель формалізує процес завантаження сцени WebAR-застосунку та дозволяє провести оцінку завантаженості CPU на кожному етапі. За результатом аналізу часового профілю, отриманого з використанням удосконаленої математичної моделі та засобів Chrome DevTools виокремлено чотири етапи завантаження: ініціалізація порожньої сцени, завантаження 3D-об'єкта із сервера, парсинг моделі та її рендеринг. Встановлено нелінійний характер навантаження на процесор. Найбільш ресурсоемним виявився етап парсингу, під час якого зафіксовано пікове завантаження CPU і найдовша блокуюча задача, спричинена синхронною десеріалізацією бінарних даних. **Висновки.** Запропонований підхід до підвищення продуктивності завантаження сцени WebAR-застосунку на мобільному пристрої шляхом балансування навантаження CPU. Також дослідження підтвердило, що критичною фазою ініціалізації є парсинг 3D-моделі. З огляду на результати моделювання та проведених емпіричних досліджень, перспективним напрямом є впровадження поетапного завантаження та обробки 3D-контенту.

Ключові слова: web; augmented reality; WebAR-застосунок; performance evaluation; парсинг; рендеринг; Angular.

Бібліографічні описи / Bibliographic descriptions

Кучук Г. А., Матвєєв М. І., Косенко Н. В., Лисиця Д. О., Левченко А. О. Оцінювання продуктивності завантаження сцени WebAR-застосунку на мобільному пристрої. *Сучасний стан наукових досліджень та технологій в промисловості*. 2026. № 2 (36). С. 52–69. DOI: <https://doi.org/10.30837/2522-9818.2026.2.052>

Kuchuk, H., Matvieiev, M., Kosenko, N., Lysytsia, D., Levchenko, A. (2026), "Evaluating the performance of a webar application's scene loading on a mobile device", *Innovative Technologies and Scientific Solutions for Industries*, No. 2 (36), P. 52–69. DOI: <https://doi.org/10.30837/2522-9818.2026.2.052>

Mykyta Lapin, Yurii Parzhyn, Kostiantyn Bokhan, Kyrylo Perevoznyk, Tetiana Aleksandrova

INVARIANT STRUCTURAL LEARNING: CONCEPT FORMATION AS HYPERGRAPH ATTRACTOR DYNAMICS

Subject of study. The subject of study is the formation of object-class concepts in machine-learning systems as a process of structural and parametric reduction of hypergraph representations, rather than as the optimization of a loss function. **Objective.** To construct an alternative learning framework for artificial neural networks – without an error functional (back-propagation) – that realizes invariant structural learning, where a class concept is defined as a structural attractor (the fixed point of a monotone reduction operator on a partially ordered space of hypergraphs), and to validate this theory on the recognition of handwritten digits. **Objectives.** 1) Formalize the framework with axioms of segmentation stability, strict reductivity, positive-only training, and locality of attention; 2) prove convergence and uniqueness of the attractor; 3) establish order-invariance of learning; 4) decompose the attractor into structural and parametric levels; 5) evaluate the transfer on a complete-contour subset of MNIST. **Methods.** Hypergraphs are obtained by skeletonizing binary contours (Growing Neural Gas + Ramer–Douglas–Peucker); concept attractors form via graph reduction over critical-point anchors. Classification uses a graph-edit-distance comparator with property costs and a complexity-adjusted log prior. The training set consists of 76 hand-picked MNIST originals, augmented (rotation $\pm 10^\circ$, shift $\pm 10\%$) to 805 instances. **Results.** The transfer learns 13 concept-attractors with node counts ranging from 3 to 15. Of 8,707 admissible complete-contour MNIST images, 8,685 yielded valid skeletal graphs; on this valid set, the transfer achieves an accuracy of 85.80%, weighted precision of 89.31%, recall of 85.80%, and an F1-score of 86.66%. The error structure is interpretable per concept – every misclassification traces back to a named attractor and feature range. **Conclusions.** The reported performance, achieved without a loss function and using three to nine originals per concept-attractor, supports the claim that learning consists of constructing a structural attractor rather than minimizing an error functional. The framework offers a principled approach to interpretable few-shot classification and a constructive alternative to gradient-based learning.

Keywords: invariant structural learning; structural attractor; concept formation; graph edit distance; few-shot recognition; explainable AI; MNIST; hypergraph reduction.

1. Introduction

Modern machine learning methods, and neural network models in particular, rely heavily on the optimization of loss functions and statistical generalization. These approaches involve an external quality criterion, and the learning process is formalized as the minimization of the discrepancy between the prediction and a given target value. Despite the empirical success of statistical learning, this paradigm has fundamental limitations: dependence on the choice of loss function, sensitivity to data distribution, and the need for large training samples.

The relevance of this issue is heightened by three concurrent factors.

First, the EU Regulation on Artificial Intelligence (Regulation 2024/1689) [1], together with Article 22 of the GDPR [2], establishes a legally binding requirement: automated decisions in regulated areas must be accompanied by an explanation that corresponds to the model's internal logic – a standard that post-hoc surrogates fail to meet in practice.

Second, the energy and computational costs of training large statistical models are growing faster than the volume of available labeled datasets, creating economic pressure in favor of methods capable of learning from a handful of examples.

Third, industrial areas with a chronic shortage of labeled data – medical OCR, defect detection in contour structures, and low-volume document processing – remain underserved by gradient-based classifiers. These three factors motivate the search for a non-statistical alternative, developed below.

An alternative approach involves viewing learning as a process of endogenous structural self-organization, in which the system does not minimize external error but instead converges to an internally consistent state. This paper proposes a formalization of this view, which extends the "information architecture" approach [3] based on hypergraph representations and dynamic structural reduction; this approach is called Invariant-Structural Learning.

The central idea is that each observed object is represented by a hypergraph, and learning involves

the sequential accumulation and alignment of such hypergraphs, followed by reduction. The alignment procedure maps elements from different hypergraphs into a common structural space, forming a cumulative hypergraph in which one can analyze the frequencies of element occurrence and their stability under the effect of reduction.

The key concept of the proposed theory is the structural attractor: an invariant structure toward which the reduction process converges. Unlike classical attractors of dynamical systems, this object is defined not through trajectories in continuous space, but through the intersection of coordinated structural invariants. It is shown below that the attractor admits an equivalent characterization as the set of elements whose frequency in the cumulative carrier is equal to one – that is, those that are always present – and this combines a dynamic interpretation with a combinatorial one.

The proposed approach differs from traditional learning methods in four fundamental ways:

- (i) there is no explicit error function and no procedure for minimizing it;
- (ii) learning is interpreted as the reduction and alignment of structures;
- (iii) invariants are formed endogenously, without external templates;
- (iv) invariance with respect to the order of example presentation is achieved.

Mathematically, the model relies on discrete structures (hypergraphs), partially ordered sets, and monotonic reduction operators that converge to fixed points.

The main contributions of this work are as follows.

1. A formal hypergraph model of learning via structural reduction is proposed.
2. The concept of a structural attractor is introduced, and it is shown that it coincides with the intersection of class invariants and the single frequency condition.
3. Theorems regarding the convergence of the reduction process and the uniqueness of the attractor are proven.
4. The invariance of the learning result with respect to the order of example presentation is established.
5. It is shown that the attractor naturally decomposes into topological (structural) and parametric levels.

Overall, the work proposes an alternative theoretical framework for machine learning, in which structural invariants and attractor dynamics play a central role, rather than error optimization. This opens the way to

constructing models focused on stability, interpretability, and learning with a limited number of examples.

Empirical validation of the proposed framework on a subset of the MNIST dataset containing closed contours is presented in Chapter 5. The concept-attractor classifier based on graph editing distance trains 13 attractors from 76 selected originals – ranging from three to nine per concept – expanded by parametric augmentation (random rotation within $\pm 10^\circ$ and translation within $\pm 10\%$ of the canvas) to 805 training instances. Of the 8,707 valid images, 8,685 formed valid skeleton graphs; on this valid set, the system achieves 85.80% accuracy, 89.31% weighted precision, 85.80% recall, and 86.66% F1-score, demonstrating the low-variance behavior predicted by the "structure-parameter" decomposition theorems in Section 4.1 and the order-invariance consequence from Theorem 2.

2. Literature Review and Problem Statement

Traditional approaches to handwritten digit recognition on MNIST [4] establish a clear benchmark against which any new method should be evaluated. Convolutional neural networks – from the canonical LeNet-5 architecture [4] to modern residual variants – consistently exceed 99% test accuracy on the full training set of 60,000 images. The Gaussian kernel support vector machine method achieves approximately 98.6% on the standardized MNIST; on raw pixels without augmentation, the result is more modest. Multilayer perceptrons with moderately wide hidden layers achieve 97–98% when trained with deslanting and affine augmentation [4]. A common feature of these methods is the need for tens of thousands of labeled examples and the absence of any built-in explanation mechanism: the decision of a convolutional network is based on linear combinations of feature maps, the interpretation of which by a human is possible only through post-hoc visualization surrogates.

A separate line of research focuses on the few-shot setting, where only a small number of labeled examples per class are available. Prototypical Networks [5] construct a class prototype as the average embedding of support examples in a latent space, trained via episodic optimization. Matching Networks [6] use an attention mechanism to weight the contribution of each support example, while Model-Agnostic Meta-Learning [7] seeks a parameter initialization from which the model adapts to a new task in a few gradient steps.

These methods are typically evaluated on the episodic protocol on Omniglot and miniImageNet, rather than on the open MNIST protocol; therefore, a direct numerical comparison with the 85.80% accuracy presented here is best interpreted as a conceptual positioning within the niche of explainable few-shot classifiers, rather than as a direct benchmark. A common limitation of this line of research is that, despite reducing the need for labeled data during adaptation, these methods do not eliminate the dependence on a large base corpus during pre-training, and the prototype in Prototypical Networks remains an opaque vector in a multidimensional space rather than an interpretable object suitable for expert verification.

The dominant approach to achieving interpretability for pre-trained black-box models has been the construction of post-hoc surrogates. Minh et al. [8], Hooshyar and Yang [9], Rajabi and Etmiani [10], and Nawaz et al. [11] survey this landscape and classify methods by locality, model specificity, and type of approximation. The most common examples are LIME [12], which locally approximates a nonlinear model with a linear surrogate, and SHAP [13], based on Shapley value attribution. Slack et al. [14] experimentally demonstrated that LIME and SHAP are vulnerable to targeted attacks: a classifier with distinctly discriminatory behavior can be made to appear neutral under both methods. Post-hoc explanations therefore offer no guarantee of correspondence to the model's actual logic, which motivates interest in models that are explainable by design – that is, models in which the internal representation itself is an interpretable description.

Graphs as image representations have a long history in computer vision: skeleton graphs preserve the topology and basic geometry of an object while discarding redundant pixels, and demonstrate the competitiveness of graph representations in pattern recognition tasks. The current trend in graph neural networks, particularly Graph Prototypical Networks [15], extends few-shot learning to attributed graphs, but reverts to latent-vector representations and does not preserve a directly inspectable representation of the concept. In a previous study by the same team [16], the line of interpretable graph attractors was tested on a restricted six-class MNIST dataset (digits 1, 2, 3, 6, 7, 9; 5,467 test images, 8 attractors) and achieved 82.35% accuracy with exact calculation of graph editing distance under a 60-second comparison budget. This work extends the methodology to the full ten-class MNIST alphabet with contour integrity filtering, expands the alphabet to 13 attractors,

introduces diagnostic feature weighting based on the width of the reduced range, and transitions to an iterative upper bound for graph editing with progressive refinement under a significantly smaller time budget – a combination that yields the reported 85.80% accuracy from Section 5.

The matching of attributed graphs formally boils down to finding the minimum sequence of editing operations (insertion, deletion, and replacement of vertices and edges) that transforms one graph into another. Sanfeliu and Fu [17] introduced the concept of edit distance for attributed relational graphs and established it as a similarity metric for pattern recognition. Conte et al. [18] summarized three decades of graph matching methods and highlighted the trade-off between exact and approximate computation; the exact graph edit distance remains NP-hard. Riesen and Bunke [19] proposed an efficient binary approximation with assignment, yielding a suboptimal upper bound in polynomial time, enabling the use of the metric on graphs of practically useful sizes at the cost of optimality. More recent neural variants replace explicit edit costs with learned embeddings: Piao et al. [20] compute graph edit distance via neural graph alignment, predicting the approximate distance from training pairs; Moscatelli et al. [21] formulate the basic node matching problem with an explicit emphasis on metric properties and assignment interpretability; Xie et al. [22] extend the paradigm to federated few-shot graph classification. An additional direction within the same class of problems is the reconstruction of incomplete spatial graphs using graph neural networks; an example is the recent work [23] on completing 3D point cloud models in the same journal. These methods significantly improve computational efficiency and scalability; however, most replace interpretable elementary editing costs with latent embeddings, reducing the readability of decisions at the level of "which vertex was matched with which and with what penalty". This leaves room for architectures in which the operations themselves remain interpretable by design – this is the motivation behind the proposed approach.

The proposed framework is based on four well-established lines of evidence from cognitive science. The binding problem – how the brain combines distributed representations of attributes into coherent perceptual wholes – was formalized by von der Malsburg [24] and remains a central problem in neural coding, with extensions in the form of neural

syntax [25]; the structural attractor model solves it directly by encoding attributive relations as anchor links (hyperedges) within a single representation. The literature on active perception, originating with Yarbus [25] and developed by Bajcsy et al. [26] and Rucci and Victor [27], shows that the trajectory of a saccade is determined by the cognitive task no less than by the image itself, which motivates the selection of anchors based on attention rather than exhaustive global alignment. Single-unit recordings [28, 29] have revealed concept-cell responses invariant to surface transformations of the stimulus, empirically supporting the existence of class-level invariant representations postulated in Definition 3 (concept). The dendritic computation hypothesis [30, 31, 32, 33, 34] asserts that a neuron's dendritic tree performs substantive logical computation rather than acting as a passive integrator, providing a biological substrate in which the algebra of structural reduction can be physically realized. These four lines of reasoning together justify the non-statistical orientation of the proposed framework and motivate the architectural decisions in Section 4.

To summarize this review, three desirable properties stand out, none of which currently have a satisfactory joint solution:

- a) training without backpropagation and without a large labeled corpus;
- b) operation with a small number of examples per class;
- c) local explainability of each decision at the model level.

Classical convolutional and support vector models provide only accuracy; few-shot meta-learning methods provide adaptation with a small training set but remain opaque; post-hoc methods of explainable AI provide an approximation of behavior rather than the logic itself. The graph attractor architecture proposed in this work, trained without backpropagation, formally satisfies all three properties simultaneously – and this defines the unsolved part of the problem and justifies the relevance of this research.

3. Purpose and Objectives of the Study

The aim of this study is to empirically validate the foundation of Invariant-Structural Learning, formalized in Section 4.1, and in particular Theorems 1–3, by constructing an end-to-end classifier of concept attractors and evaluating it on a subset of MNIST with intact digit

contours to demonstrate that few-shot learning of structural attractors without backpropagation can simultaneously achieve competitive accuracy and local interpretability at the decision level. The empirical task is intentionally limited to images with intact digit contours, which corresponds to a psychophysiologically grounded baseline recognition mode based on a complete structural description; images with broken or fragmented contours are excluded from this study, as their recognition requires an additional associative completion mechanism of a different mathematical nature, which warrants separate consideration.

To achieve this goal, this paper addresses the following five tasks.

1. Formalize the "bitmap $\rightarrow$ attributed graph" transformation transfer as a constructive implementation of the basis representation: transform each MNIST raster into an attributed graph that preserves the object's skeleton along with the coordinates of critical (anchor) points and the geometric parameters of contour segments.

2. Implement the concept-attractor formation operator as an iterative reduction of a set of graph examples to a stable generalized structure that realizes $R = C \circ S \circ R_p$ (Section 4.1.3 of the theoretical framework), and verify the parametric stratification predicted by augmentation theory through controlled augmentation by random rotation within $\pm 10^\circ$ and translation within $\pm 10\%$ of the canvas.

3. Implement the classification step as a comparison based on graph editing distance between the invisible image and each concept-attractor, with diagnostically weighted costs for node replacement across fourteen structural and geometric features and a graded cost scale for edge editing operations, returning the winning attractor along with a similarity score and the logarithmic prior complexity used for its selection.

4. Conduct an experimental evaluation of the proposed classifier on the MNIST dataset with intact contours (8,707 valid images, 13 attractor concepts, 10 classes) and report accuracy, weighted precision, recall, and F1-score along with per-concept metrics, top mismatch pairs, and the error structure by attractor name and feature range.

5. Compare the obtained metrics with published results from major classes of competing methods – classical statistical baselines, few-shot meta-learning methods, and the team's previous six-class study – and discuss the position of this work within the niche of explainable few-shot classifiers.

4. Materials and Methods

4.1. Theoretical Framework: Invariant-Structural Learning

Section 4 is divided into two parts. Subsection 4.1 presents the theoretical foundation of invariant-structural learning over hypergraph representations, defining the base spaces, the reduction operator, the learning dynamics, and the central convergence and uniqueness theorems. Subsection 4.2 describes an experimental implementation that validates the framework on a subset of MNIST with connected contours, including a "bitmap $\rightarrow$ graph" transformation transfer, a concept-attractor formation operator, a graph editing distance comparator, diagnostic feature weighting, and an example of attractor formation for the digit 7.

4.1.1. Basis Spaces and Structures

The input data space X consists of original objects (e.g., MNIST images, geometric shapes, contours, 2D and 3D models). The space P of atomic structural primitives contains the basic building blocks of each representation; for MNIST, these are line segments approximating the contour and contour branching points: endpoints, corner points, intersection points, and convergence points of line segments. The hypergraph space Z formalizes the representation of a single detector neuron. An element of Z is a hypergraph H defined by relation (1), where V is a finite set of vertices – structural primitives, E_s are spatiotemporal edges encoding basic connectivity, and E_p are parametric hyperedges encoding the relationships between sets of parameters of different vertices. The hypergraph is constructed as a layer over the underlying spatial graph; each vertex carries its own set of parameters, so the parametric space is stratified over the structure.

$$H = (V, E_s, E_p) \quad (1)$$

4.1.2. Extraction and measurement

The primitive detector D simulates the operation of a sensory system: it takes the original input and generates structural primitives. Its implementation is not the subject of this study – D is considered to be given [34]. Each measurement function m_i returns the i -th parameter of a structural element, and m collects all the parameters of the graph. Measurement induces parametric modes for each vertex; each vertex is associated with its modal group – a set of parameters measured for it.

The resulting multidimensional parametric space is defined exogenously via coordinate systems and measurement scales; the segmentation of these parametric spaces during training induces a metric. Examples for MNIST skeletonized contours include segment length, segment orientation, the angle between segments, and point coordinates.

4.1.3. Reduction

R_p is an ordering/compression operator on metrics.

It does not alter the spatial structure of the hypergraph but acts on the space of parametric relations: continuous values are replaced by discrete classes (segments). The clustering/quantization operator Q with stability parameter ε partitions the parametric space into stable segments σ_i . Parameters may be reduced to qualitative categories or disappear entirely from the hypergraph. The metric d_i in the parametric space determines the segmentation structure. Axiom 1 (segmentation stability) states that the segmentation is idempotent: $Q \circ Q = Q$ on the parametric space for a chosen ε . The structural reduction R_s is performed in three stages.

First, the frequency measure f_x of an element x is defined as the proportion of class examples in which x is present after alignment, where $f_x \in [0,1]$ and N are the number of accumulated class examples.

Second, the structural selection operator S retains only those elements that are reduction-invariant – those that appear (almost) in all class examples. The set A of such elements (anchors) is dynamically updated with each new instance. If the training set contains instances with an incomplete attractor structure, the strict equality $f_x = 1$ is relaxed to $f_x \geq \theta$ with θ – a hyperparameter that typically belongs to $(0.7,1]$; for clarity of presentation, the theory below assumes $\theta = 1$.

Third, the closure/connectivity restoration operator C reconstructs a single hypergraph from the surviving elements (anchors), preserving spatial connectivity.

$$f_x = \frac{1}{N} \sum_{i=1}^N 1[x \in A_i] \quad (2)$$

Definition 1 (anchor). An element $x \in V$ is called an anchor if its frequency measure satisfies $f_x = 1$ (equivalent to $f_x \geq \theta$ for the weakened regime). An anchor is a structural element (or parameter value) that is consistently present in (almost) every instance of the

class. It serves as a fixed point around which hypergraphs are aligned; anchors are not removed by reduction and are restored by the operator C to preserve the global structure after removing non-critical vertices ("attractor noise"). For MNIST, anchors are critical branching points of the contour: endpoints and corner points.

The complete reduction operator R , defined by relation (3), is applied first parametrically, then structurally. Axiom 2 (strict reductivity and minimality) states that under the action of R , the size of the hypergraph (in terms of the number of elements and parameters) does not increase, and the resulting attractor is minimal: it contains no elements from $f_x < 1$.

$$R = C \circ S \circ R_p \quad (3)$$

4.1.4. Learning dynamics – cumulative, invariant with respect to order

For a new positive example x' , the system constructs the following objects. The anchor alignment μ is a partial mapping $\mu: A' \rightarrow A_t$ between the anchors of the new example and the accumulated anchors. The family μ is consistent if the sequential alignment of anchors between examples is order-independent, i.e., there are no conflicting identifications of the same anchor. For cumulative aggregation, given a preserved hypergraph H_t and a new instance H' , vertex aggregation identifies vertices via μ ; spatial and parametric edges are transferred by the operator μ . The merged hypergraph $H_t \oplus_\mu H'$ merges the identified anchors and adds the remaining elements. The training iteration is defined by relation (4): the operator $\oplus$ expands the structure, while R reduces it. Given a consistent μ and a fixed selection rule μ , the result of the cumulative process is associative and commutative up to structural equivalence E : two hypergraphs are equivalent if their segmented parameters coincide, and the minimal structures are locally isomorphic around common anchors. Class-detector neurons are not required to construct hypersurfaces separating different classes; instead, an invariant hypergraph is constructed independently for each class. The algebraic reduction of hypergraphs allows for an interpretation as the structural plasticity of active dendrites, where patterns of synaptic inputs on dendritic branches form complex logical elements encoding structural invariants [35, 36]

$$H_{t+1} = R(H_t \oplus_\mu H'). \quad (4)$$

Axiom 3 (learning from positive examples). A class detector neuron is trained only on positive examples of its class. When an input arrives, the system does not "compute" a response but evolves until it reaches a stable state; its dynamics depend on the measurement of its own structural components, modeling internal consistency without an external criterion of optimality. All subsequent statements are derived from the axioms above.

4.1.5. The Principle of Invariant Reduction

The system evolves toward reducing structural redundancy while preserving invariants until it reaches a minimal fixed structure. This principle does not rely on gradient optimization or classical variational principles. Instead, the learning is guided by a monotonic reduction operator on a partially ordered space of structures, which converges to minimal invariant fixed points. A unified hierarchy of coordinate system and measurement scale segmentations is introduced, $\sigma^{(k)}$, which is specified exogenously (modeling the self-organization of segments in biological neural structures); a higher k denotes coarser ("more general") segments. The transition operators between levels are quantization operators at each level. A partial order $|$ is defined on this space (Davey, Priestley, 2002): it is reflexive and transitive (transitivity follows from the nesting and the order on the segmentation hierarchy). The invariance functional Inv induces a closure analogous to that in Formal Concept Analysis [37]. Reduction operators are expressed in terms of segments: parametric reduction moves to a coarser level; structural reduction is the composition $C \circ S \circ R_p$. Classical neural networks lose the topological structure of the image, whereas the proposed framework solves the classical neurobiological problem of binding by preserving anchor connections (hyperedges) within the attractor [24]. The discrete invariance gradient – geometrically a vector in the segmentation hierarchy – specifies the direction of generalization growth. Invariants can be added or removed (false), but true invariants are preserved: the dynamics are not monotonic, but the sequence of invariant sets converges to a stable structural core E^* . Extremality is induced by the operator and the order themselves, rather than by an external optimality functional – unlike loss functions in statistical neural networks (Buzsáki, 2010).

4.1.6. Structural Attractor and Concept

The central object is the structural attractor: an invariant structure that is a fixed point of the reduction operator and defines the internal model of the class.

Definition 2 (structural attractor). Let K be a class of objects, Z be a hypergraph space, and R be a reduction operator.

$$R(C_K) \cong C_K, C_K \preceq C \text{ whenever } R(C) \cong C, \exists t < \infty : R^t(H) \cong C_K \quad (6)$$

The attractor C_K is a fixed point R , the minimal carrier of the invariant structure, the boundary of the training process, and the canonical form of the class's objects. Therefore, training is the reduction of various representations of an object to a single structural normal form.

Definition 3 (concept). A concept is a structural attractor of a perceived holistic image, that is, its internal model. Therefore, a concept is an endogenous ontological object for the system. Models with memory of structural attractors form a separate class of self-organized dynamic systems [38, 39]; biological neurons and neural structures are one example. The proposed model suggests a neurobiological interpretation in which

- a) the attractor corresponds to a stable activation pattern;
- b) anchor elements are realized as stable synaptic configurations;
- c) reduction reflects the selection of stable structures in dendritic trees.

This interpretation is illustrative, not a direct biological statement. Self-organization here refers to a modeled process, not a physical one; in particular, feature detectors, coordinate systems, and scales are specified exogenously, whereas reduction operators model the energy dynamics of natural self-organization.

Theorem 1 (Convergence of R)**Conditions:**

- 1) Z is a partially ordered hypergraph space;
- 2) $R: Z \rightarrow Z$ satisfies Axioms 1–2 (segmentation stability, strict reductibility, and minimality);
- 3) the complexity functional Φ , defined by relation (5), takes integer values and is bounded from below by 0.

Proof outline.

Step 1 (decreasing complexity): at every non-stationary iteration, $\Phi(R(H)) < \Phi(H)$ by Axiom 2.

A structural attractor of the class K is a hypergraph C_K that satisfies:

- (i) invariance – $R(C_K) \cong C_K$;
- (ii) structural minimality – for every $C \in Z$ in $R(C) \cong C, C_K \mid C$ holds;
- (iii) attraction – for every representation $H \in Z$ of the class K , there exists a finite t such that $R^t(H) \cong C_K$.

Step 2 (finiteness): a strictly decreasing sequence of natural numbers is finite.

Step 3 (stabilization): the sequence $(R^t H)_{t \geq 0}$ reaches a fixed point in a finite number of steps.

Conclusion: there exists an C_∞ such that $C_\infty = R(C_\infty)$, where $t^* \leq \Phi(H_0)$ limits the number of steps to convergence by the initial complexity. Corollary 1 (convergence to a class attractor). If a unique structural attractor C_K is induced by the class K , then for every $H \in Z$ of the class K there exists a finite t such that $R^t(H) \cong C_K$.

$$\Phi(H) = |V| + \lambda |E| \quad (5)$$

Theorem 2 (Uniqueness of the C_K).**Conditions:**

- 1) all examples of the class are anchor-connected, i.e., for any pair (H_i, H_j) of examples of the class, there exists a chain of common anchors connecting them μ -consistent mappings;
- 2) R is monotonic and minimal (Axioms 1–2);
- 3) the dynamics of cumulative composition is order-independent.

Proof outline.

Step 1 (convergence): by Theorem 1, $R^t(H)$ converges to a fixed point.

Step 2 (order independence): the frequency function f_x depends only on the set of accumulated examples, not on the order of their presentation; therefore, the set $x: f_x = 1$ is order-independent.

Step 3 (formalization of the fixed point): by the selection axiom, the set of vertices of the fixed point coincides with the set of elements from $f_x = 1$; the complete attractor is the closure of this set of vertices.

Step 4 (uniqueness): anchor connectivity guarantees the consistency of the μ mapping across the entire class (analogous to local-global consistency in bundle theory [40]), so the resulting structure C_K is unique up to E .

Conclusion: there exists a unique C_K such that $C_K = R(C_K)$ and C_K are independent of the order of presentation.

$$\exists! C_K : C_K = R(C_K) \quad (7)$$

4.1.7. Few-shot learning

Consider the subset $T \subset K$. The coverage condition requires that every element of the attractor be present in at least one instance of T . Definition 4 (sufficient covering set): T is a sufficient covering set if (1) the coverage condition holds and (2) anchor connectivity is preserved on T – the hypergraph of the intersection of anchors T is connected, so all examples from T can be matched via common anchors. Corollary 2: If T is a sufficient covering set, then all elements of the attractor survive reduction to T , and the attractor obtained from T coincides with C_K . From this follows the property of embeddability: for every $x \in K$ there exists a mapping $\iota: C_K \rightarrow H(x)$ that preserves structural relations – the attractor is embedded in every instance. Thus, an attractor can be defined not by the entire class, but by a small subset of its carriers, which corresponds to the psychology of human perception and opens the way to unsupervised learning. A sufficient covering set ensures the reconstructability of the attractor's structure, but, in general, does not guarantee its uniqueness. Therefore, the stronger concept of an identifying set is introduced. Definition 5 (identifying set): T is an identifying set if (i) it covers all elements of the attractor; (ii) it covers all anchor relations (edges); (iii) anchor connectivity holds. T is then a minimal cover: it induces a unique fixed point of the operator R . Every identifying set is a sufficient cover; the converse generally does not hold. Corollary 3 (Reconstruction): If an attractor is embedded in every instance of the class and there exists an identifying set T^* , then the attractor obtained from T^* is equal to C_K up to E . It is not the cardinality T that matters, but its covering power. Two or three examples may be sufficient, whereas a hundred may not. Given a finite number of anchors and finite connectivity, there exists a finite minimal subset of examples that reconstructs the attractor of the class.

4.1.8. The Role of Augmentation

Augmentation is interpreted as an operator for parametric analysis of structure. Given the example $H_x = (V, E_s, E_p)$, augmentation generates the family $H_{x_i}^{(i)}$ with structures that are equivalent up to isomorphism and with continuously varying parameters. Augmentation constructs the segmentation σ_i from observed values: it is assumed that parametric variations do not depend on the structure – variations obtained from a single example are representative of the entire class in terms of parameters. Therefore, training is divided into two levels.

1. Structural level: from a small subset T , the structural skeleton is reconstructed – anchors and their relationships.

2. Parametric level: augmentation generates parametric stratification – for each anchor, an interval/segment of valid values. The full attractor combines structural coverage (the small T) and parametric coverage (augmentation). Unlike classical models, where augmentation increases the sample size to reduce overfitting, here augmentation constructs the parametric structure (and thus the metric) through segmentation. Structural examples define the topology; augmentation defines the geometry (metric). Thus, augmentation acts as a metrization operator. The induced metric on the attractor is defined locally for each parameter via its segmentation σ_i , and globally via aggregation by anchors.

Properties of this metric:

- 1) it is induced by segmentation;
- 2) it respects invariance – invariants contribute zero to the variation;
- 3) it is consistent – defined only on anchors and consistent via μ .

Therefore, this metric is not external (exogenous) but arises endogenously: structure $\rightarrow$ segmentation $\rightarrow$ metric. This is the opposite of metric learning, contrastive learning, and kernel methods.

4.1.9. The Attention Mechanism and Anchors

The most challenging practical problem in structural learning is the alignment of hypergraphs. Here, it is solved using an attention operator over anchor structural points.

Definition 6 (attention operator). The **attention operator** F is a local selector that chooses a subset $F(H) \subset V$ to construct a matching based on local

information – for example, the "capture" points of the MNIST contour. A typical example is given by equation (8), where I is a measure of local information, and τ is a threshold. The operator selects elements with extreme parameter values for a given structure or with rarity/contrast (points of maximum informativeness). Instead of solving the global hypergraph isomorphism problem (which is based on graph editing distance [20]), the attention operator restricts the alignment to the set $F(H)$, transforming the problem from a global combinatorial optimization problem into a local alignment problem and circumventing the NP-hardness of the exact graph edit distance [18]. The alignment μ is therefore defined only on $F(H)$.

$$F(H) = \{x \in V : I(x) \geq \tau\} \quad (8)$$

Axiom 4 (locality of attention and consistency anchors)

1. The attention operator F is local – it depends only on H , restricted to $F(H)$.

2. Anchors are invariant under reduction: $F(R(H)) = R(F(H))$ up to structural equivalence.

The attention operator depends on task-induced top-down modulation; it models the brain's active perception process [26, 27, 28]. Yarbus experimentally demonstrated that the trajectory of saccades is determined not only by image properties but also by the cognitive task; in our context, this confirms that the selection of anchor points is a goal-directed action F that minimizes perceptual redundancy.

4.1.10. Neural Map

The class " K " is represented by a self-organizing attractor map (9) – a finite set of structural attractors (detector neurons) – corresponding to subclasses or stable variations within the class. Each attractor of a subclass $C_K^{(j)}$ satisfies $C_K | C_K^{(j)}$: the class attractor is a minimal invariant substructure embedded in every subclass attractor. The pseudo-distance d on the structural space is bounded from below between different attractors – see (10) – ensuring structural distinguishability. The class map has a strict hierarchy: the central neuron is the conceptual class detector; d measures the edit distance between this central hypergraph and the hypergraphs of subclass detectors or

memorized individual examples. Subclass attractors are formed under specific novelty conditions (their formulation is beyond the scope of this work [41]).

$$M_K = C_K^{(j)}_{j=1,\dots,n}, n < \infty \quad (9)$$

$$d(C_K^{(i)}, C_K^{(j)}) \geq \delta > 0 (i \neq j) \quad (10)$$

Theorem 3 (Self-organization of the attractor map) Conditions:

1) The reduction operator R satisfies the axioms (monotonicity, invariant selection, and connectedness closure);

2) the learning dynamics is given by $H_{t+1} = R(H_t \oplus_\mu x_{t+1})$ and for each class K , according to Theorem 2, there exists a structural attractor C_K ;

3) a pseudometric d is defined on the structural space, distinguishing structures with accuracy up to ϵ ;

4) a separability threshold $\delta > 0$ is fixed.

Conclusions: During the learning process, a finite map $M_K = C_K^{(j)}_{j=1,\dots,n}$ is endogenously formed with the following properties:

(i) structural embeddability ($C_K | C_K^{(j)}$ for every j);

(ii) separability ($d(C_K^{(i)}, C_K^{(j)}) \geq \delta$ for $i \neq j$);

(iii) attraction (every example of the class converges under the action of R to some $C_K^{(j)}$); (iv) self-organization (M_K is not specified a priori, but arises as a set of stable fixed points); (v) finiteness ($n < \infty$).

Proof outline.

Step 1 (convergence to fixed points): By Theorem 1, every learning trajectory stabilizes at a fixed point R . Step 2 (nature of fixed points): by the selection axiom, every limit structure is equal to a closure $x : f_x = 1$ on its trajectory; different subsets of examples can induce different attractors.

Step 3 (emergence of new attractors): a new example x' with an anchor system A' , compatible with the current attractor $C_K^{(j)}$, yields the same attractor; otherwise R generates a new fixed point $t C_K^{(j+1)} \neq C_K^{(j)}$.

Step 4 (endogenous map formation): iterating over the sequence of examples yields a set of fixed points; no attractor is specified exogenously.

Step 5 (structural nesting): by Theorem 2, $C_K | C_K^{(j)}$ for every j (an attractor of the class is a common substructure).

Step 6 (separability): by the threshold condition $d \geq \delta$, attractors do not merge under the action of R .

Step 7 (finiteness): the structural space is finite due to the finiteness of the number of vertices and parameter segmentations; therefore, $n < \infty$.

Conclusion: the attractor map $M_K = C_K^{(j)}$ arises as a set of stable fixed points $(R, \oplus_\mu, d)$, and its structure is determined endogenously by these three operators, rather than by an external function. Therefore, the map is the result of a dynamic partition of the example space into basins of attraction.

4.1.11. Competition between attractors

The distance between attractors enables "winner-take-all" competition both within a single-class map and between maps of different classes. To achieve this, the response of each detector neuron must aggregate information about the structural and parametric composition of its attractor. This view, in our opinion, sheds light on the question of the neural code – in particular, on the problem of binding [24] and the formation of invariant conceptual representations [29, 30] – assuming that the basic element of the neural code is not the temporal frequency of spikes, but the topology of connections in the dendritic tree [33]. For a finite set of attractors, the distance d is a mixed pseudometric (11) that combines discrete structures and parameters, where d_s is the structural editing distance, and d_p is the parametric distance after alignment. The weights α, β are determined empirically and play a decisive role. According to Riesen [42], the structural distance is defined via partial alignment $\mu \in \Pi(C, C')$ as the sum of three terms – unaligned elements C , unaligned elements C' , and the size of the aligned substructure – which yields the number of elements that could not be matched. The parametric distance is computed after matching – either directly or via segments σ_i . The function d satisfies non-negativity, symmetry, and identity ($d(C, C) = 0$); $d(C_1, C_2) = 0$ implies $C_1 \cong C_2$ (a pseudometric, since the identification is up to E). Together, these properties formally define a metric between attractors that supports competition on a "winner-takes-all" basis.

$$d = \alpha d_s + \beta d_p \quad (11)$$

4.2. Experimental Implementation

Section 4.2 describes how each abstract construct from Section 4.1 – the dictionary of structural primitives,

the reduction operator "R", parametric augmentation stratification, and attention-guided anchor selection – is implemented in code and empirically validated. The theoretical foundation of Section 4.1 is formulated over hypergraphs $H = (V, E_s, E_p)$ with parametric hyperedges E_p that connect parametric sets of different vertices. In this implementation, the parametric layer is reduced to a set of attributes per vertex rather than to an explicit hyperedge: each vertex $v \in V$ carries a vector of fourteen continuous and categorical attributes (Table 1) instead of participating in a separately stored relation E_p . The base graph (V, E_s) is therefore a directed attribute graph in NetworkX, and the hyperedge layer emerges implicitly through the diagnostically weighted cost of node substitution, described later in this subsection, which links attributes between matched pairs of vertices at the time of comparison. This is a representative simplification of Section 4.1, not a theoretical deviation: hypergraph terminology is retained in the general formalism of Section 4.1, and thereafter, starting with Section 4.2, the term "graph" is used where this simplification is made.

The conversion of a raster image into an attributed graph is implemented as a three-stage transfer.

First, the image is binarized using a fixed intensity threshold that separates foreground pixels from the background.

Second, the Growing Neural Gas algorithm [43] adaptively places vertices along the object's median line, forming a topologically correct skeleton that respects the local distribution of foreground pixels.

Third, the Ramer–Douglas–Peucker algorithm [44] removes intermediate vertices whose deviation from the chord between neighbors does not exceed a specified tolerance, leaving only structurally significant points: endpoints, corner points, and convergence (branching) points.

The resulting representation is a simple directed attributed graph. Vertices of the first type – Point nodes – represent critical points of the skeleton and implement the V_{anchor} from Definition 1. Vertices of the second kind – Vector nodes – represent oriented segments between adjacent Point nodes and implement structural edges E_s , elevated to first-class objects (so that each segment can carry its own set of parameters, as required by the stratification of the parametric space over the structure from Section 4.1). The spatial coordinates of all vertices are normalized relative to the center of mass of

the foreground mask within a symmetric range, ensuring invariance under parallel translation in the image plane.

Each vertex carries a 14-element feature vector. The features are divided into four conceptual categories.

- (i) Spatial-geometric: *normalized_x*, *normalized_y* – horizontal and vertical coordinates relative to the centroid; *distance_to_centroid* – radial distance.
- (ii) Orientational, defined only for Vector nodes: *horizontal_direction*, *vertical_direction* – categorical orientation along the two coordinate axes; *angle_with_ox* – orientation angle in degrees (also defined on Point nodes as the angle of convergence of incident segments).
- (iii) Structural-functional, defined only for Point nodes: *is_endpoint*, *is_corner* – Boolean flags indicating whether the node terminates a single segment or marks a change in direction exceeding a fixed threshold, respectively; *junction_angle_min* – the smallest internal angle at

a branching point. (iv) Topological (per neighborhood): *length_ratio_to_max* – the length of a segment normalized by the longest segment in the graph; *eccentricity* – graph-theoretical eccentricity of a node (maximum distance in edges to any other node); *avg_neighbor_vector_length* – average normalized length of incident edges; *neighbor_endpoint_count*, *neighbor_junction_count* – number of neighboring Point nodes by role. The complete set is summarized in Table 1. Attributes defined only on one node type – for example, the orientation triplet on Vector or the structural-functional triplet on Point – are implicitly omitted during comparison when matching with a node of a different type, so the decomposition of values presented later in this subsection always operates on the intersection of features actually present on both sides of the candidate substitution.

Table 1. Feature vector of a node in the graph representation of an image, grouped by conceptual categories

Name	Description
Normalized <i>x</i>-coordinate	The horizontal coordinate of the node, normalized relative to the center of mass of the foreground mask; localizes the critical point of the skeleton along the horizontal axis of the image plane.
Normalized <i>y</i>-coordinate	The vertical coordinate of the node, normalized relative to the center of mass; localizes the critical point along the vertical axis.
Distance to the centroid	The normalized radial distance of the node from the center of mass; distinguishes between central and peripheral skeleton points.
Angle with the OX axis	The orientation angle of the segment in degrees for Vector nodes; the angle of convergence of incident segments for Point nodes; an invariant shape descriptor of the segment.
Endpoint flag	Boolean attribute; set to true if the node has degree 1 and terminates a single skeleton segment.
Corner point flag	Boolean attribute; positive when the node marks a change in direction exceeding a fixed angular threshold.
Minimum branching angle	The smallest internal angle at a branching point between converging segments; distinguishes between acute and obtuse convergence points in the skeleton.
Number of neighboring terminal points	The number of neighboring Point nodes marked as endpoints; characterizes the local density of dead-end branches near the node.
Number of neighboring branching points	The number of neighboring Point nodes marked as convergence points; indicates proximity to complex structural connections.
Average length of neighboring vectors	The average normalized length of segments incident to a Point node; characterizes the local geometric scale around the point.
Horizontal direction	A categorical attribute for Vector nodes; the orientation of the segment along the OX axis (left or right).
Vertical direction	A categorical attribute for Vector nodes; the orientation of the segment along the OY axis (up or down).
Length-to-maximum ratio	The ratio of a Vector node's length to the longest segment in the graph; distinguishes between main segments and short connecting segments.
Eccentricity in a graph	The maximum distance in the edges from a vertex to any other vertex in the graph; endpoints have high eccentricity, while central vertices have low eccentricity.

A concept attractor is formed by inductively and iteratively applying the reduction operator " $R = C \circ S \circ R_p$ " from Section 4.1 to a small set of skeleton example

graphs of the same class. At each step, the operator takes the current state of the concept and a fresh example graph and returns an updated concept that generalizes both inputs. The schematic operation of the operator is given

by the equation below: the new state of the concept is the result of the reduction of the pair (current state of the concept, new example).

$$C_{i+1} = R(C_i, G_{i+1}) \quad (12)$$

Generalization is performed on a per-component basis. For continuous attributes, a valid interval is constructed, bounded by the minimum and maximum values observed among the absorbed instances, and the typical value is stored at the center – this implements the parametric segmentation σ_i from Section 4.1.3. For categorical attributes, the intersection of valid values is stored. For list attributes (where present), set intersection is applied. Structural alignment of nodes between the current concept and the new example uses the same graph edit distance comparator as during classification, ensuring a common core of similarity for the training and classification procedures. After several examples, the iteration stops adding new vertices and only expands the intervals of existing ones, signaling the attainment of a fixed point of the operator R – a structural attractor C_K of the class K (Definition 2).

The role of augmentation in this work is to provide a practical implementation of Section 4.1.8 of the theoretical framework, rather than to address the classic issue of overfitting. Each selected original example H_x is expanded into a family $H_x^{(i)}$ of variants that are structurally equivalent up to anchor alignment and with continuously varying parameters. Variants are generated by an internal helper that rotates each original by an angle uniformly distributed over $[-10^\circ, +10^\circ]$ and shifts it by an offset uniformly distributed over $[-10\%, +10\%]$ of the canvas dimensions, with black fill outside the rotated mask. The number of variants per original ranges from five to ten depending on the concept; the total number of augmentations per concept is listed in Table 2.

This gives rise to three theoretical effects described in Section 4.1.8. First, the augmented family tightens the bounds of the admissible coordinate intervals at each anchor and narrows the estimates of the diagnostic weights (see below); without augmentation, the interval estimated from three to nine selected originals reduces to an almost degenerate width at most coordinates, and the diagnostic weight is uniformly saturated across all features. Second, the segmentation σ_i of each parameter is forced to refer to a stable interval rather than a single

point, which directly implements the metric-from-segmentation principle (structure $\rightarrow$ segmentation $\rightarrow$ metric): the width of the post-reduction interval governs each feature's contribution to the substitution cost, and a narrower width receives a higher weight during comparison. Third, the augmented family empirically ensures order invariance in training (Theorem 2): after merging the originals in any order with their variants, the resulting attractor is invariant under permutation of the merge order with accuracy up to ϵ , since the set of elements with frequency one on the cumulative carrier does not depend on the order. Therefore, the augmentation contract is not to "increase the training set", but to "make the metric measurable from the data".

Classifying an unseen image boils down to calculating the upper bound of the graph edit distance between the skeleton graph of the test image and each concept-attractor in the alphabet. The implementation uses the iterative upper bound generator "optimize_graph_edit_distance" from NetworkX [45], which is initialized with a polynomial approximation via the Biedel assignment [19] and monotonically refines the upper bound by iterating over alternative vertex mappings. After a fixed time budget – 15 seconds per concept comparison in our campaign – the procedure stops and returns the best upper bound obtained. Therefore, the returned value is a time-limited estimate: it does not guarantee the global minimum of the edit distance, but on graphs with the size and structure of the trained alphabet (3 to 15 anchor nodes), it stabilizes near the optimum within the time budget for over 99% of comparisons. Edge processing is deliberately simplified: any pair of edges is considered compatible, edge removal costs MINOR, edge insertion is free, and explicit attribute substitution at the edge level is not computed. Therefore, the entire structural penalty is concentrated at the vertex level via the cost scale below and the diagnostic weighting below.

Editing operations are evaluated on a six-level monotonically increasing scale that assigns a single cost class to each quality category of correspondence: NO_COST = 0.00 – exact attribute match, MINOR = 0.65 – missing vertex or edge (deletion/insertion on one side), GENERAL = 0.75 – moderate attribute deviation, SEVERE = 1.00 – severe attribute deviation, NO_MATCH = 1.50 – when the attribute is present on only one node, IMPOSSIBLE = 10.00 – label mismatch or prohibited insertion. The scale values are round numbers selected by the engineer within the range

MINOR < GENERAL < SEVERE < NO_MATCH < IMPOSSIBLE; only MINOR = 0.65 is empirically established as a key finding of the preliminary tuning study. The ratio NO_MATCH: MINOR $\approx$ 2.3 makes the presence of a feature on one side and its absence on the other cost more than two structural removals, motivating the comparator to prefer like-to-like matches over forced cross-type substitutions. The IMPOSSIBLE guard condition acts as a strict barrier, completely eliminating Point $\leftrightarrow$ Vector substitutions from the optimization.

Within each candidate vertex substitution, the comparator must aggregate the feature-wise contributions into a single substitution score. Uniform aggregation ($1/N$ s per feature), as previous ablations have shown, dilutes the contribution of highly diagnostic features as soon as the list of features exceeds five elements. A mitigation is diagnostic feature weighting based on interval width – calculated concept-by-feature from the post-reduction interval width (defined below). Here, " $SCALE_STRENGTH(f)$ " is a fixed per-feature multiplier (default 1.0; reduced to 0.3 for redundant features such as "cycle_count"), " $[range_width\left(\{f\}_{c} \right)]$ " is the post-reduction interval width of the feature " f " on the concept " c ", and " ε " is a stabilization parameter ("1.0" in the current series) that prevents division by zero on features with degenerate widths. Categorical features are assigned a constant weight, independent of width. The final value of a node's substitution is the weighted sum of the feature-specific attribute values, as defined by the value scale. A narrow interval – that is, a feature whose value is stable across all absorbed originals and their augmented variants – yields a high weight, causing the comparator to focus on those features that the reduction operator has already identified as diagnostic.

$$w(f, c) = \frac{SCALE_STRENGTH(f)}{range_width(f, c) + \varepsilon}, \sum_f w(f, c) = 1 \quad (13)$$

The final similarity metric is constructed by bounded normalization of the upper bound of the edit distance relative to the size of the larger of the two compared graphs (the formula is given below), where n is equal to the sum of the number of vertices and edges of the corresponding graph. The denominator $GED_{upper} + \max(n_{test}, n_C)$ ensures that the mapping is bounded in "[0,1]" for any non-negative cost, with $\mathrm{sim} = 1$ for exact matches, and $\mathrm{sim} \rightarrow 0$ for graphs with incompatible sizes or label structures.

This mapping partially compensates for the asymmetry in graph size, but does not eliminate the bias in favor of structurally minimal concepts – a phenomenon that was empirically observed in the excessive activation of the compact attractor of the alphabet (Section 5). The predicted class is the concept with the maximum of the scoring function (formula at the end), where " $\Phi(C)$ " is the structural complexity of the concept (the sum of the number of vertices and edges), and " λ " is a small positive constant. This is a Bayesian prior in the spirit of the principle of minimal description length: a more complex concept is treated as a priori more specific, partially compensating for the bias of the raw similarity metric in favor of small attractors. Along with the predicted class, the classifier returns an explanatory artifact with four fields: the winning concept's identifier, the similarity score, and two complexities: $\Phi(C)$ and $\Phi(H_{test})$. These four fields constitute a complete local explanation of the decision: any subsequent audit of the prediction boils down to determining why the winning concept outperformed its closest competitor, without the need for a post-hoc surrogate model.

$$\mathrm{sim}(H_{test}, C) = 1 - \frac{GED_{upper}(H_{test}, C)}{GED_{upper}(H_{test}, C) + \max(n_{test}, n_C)} \quad (14)$$

$$\mathrm{score}(C) = \mathrm{sim}(H_{test}, C) + \lambda \cdot \log \Phi(C) \quad (15)$$

To illustrate the operation of the reduction operator, the formation of attractor 7_1 (the only class-7 attractor in the alphabet) is traced step by step. The initial state consists of nine selected originals of the digit 7 from the MNIST training set following Growing Neural Gas skeletonization and Ramer–Douglas–Peucker simplification. Due to handwriting variability, the graph examples differ in the presence of intermediate corner points along the horizontal segment, slight bends in the diagonal segment, and variations in the upper corner angle. The same graph editing distance kernel used during classification identifies an invariant kernel of three critical points in each example: the starting endpoint of the horizontal segment in the upper-left part of the field, the upper corner point in the upper-right part where the horizontal segment transitions into the diagonal segment, and the endpoint of the diagonal segment in the lower part of the field. These three Point vertices are connected by two Vector vertices: one for the horizontal segment between the starting point and the upper corner point, and one for the diagonal segment between the upper corner point and the end point. All other vertices found only in

the subset of originals – additional intermediate corners on the horizontal segment, micro-bends along the diagonal – are removed in the next step as elements with frequency " $f_x < 1$ ", in accordance with Definition 1.

For each remaining vertex of the invariant core, the reduction operator determines the range of valid values for each coordinate feature among the absorbed originals. The intervals reflect the expected topology of the digit: the initial point is located in the upper-left part, the upper corner point in the upper-right, and the final point in the lower half with a relatively wide horizontal spread. Vector nodes inherit similar intervals: the horizontal segment is consistently located in the upper band with a small positional spread, while the diagonal segment crosses the field from the upper-right section to the lower band. The width of the resulting interval directly determines the diagnostic weight: coordinates whose values consistently repeat between originals carry greater weight during classification. After several iterations, the structure stabilizes into a linear bidirectional chain alternating between Point and Vector vertices: initial Point – horizontal Vector – upper corner Point – diagonal Vector – final Point. Subsequent iterations expand the intervals but do not introduce new vertices – this is a sign of reaching an attractor. The resulting compact attractor 7_1 has 5 Point + Vector vertices and a complexity of $\Phi(C) = 5 + 4 = 9$, making it one of the smallest non-trivial attractors in the alphabet. The compact structure encodes the topological invariant of the digit 7, yet that very compactness is the source of the overactivation pattern described in Chapter 5: a small, highly segmented attractor wins many ambiguous matches against larger concepts.

5. Research results

The experiment was conducted on a subset of MNIST images with intact contours. The original images, which were 28×28 pixels in size, were scaled to 100×100 using the Lanczos method so that the skeleton representation would have sufficient resolution to detect endpoints and convergence points. From the initial test set of 10,000 images, the manifest contour integrity filter (structure = complete) retained 8,707 images; the class-wise counts of valid images are shown in Table 3 (the "Volume" column). Images with discontinuous or fragmented contours were deliberately excluded from the scope of this study, as their recognition requires

an additional associative completion mechanism of a different mathematical nature, as already noted in Section 3 (Objective). The training set was formed from 76 selected originals, distributed across 13 concepts; each original was augmented using parametric augmentation (random rotation within $[-10^\circ, +10^\circ]$ and translation by $[-10\%, +10\%]$ of the canvas dimensions) with a per-concept multiplier ranging from five to ten (the full per-concept count is given in Table 2), resulting in 729 augmented instances and 805 training instances in total. No test images were used during concept formation.

The classification parameters are specified in the experiment configuration: threshold for skeletonization 110, simplification $\varepsilon = 4.55$; a 14-dimensional feature vector, as shown in Table 1; a graph edit distance comparator with iterative refinement of the upper bound, a time limit of 15 seconds per concept comparison; value scale NO_COST / MINOR / GENERAL / SEVERE / NO_MATCH / IMPOSSIBLE = 0/0.65/0.75/1.0/1.5/10; diagnostic weighting with stabilization parameter $\varepsilon = 1.0$; small positive logarithmic a priori coefficient λ on the logarithm of complexity. The full transfer was run end-to-end in batch mode on a multi-instance Nuclio deployment (8 classification consumers, 2 skeletonization consumers, 2 contour analysis consumers) with shared Kafka producers and concept-based graph caching (concept graphs are loaded once during function initialization and reused across messages). The campaign yielded 8,685 successfully classified images and 22 entries in the queue of unresolved messages ($\approx 0.25\%$) from images where the skeletonization stage was unable to construct a connected graph at the selected threshold. Failed images were excluded from metric calculations.

The trained alphabet contains 13 attractor concepts covering ten classes of digits; classes 1, 2, and 4 each have two attractors that capture different stylistic variants, while the remaining seven classes have one each. The sizes of the concepts (the number of preserved structural primitives) range from three nodes for the simplest variant of the digit 1 to fifteen nodes for the digit 8: the latter is the only digit in the alphabet with a closed-loop topology and a double connection, which is reflected in the largest number of preserved primitives. The compact attractor 7_1 (5 nodes, complexity $\Phi(C) = 9$), one of the smallest non-trivial attractors, is structurally a linear chain of two segments meeting at a single corner point. The largest – 8_1 (15 nodes,

complexity ≈ 29) – encodes the topology of a double loop of the number 8. The discrepancy in sizes is a property of the number shapes themselves; it is not a tuning artifact, and it forms the structural background of the overactivation pattern described below. The total number of originals, augmented variants, and training sets is shown in Table 2.

Of the 8,707 images fed into the transfer, 8,685 formed a valid skeleton graph and were included in the metric calculations; the remaining 22 images ($\approx 0.25\%$) did not form a connected graph at the binarization stage and were excluded from the per-class average metrics without a significant impact on the overall accuracy estimate. Accuracy on the valid set is 85.80%, weighted precision is 89.31%, recall is 85.80%, and F1 score is 86.66%. The 3.51 percentage point gap between weighted precision and recall is informative: it indicates uneven coverage of individual classes – primarily class 2 with a coverage of 60.0% and class 8 with a coverage of 74.7%; both classes are characterized by high precision and low coverage, which is a typical marker of an under-covering attractor alphabet. Class-specific precision, completeness, F1-score, and volume are presented in Table 3.

The highest F1 scores belong to classes 0, 9, and 1 – the digits with the most stable skeletal topology in handwritten form. Moderate F1 scores are observed for classes 3, 4, 5, 6, and 8. Two classes clearly deviate from the trend and account for the residual error budget of the alphabet: class 2 with an F1 of 71.1% (due to a completeness of 60.0% with high precision) and class 7 with an F1 of 78.8% (due to a precision of 66.2% with very high completeness). The excessive activation in class 7 as a complementary counterpart to the coverage deficit in class 2 is not a coincidence – it has a single common structural cause, which is explained below.

Residual errors are dominated by two mismatch patterns, which together account for the majority of the 14.20% class error budget. The complete mismatch matrix is shown in Figure 4; Table 4 lists the ten largest mismatch pairs by absolute count. The dominant pattern – mismatch $2 \rightarrow 7$ (252 instances, 66.0% of class 2 errors) – and the secondary patterns $3 \rightarrow 7$, $4 \rightarrow 7$, and $1 \rightarrow 7$ share a single structural cause. The compact attractor 7_1 with five Point + Vector vertices and a complexity of 9 encodes a linear topology consisting of two segments and a single angular bend. This topology is locally consistent with the stylistic variant of the digit 2 with an open tail (a corner in the upper-left part, followed by

an endpoint with an open loop and a diagonal segment to the lower endpoint), with elongated forms of the digit 3 without a closing arc, with segments of the digit 4 converging at a single branching point, and with thin variants of the digit 1 with a slanted upper stroke. None of these classes has a competing attractor in the alphabet that would favor an angle in the upper-left zone with a narrower coordinate interval, so the comparator selects 7_1 as the highest-similarity match for these images. The advantage of a small attractor in mapping is a direct empirical consequence of the "structure–parameter" decomposition from Section 4.1.8: when only the structural skeleton matches and the parametric mapping is broad, a small and widely segmented attractor wins many ambiguous matches against larger concepts.

The second pattern – Class 2 under-coverage – is the dual aspect of the same problem from the perspective of Class 2. The two attractors, 2_1 and 2_2, cover the closed-loop variant and the variant with a wide base of the digit 2, respectively; however, the stylistic variant with an open tail – common among left-handed people and in fast writing – falls outside the convex envelope of the intervals of both attractors but remains within the interval envelope of 7_1. This can be corrected by extending the alphabet with a third stylistic attractor, 2_3 (Section 6, correction (i); Section 7, prospective point (i)). The proportion of unclassified images in class 8 (127 out of 580 class 8 images, $\approx 21.9\%$ of class 8 and 86.4% of class 8 errors according to Table 4) – a separate phenomenon: 8_1 has the largest concept size in the alphabet and requires significant structural matching before evaluation becomes possible. When the binarization threshold breaks the closed loops of the number 8, the resulting graph either fails the connectivity check during the skeletonization stage (and is placed in the queue of unsolved messages, calculated as 22 entries for the entire corpus), or it passes this stage as a fragmented graph that no residual concept can align above the comparator threshold – in this case, the image is rejected before a class label is assigned. This is a pre-classification rejection mode, not a comparator rejection, and it can be mitigated either by adjusting the threshold for class 8 candidates (correction at the extraction stage) or by introducing a fragment completion stage at the graph level (a promising direction beyond the scope of the current study).

Direct numerical comparisons between architectures should be made with caution, as the leading candidate baseline models were evaluated under conditions

different from the current ones. Convolutional neural networks from the LeNet-5 family [4] and modern residual variants achieve over 99% accuracy on the standard MNIST test split, but they require thousands of labeled examples per class, a full training cycle with backpropagation, and do not offer built-in explanations – the network’s decision is based on linear combinations of feature maps, which can only be interpreted by humans through post-hoc visualization surrogates, the fragility of which is documented by Slack et al. [14]. Support vector machines with Gaussian kernels and invariance-aware preprocessing achieve approximately 98.6% accuracy on the standard MNIST dataset; on raw pixels without invariance augmentation, the result is more modest. Multilayer perceptrons with skewing and affine augmentation achieve 97–98% [4]. The prototype networks by Snell et al. [5], following the canonical 5-way 5-shot protocol on Omniglot or miniImageNet, achieve approximately 95–97% within the protocol benchmark; this figure is not directly comparable to the current 85.80% under the open MNIST test protocol. A previous study by the graph attractors team on six classes [16] on the restricted MNIST alphabet (digits 1, 2, 3, 6, 7, 9; 5,467 test images, 8 attractors, 60-second comparison budget, exact graph editing distance) achieved 85.42% accuracy. The current study extends the methodology to the full ten-class alphabet using contour integrity filtering, expands the alphabet to 13 attractors, introduces diagnostic feature weighting based on post-reduction interval width, and transitions to an iterative upper bound

of the graph editing distance with progressive refinement within a significantly smaller time budget – a combination of which yields the 85.80% accuracy reported here.

Therefore, the proposed architecture does not compete with deep convolutional networks in terms of absolute accuracy on the standard MNIST benchmark. Its niche lies in the simultaneous fulfillment of three properties, which in Chapter 2 are described as having been solved separately but not together: training without backpropagation, operation with a small number of examples per class, and local interpretability of each decision at the model level. The result of 85.80%, achieved by thirteen attractor concepts trained on 76 selected originals (from three to nine per concept), demonstrates that this trio is achievable in practice on a non-trivial recognition task – this is empirical confirmation of Tasks (4) and (5) of Section 3 and closes the gap stated in Section 2.

Figures 1–4 provide visual support. Figure 1 shows the complete transfer for a single image: the raw raster, the binary mask, the skeleton after algorithmic approximation of the medial line, the simplified graph, and the selected concept attractor. Figure 2 shows a graph representation of a typical handwritten digit 7 with annotations of critical points. Figure 3 shows the resulting attractor 7_1 after reduction with superimposed valid coordinate intervals. Figure 4 shows the mismatch matrix; the visual distinctiveness of column 7 quantitatively corresponds to the excessive activation of attractor 7_1, documented in Table 4.

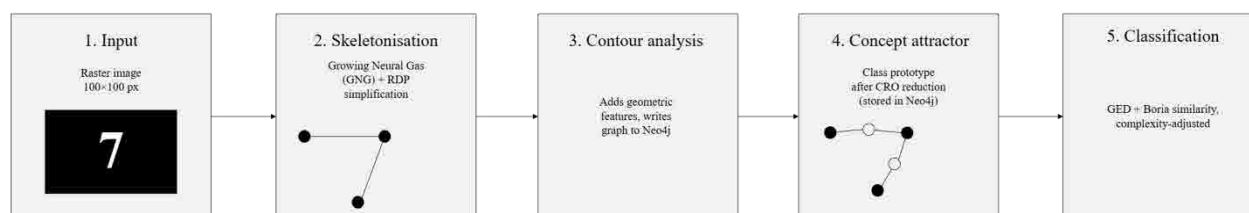


Fig 1. End-to-end transfer from a raster image to classification by a concept-attractor based on graph editing distance

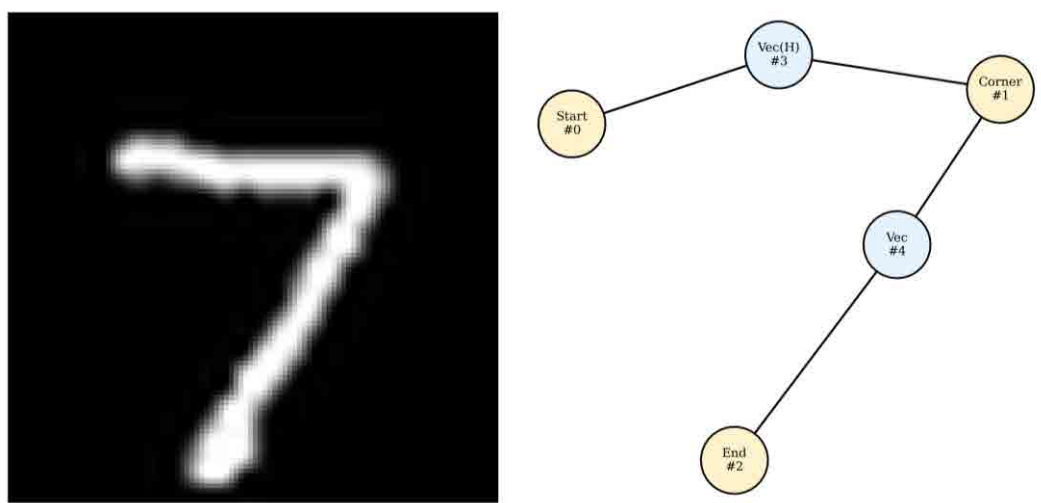


Fig 2. Graph representation of the digit 7: on the left – the original MNIST raster; on the right – the corresponding skeleton graph with Point and Vector nodes

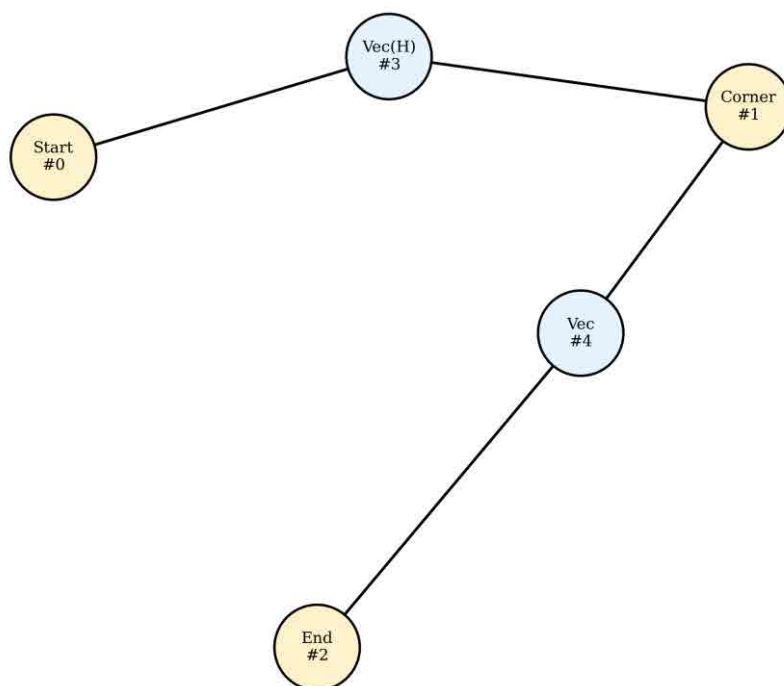


Fig 3. Concept attractor 7_1 after structural reduction: a linear chain consisting of three Point nodes and two Vector nodes – one of the smallest non-trivial attractors in the alphabet

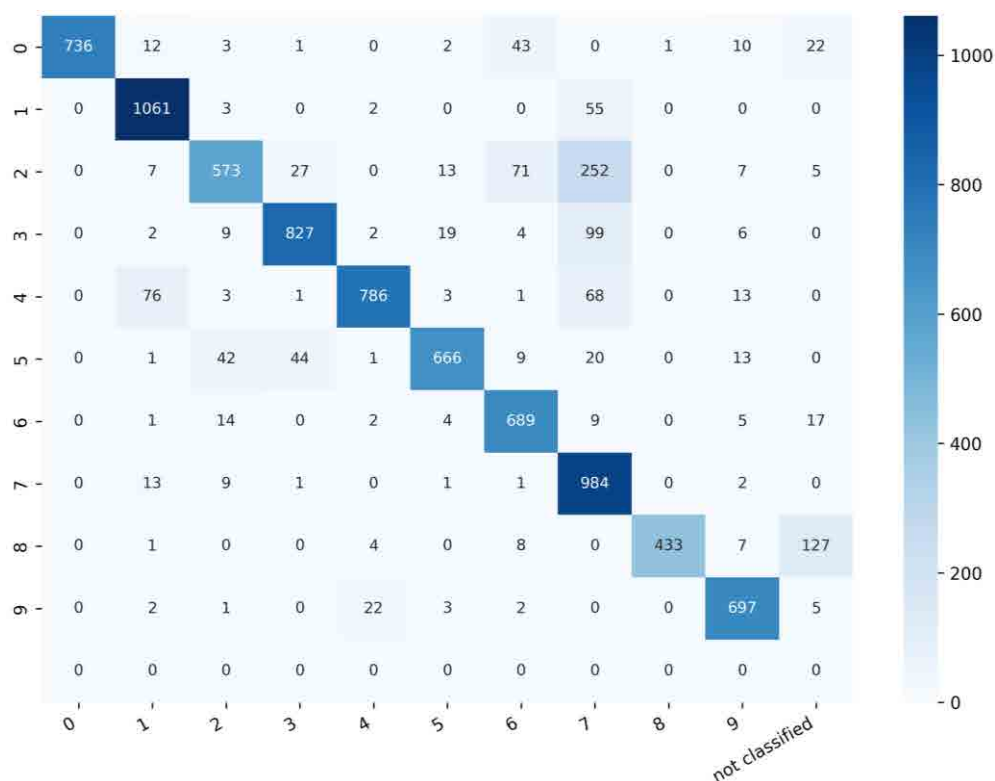


Fig 4. Mismatch matrix on a subset of MNIST with closed contours (8,685 successfully classified images across 10 classes). The "not classified" column contains images that did not match any concept

Table 2. Concept alphabet: number of originals per concept, number of augmented variants, training set size after augmentation, number of concept nodes, and class membership

Concept	Class	Originals	Augmented	Training instances	Number of concept nodes
0 1	0	3	30	33	12
1 1	1	3	30	33	7
1 2	1	4	40	44	3
2 1	2	6	60	66	7
2 2	2	3	30	33	10
3 1	3	6	30	36	11
4 1	4	7	70	77	8
4 2	4	6	59	65	9
5 1	5	6	60	66	9
6 1	6	7	70	77	10
7 1	7	9	90	99	5
8 1	8	8	80	88	15
9 1	9	8	80	88	8
Total	–	76	729	805	–

Table 3. Class-wise classifier metrics on the MNIST subset with connected contours: precision, recall, F1, and volume

Class	Precision, %	Recall, %	F1, %	Volume
0	100.0	88.7	94.0	830
1	90.2	94.6	92.4	1121
2	87.2	60.0	71.1	955
3	91.8	85.4	88.5	968
4	96.0	82.6	88.8	951
5	93.7	83.7	88.4	796
6	83.2	93.0	87.8	741
7	66.2	97.3	78.8	1011
8	99.8	74.7	85.4	580
9	91.7	95.2	93.4	732

Table 4. Top 10 pairs of discrepancies (true → predicted): absolute count and proportion of class errors

True class	Predicted class	Count	% of class errors
2	7	252	66.0
8	unclassified	127	86.4
3	7	99	70.2
4	1	76	46.1
2	6	71	18.6
4	7	68	41.2
1	7	55	91.7
5	3	44	33.8
0	6	43	45.7
5	2	42	32.3

6. Discussion of results

This work is conceptual in nature and requires further research. The following discussion addresses six points that arose during the theoretical exposition, followed by the experimental section on the failure modes observed in Chapter 5 and the areas of application where the proposed approach offers a competitive advantage.

First, exogenous primitives detectors. Although the work focuses on self-organizing systems, primitives detectors are specified exogenously rather than learned. This is a model-specific limitation; in the brain, the formation of corresponding sensory areas is genetically programmed rather than learned. Future research should include a model of endogenous formation of primitives – for example, through hierarchical autoencoders or self-organizing maps.

Second, exogenous coordinates and scales. Coordinate systems and measurement scales are also specified exogenously. Endogenous scales and coordinates are likely formed in the brain based on both genetics and experience; a full exploration of this lies beyond the scope of the current study. In the experiment, the spatial coordinate system is fixed on the MNIST dataset, and the feature scales (e.g., normalized coordinates in $[-1, +1]$) are specified by deterministic normalizers.

Third, the selection of anchors. The central problem of the theory is the search for critical (anchor) structural points that constitute the attractor. The theory assumes that these points are determined by frequency with respect to a hypergraph representation; the efficient search among many candidates is a separate subfield. The current implementation uses a deterministic heuristic for identifying critical points (endpoints, corner points, merge points) as post-hoc preprocessing – this choice is convenient but does not model the endogenous nature of anchor selection.

Fourth, attention as a heuristic. The formal theory of the attention operator is critical, as it is the main mechanism for circumventing the NP-hardness of global hypergraph matching. This work uses only heuristic anchor attention.

Fifth, endogenous ontology and scaling. A key strength of this approach is the natural formation of an endogenous ontology – the system's internal model of the external world. This ontology is formed both by the maps of neurons of individual classes and by their hierarchy [41]. The scalability problem is solved automatically: detector neurons are trained independently, and their hierarchy generates a hierarchy of concepts that has neurobiological support [28].

Sixth, the dendritic computation hypothesis. The key idea is that a neuron's dendritic tree is the primary "computational" system that constructs a structural attractor; the neuron's response is its aggregated characteristic. This is consistent with neurobiological evidence that individual dendritic branches act as independent computational units and that a neuron's firing is determined by the activity of a limited subset of synapses [30, 31, 32].

From an experimental perspective, three failure modes warrant discussion. First, the completeness for class 2 is 60.0%, indicating insufficient coverage by concepts 2_1 and 2_2: the stylistic variability of handwritten 2 (forms with a closed loop versus forms with an open tail) is not captured by the two attractors and would require a third stylistic attractor. This is consistent with the statement in Section 4.1 that intraclass variation should be expressed by additional subclass attractors when the conditions of separability are met. Second, the accuracy on class 7 is 66.2%, reflecting the overfitting of the compact attractor 7_1 (5 nodes, complexity 9), one of the smallest non-trivial attractors in the trained alphabet. The advantage of small attractors in terms of dimension-dependent coverage is a direct

consequence of the "structure–parameter" decomposition: when the structural skeleton coincides and the parametric coverage is broad, a small, widely segmented attractor wins many ambiguous matches. Third, the proportion of unclassified items in class 8 is non-trivial, because attractor 8_1 has the highest complexity (15 nodes) in the alphabet, requiring significant structural matching before classification becomes possible. These three patterns confirm that the structure of classification errors based on attractors is interpreted conceptually and across a range of features, satisfying the requirement for explainability set forth in Section 1.

Areas of application. The framework's traceability – where each classification decision points to a named concept attractor and specific ranges of anchor features – makes it a candidate for high-confidence optical symbol recognition (medical forms, legal documents), for industrial defect detection in contour structures (welds, stampings, printed circuit boards), where the number of defect classes is small and labeled examples are scarce, and for educational analytics, where pedagogical explanations are required. Since training requires only a handful of positive examples per class, the approach is also relevant for rapid re-targeting in engineering applications – adapting an existing alphabet to a new font, a new sensor configuration, or a new manufacturing operation – without retraining a large optimization model. In regulated areas covered by the EU AI Act (EU AI Act) [1] and the General Data Protection Regulation (GDPR) [2] regarding the transparency of automated decisions, the ability to automatically generate an interpretable decision log for auditing is a decisive advantage – this aligns with the journal's priority regarding the industrial application of results.

7. Conclusions

This paper proposes a formal learning model based on the structural reduction of hypergraphs and the identification of invariants during the process of aligning a set of observations. Unlike traditional statistical approaches, in which learning is formulated as the minimization of an error functional, the proposed model views learning as convergence to a structural attractor.

The main result is the introduction and formalization of the concept of a structural attractor as a fixed point of a reduction operator acting on the space of coordinated hypergraphs.

It is shown that the attractor has the following properties:

- (i) it is the result of a finite reduction process;
- (ii) it is unique for a given class of objects;
- (iii) it admits an equivalent characterization as the intersection of structural invariants and as the set of single elements of frequency in the cumulative carrier;
- (iv) it is invariant with respect to the order of data presentation.

Furthermore, the attractor naturally decomposes into two levels: the structural level (topology of relations) and the parametric level (segments and metrics), which separates the tasks of identifying structural invariants and their parameterization. An alternative formulation based on energy minimization for related neural network models was developed in [46].

The fundamental difference from classical machine learning models lies in the contrast between operators: while statistical neural network models use the minimization operator of the error/divergence functional, the proposed approach employs the algebra of structural invariants – specifically, the fixed-point equation $R(C) = C$ on a partially ordered set of hypergraphs. Learning is interpreted not as a search for the minimum of a functional, but as the computation of the fixed point of the operator on the space of structures. This change in formalism has several consequences.

First, there is no longer a need to specify an error functional or an external optimality criterion.

Second, learning becomes endogenous: it is determined by the data structure.

Third, a natural mechanism for robustness to variations in input representations arises from the identification of invariants.

The model opens up prospects for further research: conditions for the existence and stability of attractors, the design of efficient hypergraph matching algorithms, and the systematic study of learning with a limited number of examples (few-shot and zero-shot modes) without statistical optimization procedures. In summary, this work lays the foundations for an alternative theoretical paradigm of learning, in which structural invariants, reduction, and the dynamics of attractors play a central role.

The empirically proposed transfer achieved 85.80% accuracy (89.31% weighted precision, 85.80% recall, 86.66% F1-score) on the MNIST subset with complete contours (8,707 valid images), having been trained on 76 selected originals – ranging from three to nine

per concept – expanded to 805 instances via using parametric rotation ($\pm 10^\circ$) and translation ($\pm 10\%$). The trained alphabet of 13 attractors had a conceptual number of nodes ranging from 3 to 15. These metrics confirm the "structure-parameter" decomposition hypothesis from Sections 4.1.7–4.1.8: a structural skeleton extracted from a small sample is sufficient to cover the class, whereas the metric is induced endogenously through augmentation-driven segmentation rather than learned through statistical regularization.

Three directions for further work emerge from the experimental analysis.

(i) Restore the incompleteness of class 2 by introducing a third stylistic attractor, 2_3 – so that open-tail and closed-loop forms become two distinct subclass detectors rather than sharing a single common attractor.

(ii) Replace heuristic attention based on anchors with an endogenous attention operator that searches for anchors dynamically rather than computing them via deterministic post-hoc preprocessing; this directly addresses the fourth point of Section 6 and operationalizes Definition 6 / Axiom 4 in the implementation.

(iii) Extend the framework to EMNIST and Omniglot to verify whether trained concept-attractors transfer across alphabets – this verifies the claim made in Section 6 that the framework forms a hierarchy of endogenous concepts.

Conflict of interest

The authors declare that they have no conflicts of interest, including financial, personal, authorial, or any other conflicts that could influence the research or the results published in this article.

Funding

The study was conducted without financial support.

Data availability

Data will be provided upon reasonable request. The model's source code, experiment configuration, manifest of valid images, per-concept metrics, error log, and misclassification matrix will be provided upon reasonable request to the corresponding author

(Y. Parzhin, e-mail: yparzhyn@augusta.edu); the request must include a justification of the scientific purpose.

Use of artificial intelligence

The authors used AI tools (Anthropic Claude Opus 4.7, model ID "claude-opus-4-7", accessed via the Anthropic API) during four stages of this manuscript's preparation.

1. Preliminary domain analysis – queries regarding related fields (graph editing distance, structural image recognition, few-shot learning, post-hoc explainability) to delineate the scope of related work. The authors independently verified each AI-generated statement against primary sources before inclusion.

2. Literature search – generation of candidate reference lists, which were then verified by the authors via Crossref / OpenAlex / Scopus for DOI, author identity, and publication metadata before inclusion in the citations.

3. Cross-validation of results – systematic comparison of experimental findings (85.80% accuracy on the MNIST subset with connected contours, per-concept F1 scores, top-10 discrepancies) with published results for MNIST (CNN ~99%, SVM ~98%, MLP ~97–98%) and few-shot learning (Prototypical Networks ~95–97% on Omniglot/miniImageNet) to ensure the correct contextualization of statements.

4. Translation editing – correction of errors in the translation of draft materials from Ukrainian and Russian into English; the author reviewed each phrase suggested by the AI before acceptance. All conceptual content (problem formulation, theoretical framework, experimental design, analysis, and conclusions) was written by human authors. AI tools did not participate in the formulation of theorems, proofs, or experimental hypotheses and did not influence the scientific conclusions of the work.

Acknowledgments

The authors express their sincere gratitude to Alexander Schwartzman, Dean of the School of Computer and Cyber Sciences at Augusta University (Georgia, USA), for his invaluable assistance and support in conducting this research.

References

1. European Parliament, Council of the European Union (2024), "Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)", Official Journal of the European Union, L series, 12 July 2024.
2. European Parliament, Council of the European Union (2016), "Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation)", Official Journal of the European Union, L 119, 4 May 2016, pp. 1–88
3. Parzhyn, Y. (2025), "Architecture of information", *arXiv preprint arXiv:2503.21794*, 81 p. DOI: <https://doi.org/10.48550/arXiv.2503.21794>
4. LeCun, Y., Bottou, L., Bengio, Y., and Haffner, P. (1998), "Gradient-based learning applied to document recognition", *Proceedings of the IEEE*, Vol. 86, No. 11, pp. 2278–2324. DOI: <https://doi.org/10.1109/5.726791>
5. Snell, J., Swersky, K., and Zemel, R. (2017), "Prototypical networks for few-shot learning", *Advances in Neural Information Processing Systems 30 (NeurIPS 2017)*, pp. 4077–4087. DOI: <https://doi.org/10.48550/arXiv.1703.05175>
6. Vinyals, O., Blundell, C., Lillicrap, T., Kavukcuoglu, K., and Wierstra, D. (2016), "Matching networks for one-shot learning", *Advances in Neural Information Processing Systems 29 (NeurIPS 2016)*, pp. 3630–3638. DOI: <https://doi.org/10.48550/arXiv.1606.04080>
7. Finn, C., Abbeel, P., and Levine, S. (2017), "Model-agnostic meta-learning for fast adaptation of deep networks", *Proceedings of the 34th International Conference on Machine Learning (ICML 2017)*, pp. 1126–1135. DOI: <https://doi.org/10.48550/arXiv.1703.03400>
8. Minh, D., Wang, H. X., Li, Y. F., and Nguyen, T. N. (2022), "Explainable artificial intelligence: a comprehensive review", *Artificial Intelligence Review*, Vol. 55, No. 5, pp. 3503–3568. DOI: <https://doi.org/10.1007/s10462-021-10088-y>
9. Hooshyar, D., and Yang, Y. (2024), "Problems with SHAP and LIME in interpretable AI for education: a comparative study of post-hoc explanations and neural-symbolic rule extraction", *IEEE Access*, Vol. 12, pp. 137472–137490. DOI: <https://doi.org/10.1109/ACCESS.2024.3463948>
10. Rajabi, E., and Etmiani, K. (2024), "Knowledge-graph-based explainable AI: a systematic review", *Journal of Information Science*, Vol. 50, No. 4, pp. 1019–1029. DOI: <https://doi.org/10.1177/01655515221112844>
11. Nawaz, U., Anees-ur-Rahaman, M., and Saeed, Z. (2025), "A review of neuro-symbolic AI integrating reasoning and learning for advanced cognitive systems", *Intelligent Systems with Applications*, Vol. 26, Article 200541. DOI: <https://doi.org/10.1016/j.iswa.2025.200541>
12. Ribeiro, M. T., Singh, S., and Guestrin, C. (2016), "Why should I trust you?: Explaining the predictions of any classifier", *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD 2016)*, pp. 1135–1144. DOI: <https://doi.org/10.1145/2939672.2939778>
13. Lundberg, S. M., and Lee, S.-I. (2017), "A unified approach to interpreting model predictions", *Advances in Neural Information Processing Systems 30 (NeurIPS 2017)*, pp. 4765–4774. DOI: <https://doi.org/10.48550/arXiv.1705.07874>
14. Slack, D., Hilgard, S., Jia, E., Singh, S., and Lakkaraju, H. (2020), "Fooling LIME and SHAP: adversarial attacks on post hoc explanation methods", *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society (AIES '20)*, pp. 180–186. DOI: <https://doi.org/10.1145/3375627.3375830>
15. Ding, K., Wang, J., Li, J., Shu, K., Liu, C., and Liu, H. (2020), "Graph prototypical networks for few-shot learning on attributed networks", *Proceedings of the 29th ACM International Conference on Information and Knowledge Management*, pp. 295–304. DOI: <https://doi.org/10.1145/3340531.3411922>
16. Lapin, M., and Bokhan, K. (2025), "Few-shot learning of a graph-based neural network model without backpropagation", *Management Information System and Devices*, No. 187, pp. 103–122. DOI: <https://doi.org/10.30837/0135-1710.2025.187.103>
17. Sanfeliu, A., and Fu, K.-S. (1983), "A distance measure between attributed relational graphs for pattern recognition", *IEEE Transactions on Systems, Man, and Cybernetics*, Vol. SMC-13, No. 3, pp. 353–362. DOI: <https://doi.org/10.1109/TSMC.1983.6313167>
18. Conte, D., Foggia, P., Sansone, C., and Vento, M. (2004), "Thirty years of graph matching in pattern recognition", *International Journal of Pattern Recognition and Artificial Intelligence*, Vol. 18, No. 3, pp. 265–298. DOI: <https://doi.org/10.1142/S0218001404003228>
19. Riesen, K., and Bunke, H. (2009), "Approximate graph edit distance computation by means of bipartite graph matching", *Image and Vision Computing*, Vol. 27, No. 7, pp. 950–959. DOI: <https://doi.org/10.1016/j.imavis.2008.04.004>
20. Piao, C., Xu, T., Sun, X., Rong, Y., Zhao, K., and Cheng, H. (2023), "Computing graph edit distance via neural graph matching", *Proceedings of the VLDB Endowment*, Vol. 16, No. 8, pp. 1817–1829. DOI: <https://doi.org/10.14778/3594512.3594514>
21. Moscatelli, A., Piquenot, J., Berar, M., Heroux, P., and Adam, S. (2024), "Graph node matching for edit distance", *Pattern Recognition Letters*, Vol. 184, pp. 14–20. DOI: <https://doi.org/10.1016/j.patrec.2024.05.020>

22. Xie, Y., Liang, Y., Wen, C., Qin, A. K., and Gong, M. (2024), "Federated collaborative graph neural networks for few-shot graph classification", *Machine Intelligence Research*, Vol. 21, No. 6, pp. 1077–1091. DOI: <https://doi.org/10.1007/s11633-023-1463-3>
23. Chukhran, I., Udovenko, S., Shergin, V., and Chala, L. (2025), "Method for augmenting 3D point cloud models using graph neural networks", *Innovative Technologies and Scientific Solutions for Industries*, No. 2(32), pp. 129–150. DOI: <https://doi.org/10.30837/2522-9818.2025.2.129>
24. von der Malsburg, C. (1999), "The what and why of binding: the modeler's perspective", *Neuron*, Vol. 24, No. 1, pp. 95–104. DOI: [https://doi.org/10.1016/S0896-6273\(00\)80825-9](https://doi.org/10.1016/S0896-6273(00)80825-9)
25. Buzsáki, G. (2010), "Neural syntax: cell assemblies, synapse ensembles, and readers", *Neuron*, Vol. 68, No. 3, pp. 362–385. DOI: <https://doi.org/10.1016/j.neuron.2010.09.023>
26. Yarbus, A. L. (1967), *Eye Movements and Vision*, Plenum Press, New York, 222 p. DOI: <https://doi.org/10.1007/978-1-4899-5379-7>
27. Bajcsy, R., Aloimonos, Y., and Tsotsos, J. K. (2018), "Revisiting active perception", *Autonomous Robots*, Vol. 42, No. 2, pp. 177–196. DOI: <https://doi.org/10.1007/s10514-017-9615-3>
28. Rucci, M., and Victor, J. D. (2015), "The unsteady eye: an information-processing stage, not a bug", *Trends in Neurosciences*, Vol. 38, No. 4, pp. 195–206. DOI: <https://doi.org/10.1016/j.tins.2015.01.005>
29. Quiroga, R. Q., Reddy, L., Kreiman, G., Koch, C., and Fried, I. (2005), "Invariant visual representation by single neurons in the human brain", *Nature*, Vol. 435, pp. 1102–1107. DOI: <https://doi.org/10.1038/nature03687>
30. Quiroga, R. Q. (2012), "Concept cells: the building blocks of declarative memory functions", *Nature Reviews Neuroscience*, Vol. 13, No. 8, pp. 587–597. DOI: <https://doi.org/10.1038/nrn3251>
31. Magee, J. C. (2000), "Dendritic integration of excitatory synaptic input", *Nature Reviews Neuroscience*, Vol. 1, No. 3, pp. 181–190. DOI: <https://doi.org/10.1038/35044552>
32. Häusser, M. (2001), "Synaptic function: dendritic democracy", *Current Biology*, Vol. 11, No. 1, pp. R10–R12. DOI: [https://doi.org/10.1016/S0960-9822\(00\)00034-8](https://doi.org/10.1016/S0960-9822(00)00034-8)
33. Branco, T., and Häusser, M. (2010), "The single dendritic branch as a fundamental functional unit in the nervous system", *Current Opinion in Neurobiology*, Vol. 20, No. 4, pp. 494–502. DOI: <https://doi.org/10.1016/j.conb.2010.07.009>
34. Larkum, M. E. (2022), "Are dendrites conceptually useful?", *Neuroscience*, Vol. 489, pp. 4–14. DOI: <https://doi.org/10.1016/j.neuroscience.2022.03.008>
35. Parzhin, Y., Galkyn, S., and Sobol, M. (2022), "Method for binary contour images vectorization of handwritten characters for recognition by detector neural networks", *2022 IEEE 3rd KhPI Week on Advanced Technology (KhPIWeek)*, pp. 1–6. DOI: <https://doi.org/10.1109/KhPIWeek57572.2022.9916331>
36. Caspard, N., Leclerc, B., and Monjardet, B. (2012), *Finite Ordered Sets: Concepts, Results and Uses*, Cambridge University Press, Cambridge, 337 p. DOI: <https://doi.org/10.1017/CBO9781139005135>
37. Hanika, T., and Hirth, J. (2022), "Knowledge cores in large formal contexts", *Annals of Mathematics and Artificial Intelligence*, Vol. 90, No. 6, pp. 537–567. DOI: <https://doi.org/10.1007/s10472-022-09790-6>
38. Krotov, D. (2023), "A new frontier for Hopfield networks", *Nature Reviews Physics*, Vol. 5, No. 7, pp. 366–367. DOI: <https://doi.org/10.1038/s42254-023-00595-y>
39. Ramsauer, H., Schäfl, B., Lehner, J., Seidl, P., Widrich, M., Adler, T., Gruber, L., Holzleitner, M., Pavlović, M., Sandve, G. K., Greiff, V., Kreil, D., Kopp, M., Klambauer, G., Brandstetter, J., and Hochreiter, S. (2021), "Hopfield Networks is All You Need", *International Conference on Learning Representations (ICLR 2021)*. DOI: <https://doi.org/10.48550/arXiv.2008.02217>
40. Bodnar, C., Di Giovanni, F., Chamberlain, B. P., Liò, P., and Bronstein, M. M. (2022), "Neural Sheaf Diffusion: A Topological Perspective on Heterophily and Oversmoothing in GNNs", *Advances in Neural Information Processing Systems 35 (NeurIPS 2022)*. DOI: <https://doi.org/10.48550/arXiv.2202.04579>
41. Parzhin, Y., Kosenko, V., Podorozhniak, A., Malyeyeva, O., and Timofeyev, V. (2020), "Detector neural network vs connectionist artificial neural networks", *Neurocomputing*, Vol. 414, pp. 191–203. DOI: <https://doi.org/10.1016/j.neucom.2020.07.025>
42. Riesen, K. (2015), *Structural Pattern Recognition with Graph Edit Distance: Approximation Algorithms and Applications*, Springer International Publishing, Cham. DOI: <https://doi.org/10.1007/978-3-319-27252-8>
43. Fritzke, B. (1995), "A growing neural gas network learns topologies", *Advances in Neural Information Processing Systems 7 (NeurIPS 1994)*, pp. 625–632. DOI: <https://doi.org/10.5555/2998687.2998765>
44. Douglas, D. H., and Peucker, T. K. (1973), "Algorithms for the reduction of the number of points required to represent a digitized line or its caricature", *Cartographica: The International Journal for Geographic Information and Geovisualization*, Vol. 10, No. 2, pp. 112–122. DOI: <https://doi.org/10.3138/FM57-6770-U75U-7727>
45. Hagberg, A. A., Schult, D. A., and Swart, P. J. (2008), "Exploring network structure, dynamics, and function using NetworkX", *Proceedings of the 7th Python in Science Conference (SciPy 2008)*, pp. 11–15. DOI: <https://doi.org/10.25080/TCWV9851>

46. Parzhyn, Y., Lapin, M., and Bokhan, K. (2025), "A new approach to building energy models of neural networks", *Advanced Information Systems*, Vol. 9, No. 4, pp. 100–119. DOI: <https://doi.org/10.20998/2522-9052.2025.4.13>

Received (Надійшла) 08.08.2026

Accepted for publication (Прийнята до друку) 21.05.2026

Publication date (Дата публікації) 27.06.2026

Відомості про авторів / About the Authors

Лapin Микита Олексійович – Національний технічний університет "Харківський політехнічний інститут", здобувач вищої освіти ступеня доктора філософії кафедри системного аналізу та інформаційно-аналітичних технологій, Харків, Україна;
Mykyta Lapin – National Technical University "Kharkiv Polytechnic Institute", Postgraduate Student of the Systems Analysis and Information-Analytical Technologies Department, Kharkiv, Ukraine;
 e-mail: Mykyta.Lapin@cit.khpi.edu.ua
 ORCID ID: <https://orcid.org/0009-0003-6307-1172>
 Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=60171161300>

Паржин Юрій Володимирович – доктор технічних наук, професор, Школа Комп'ютерних та Кібер Наук Університету Огасти, постдокторант, Огаста, Джорджія, США.
Yurii Parzhyn – Doctor of Technical Sciences, Professor, School of Computer and Cyber Sciences Augusta University, Postdoctoral Fellow, Augusta, Georgia, USA;
 e-mail: yparzhyn@augusta.edu
 ORCID ID: <https://orcid.org/0000-0001-5727-1918>
 Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=57224412390>

Бохан Костянтин Олександрович – кандидат технічних наук, Національний технічний університет "Харківський політехнічний інститут", доцент кафедри системного аналізу та інформаційно-аналітичних технологій, Харків, Україна;
Kostiantyn Bokhan – Candidate of Technical Sciences, National Technical University "Kharkiv Polytechnic Institute", Associate Professor of the of Systems Analysis and Information-Analytical Technologies Department, Kharkiv, Ukraine;
 e-mail: kostiantyn.bokhan@khpi.edu.ua
 ORCID ID: <https://orcid.org/0000-0003-3375-2527>
 Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=57191592568>

Перевозник Кирило Максимович – Національний технічний університет "Харківський політехнічний інститут", здобувач вищої освіти ступеня доктора філософії кафедри системного аналізу та інформаційно-аналітичних технологій, Харків, Україна;
Kyrylo Perevoznyk – National Technical University "Kharkiv Polytechnic Institute", Postgraduate Student of the Systems Analysis and Information-Analytical Technologies Department, Kharkiv, Ukraine;
 e-mail: kyrylo.perevoznyk@cs.khpi.edu.ua
 ORCID ID: <https://orcid.org/0009-0009-2327-1501>
 Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=59564678000>

Александрова Тетяна Євгенівна – доктор технічних наук, професор, Національний технічний університет "Харківський політехнічний інститут", завідувачка кафедри системного аналізу та інформаційно-аналітичних технологій, Харків, Україна;
Tetiana Aleksandrova – Doctor of Technical Sciences, Professor, National Technical University "Kharkiv Polytechnic Institute", Head of the Systems Analysis and Information-Analytical Technologies Department, Kharkiv, Ukraine;
 e-mail: Tetiana.Aleksandrova@khpi.edu.ua
 ORCID ID: <https://orcid.org/0000-0001-9596-0669>
 Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=57189376480>

ІНВАРІАНТНО-СТРУКТУРНЕ НАВЧАННЯ: ФОРМУВАННЯ КОНЦЕПТІВ ЯК ДИНАМІКА ГІПЕРГРАФОВИХ АТРАКТОРІВ

Предметом дослідження є формування концептів класів об'єктів у системах машинного навчання як процесу структурної та параметричної редукції гіперграфового подання, а не як оптимізації функціонала втрат. **Мета роботи** – побудувати альтернативний підхід статистичного навчання штучних нейронних мереж без використання функції помилки (зворотне поширення), у якому реалізується інваріантне структурне навчання, де концепт класу визначається як структурний аттрактор (нерухома точка монотонного оператора редукції на частково впорядкованому просторі гіперграфів), і підтвердити цю теорію на прикладі розпізнавання рукописних цифр. **Завдання:** 1) формалізувати основу з аксіомами стабільності сегментації, строгої редуктивності, навчання за позитивними прикладами й локальності уваги; 2) довести збіжність і єдиність атрактора; 3) встановити інваріантність щодо порядку поглинання прикладів; 4) розкласти аттрактор на структурний і параметричний рівні; 5) оцінити конвеєр на підмножині MNIST із цілісними контурами. **Методи дослідження:** гіперграфи отримано за допомогою скелетизації бінарних контурів (*Growing Neural Gas* + Рамер – Дуглас – Пекер); концепти-атрактори сформовано редукцією за анкерами – критичними точками контуру. Класифікація використовує порівнювач за відстанню редагування графа з вартостями ознак і логарифмічним апіорі за складністю. Тренувальна множина – 76 оригіналів MNIST, доповнених аугментацією (ротація $\pm 10^\circ$, зсув $\pm 10\%$) до 805 примірників. **Результати.** Конвеєр навчає 13 концептів-атракторів із кількістю вузлів від 3 до 15. З 8 707 допустимих зображень MNIST з цілісними контурами 8 685 утворили валідні скелетні графи; на цій валідній підмножині досягнуто точність 85,80%, зважену влучність 89,31%, повноту 85,80% і F1-міру 86,66%. Структура помилок інтерпретована поконцептно: кожна помилка простежувана до конкретного атрактора й діапазону ознак. **Висновки.** Зафіксована продуктивність, досягнута без функціонала втрат і з трьох до дев'яти унікальних оригіналів на концепт-атрактор, підтверджує теоретичне твердження, що навчання є побудовою структурного атрактора, а не мінімізацією функціонала помилки. Основа дає принциповий шлях до пояснювальної маловибіркової класифікації та конструктивну альтернативу градієнтному навчанню.

Ключові слова: інваріантно-структурне навчання; структурний аттрактор; формування концепту; відстань редагування графа; маловибіркове розпізнавання; пояснювальний штучний інтелект; MNIST; редукція гіперграфа.

Бібліографічні описи / Bibliographic descriptions

Лапін М. О., Паржин Ю. В., Бохан К. О., Перевозник К. М., Александрова Т. Є. Інваріантно-структурне навчання: формування концептів як динаміка гіперграфових аттракторів. *Сучасний стан наукових досліджень та технологій в промисловості*. 2026. № 2 (36). С. 70–94. DOI: <https://doi.org/10.30837/2522-9818.2026.2.070>

Lapin, M., Parzhyn, Y., Bokhan, K., Perevoznyk, K., Aleksandrova, T. (2026), "Invariant structural learning: concept formation as hypergraph attractor dynamics", *Innovative Technologies and Scientific Solutions for Industries*, No. 2 (36), P. 70–94. DOI: <https://doi.org/10.30837/2522-9818.2026.2.070>

UDC 004.02, 004.67, 004.891.3, 616.24-002.5-02:316.342.6:316.62:314(477)

DOI: <https://doi.org/10.30837/2522-9818.2026.2.095>

Denys Nevinskyi, Yaroslav Vyklyuk

PROJECT-ORIENTED METHODOLOGY FOR PROACTIVE EPIDEMIC MANAGEMENT BASED ON ARTIFICIAL INTELLIGENCE AND SEIRS MODELS

The subject of the research presented in the article is the processes of proactive epidemic management as complex socio-medical systems under conditions of high uncertainty, limited resources, and dynamic external influences, including the use of mathematical epidemiological modeling and intelligent digital technologies. The purpose of the study is to develop a project-oriented methodology for epidemic management based on the integration of the classical SEIRS model with artificial intelligence, machine learning, and digital decision-support systems, enabling a transition from reactive response to proactive forecasting and prevention of infectious disease outbreaks. The following tasks are addressed in the article: the formation of a conceptual and process-based model of proactive epidemic management within a project-oriented approach; the development of an integrated architecture for the collection, analysis, and interpretation of epidemiological data; the adaptation of the SEIRS compartmental model to dynamic parameterization using machine learning methods; the justification of geospatial and multi-agent approaches for capturing heterogeneity in disease spread; and the identification of the role of intelligent optimization methods in managing interventions throughout the epidemic project lifecycle. The following methods are used: mathematical modeling of epidemiological processes based on systems of SEIRS differential equations; machine learning and ensemble forecasting methods for dynamic estimation of model parameters; geospatial analysis and multi-agent modeling to represent population mobility and spatial risk patterns; reinforcement learning methods for optimizing management decisions; and project-oriented management approaches for organizing and adapting epidemic response processes. The following results are obtained: principles of proactive epidemic management based on the integration of epidemiological modeling and artificial intelligence are formulated; a project-oriented methodology that aligns the project management cycle with the phases of the SEIRS model is proposed; a generalized structure of a hybrid intelligent model capable of adapting to different types of infectious diseases, including tuberculosis and COVID-19, is developed; and the potential for improving early risk detection and rational resource allocation under crisis conditions is substantiated. Conclusions: the application of the proposed project-oriented methodology and the integrated SEIRS model enhanced with artificial intelligence provides a foundation for the transition to proactive epidemic management, increases the adaptability of public health systems, and can be used as a universal methodological framework for the development of intelligent epidemic forecasting and prevention systems in complex and unstable environments.

Keywords: proactive epidemic management; project-oriented methodology; artificial intelligence; SEIRS model; machine learning; digital health; epidemiological forecasting; decision support systems.

1. Introduction

Infectious diseases remain one of the key threats to global public health security, despite significant advances in medical technology and epidemiological surveillance systems. The COVID-19 pandemic has demonstrated the limitations of traditional reactive approaches to epidemic management, which are largely focused on responding after outbreaks have been identified, rather than on preventing them. Additional challenges are posed by increasing population mobility, urbanization, socioeconomic inequality, and climate change, which complicate the prediction of the spread of infectious diseases.

For Ukraine, these problems are particularly acute amid prolonged military aggression, which has led to the disruption of the healthcare infrastructure, internal and

external population displacement, fragmentation of epidemiological data, and limited access to medical services in a number of regions. Under such conditions, maintaining the effectiveness of traditional surveillance systems and planning anti-epidemic measures is extremely difficult, necessitating the implementation of adaptive, intelligently managed, and proactive approaches to epidemic management.

Traditional epidemic management models are based on retrospective analysis of statistical data and are primarily used to describe existing disease trends. This approach is insufficiently effective in conditions of rapidly changing epidemiological situations, high uncertainty, and incomplete data. Reactive strategies lead to delayed implementation of management interventions, inefficient use of resources, and increased socioeconomic losses.

Proactive epidemic management involves shifting the focus from response to prediction, based on risk forecasting, early detection of potential outbreaks, and advance planning of containment measures. This approach requires the use of mathematical models capable of accounting for the complex dynamics of infectious processes, as well as tools for analyzing large volumes of heterogeneous data in real time.

Artificial intelligence and machine learning methods open up new possibilities for epidemiological forecasting and management. Unlike classical deterministic models, intelligent approaches allow for the adaptive assessment of key parameters of infection spread, taking into account socioeconomic, demographic, spatial, and behavioral factors. The integration of digital medical records, geospatial data, population mobility indicators, and clinical trial results forms the basis for building comprehensive management decision-support systems.

Of particular importance in this context are extended compartmental models, notably the SEIRS model, which accounts for latent stages of infection and the possibility of recovered individuals reverting to the susceptible group. Combining such models with artificial intelligence methods allows for a shift from static parameterization to a dynamic, context-dependent description of epidemiological processes.

Despite active research in the application of artificial intelligence in epidemiology, most existing studies focus on specific aspects of disease incidence forecasting or diagnostic automation. The integration of epidemiological modeling with project-oriented approaches to management – which allow for the systematic organization of decision-making, the implementation of interventions, and their adaptation in a dynamic environment – has been studied to a much lesser extent.

Thus, there is a research gap between mathematical models of infection spread and practices for managing epidemics as complex, multi-stage projects implemented under conditions of limited resources and high uncertainty.

The aim of this article is to develop and justify a project-oriented methodology for proactive epidemic management based on the integration of the SEIRS model with artificial intelligence methods and digital technologies.

To achieve this goal, the following tasks are addressed in this work:

- formalize the concept of proactive epidemic management within a project-based approach;

- develop an integrated methodological framework combining epidemiological modeling and intelligent data analysis methods;

- justify the use of dynamic parameterization of the SEIRS model using machine learning; demonstrate the adaptability of the proposed methodology to various types of infectious diseases and crisis conditions.

2 Analysis of recent studies and publications

2.1. Compartmental models and the development of SEIR/SEIRS approaches

Mathematical modeling of epidemics has its origins in the classical theory proposed by W. Kermack and A. McKendrick, which laid the foundations for the compartmentalization of populations and the formalization of infection and recovery processes [1]. Further development of these ideas led to the emergence of SEIR-type models, which account for the latent (incubation) period of the disease [2].

The SEIRS model, which additionally describes the loss of immunity and the return of recovered individuals to the susceptible group, is widely used for infectious diseases with a long-lasting or inconsistent immune response, particularly tuberculosis and COVID-19 [3]. Studies show that SEIRS models better reproduce the wave dynamics of epidemics and can describe endemic disease states [4].

Further modifications of SEIR/SEIRS models aim to account for time lags, spatial heterogeneity, and various intervention scenarios [5]. Such approaches improve forecasting accuracy but simultaneously complicate parameter calibration and require integration with additional data sources [6].

2.2. Artificial intelligence, geospatial, and multi-agent approaches in epidemiology

The growing availability of large datasets on epidemiological, socioeconomic, and mobility data has facilitated the active implementation of artificial intelligence and machine learning methods in infectious disease forecasting [7]. Review studies indicate that ML approaches are capable of identifying hidden nonlinear dependencies inaccessible to classical statistical methods and adapting to changing conditions of infection spread [8].

At the same time, numerous authors highlight the limitations of purely data-driven models, which, without mechanistic interpretation, may lose their generalization

ability and be sensitive to changes in context [9]. In this regard, hybrid approaches are being actively developed that combine compartmental models with neural networks, ensemble algorithms, and graph models [10].

Particular attention is paid to the use of graph neural networks and recurrent architectures to account for spatiotemporal dependencies and population mobility [11]. Such models demonstrate increased accuracy in short-term forecasting and can be used to assess the impact of policy interventions [12].

Geospatial models are a key tool for analyzing regional epidemic dynamics, as they allow for the consideration of population migration, transportation links, and local environmental characteristics [13]. Studies show that spatially explicit SEIR models provide a more realistic representation of infection spread compared to aggregated approaches [14].

Multi-agent systems (MAS) complement geospatial models by allowing the description of individual agents' behavior, their contacts, and individual characteristics [15]. Such approaches are particularly effective for modeling infections with long incubation periods and complex social determinants, notably tuberculosis [16].

In previous work, the authors developed geospatial multi-agent models to analyze the transmission of tuberculosis and COVID-19, allowing for the integration of demographic, medical, and spatial factors within a single modeling environment [17–20].

2.3. Optimization of management interventions and decision support systems

Proactive epidemic management involves not only forecasting but also selecting optimal intervention strategies under conditions of limited resources. In this context, reinforcement learning methods are being actively researched, enabling the formulation of management policies based on an assessment of the long-term consequences of decisions [21].

Studies demonstrate the effectiveness of RL approaches for optimizing non-pharmaceutical interventions, testing planning, and the allocation of healthcare resources [22]. Other works combine RL with compartmental models and real-world data for adaptive control of the intensity of restrictions at the regional level [23].

Decision support systems integrated with epidemiological models and analytical dashboards are viewed as a key tool for translating forecasts into

practical management actions [24]. In previous studies, the authors demonstrated the potential of using intelligent agents to support the management of global threats and complex crisis processes [25].

2.4. Summary of the literature review

The conducted literature review indicates that current research covers a wide range of approaches – from classical SEIRS models to hybrid systems based on artificial intelligence, geospatial analysis, and reinforcement learning [26–28]. At the same time, most studies focus on individual components of epidemiological analysis and do not form a comprehensive methodology for managing epidemics as complex projects. Thus, a research gap remains regarding the integration of SEIRS models, intelligent data analysis methods, multi-agent systems, and optimization approaches within a single project-oriented management cycle, which justifies the conduct of this study.

3. Purpose and objectives of the study

The aim of this work is to develop a project-oriented methodology for epidemic management based on the integration of the classical SEIRS model with artificial intelligence, machine learning, and digital decision support systems, thereby facilitating a shift from reactive response to proactive forecasting and prevention of infectious disease outbreaks.

The article addresses the following tasks: forming a conceptual and process model of proactive epidemic management within a project-based approach; developing an integrated architecture for the collection, analysis, and interpretation of epidemiological data; adapting the SEIRS compartmental model to the conditions of dynamic parameterization using machine learning methods; justifying the use of geospatial and multi-agent approaches to account for the heterogeneity of infection spread; determining the role of intelligent methods for optimizing management interventions in the project cycle of epidemic response.

4. Project-oriented methodology for proactive epidemic management

4.1. General architecture of the proposed methodology

The proposed methodology represents a systematic and cyclical approach to project management of processes for forecasting and minimizing epidemic

risks (Fig. 1). It integrates project management principles (particularly agile and hybrid methodologies) with artificial intelligence and machine learning methods to form an adaptive, data-driven decision support system.

The main goal of the methodology is to transform reactive responses to infectious disease outbreaks into proactive risk management based on predictive analytical models.

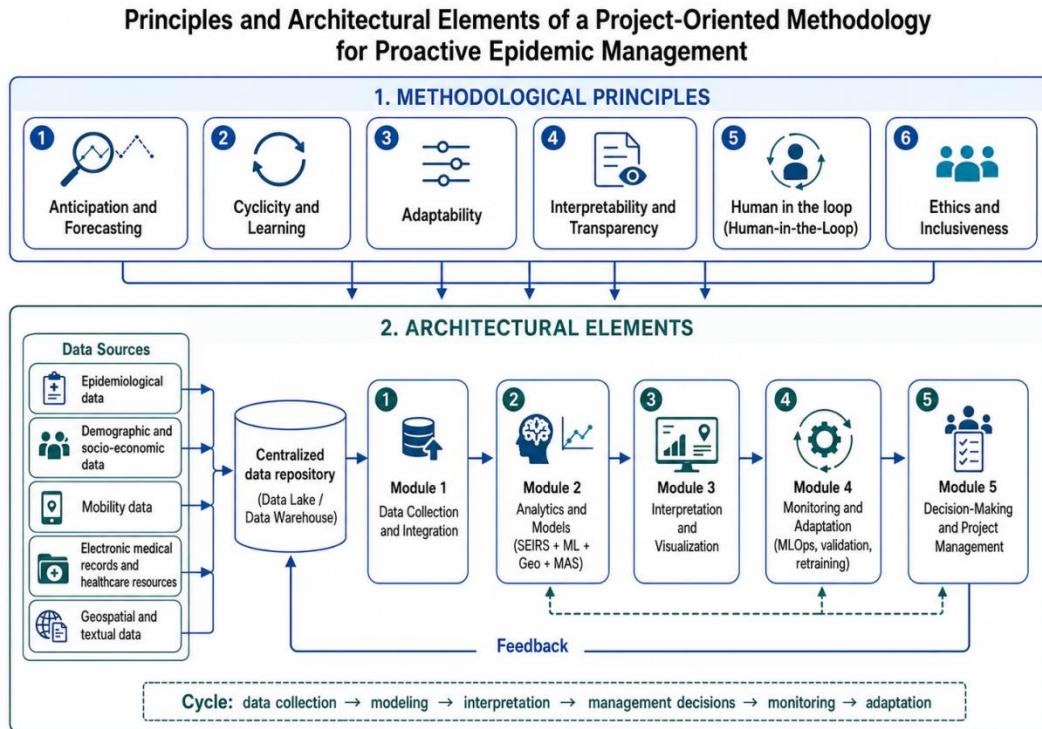


Fig. 1. Principles and Architectural Elements of Proactive Epidemic Management Based on SEIRS and AI

In Fig. 1, the proposed methodology is presented in the form of two interrelated levels: the level of methodological principles and the level of architectural elements. The principles level defines the conceptual requirements for building a proactive epidemic management system: anticipatory nature, cyclicity, adaptability, interpretability, expert involvement in the decision-making loop, and ethical use of data and artificial intelligence algorithms. The architectural level reflects the functional structure of the system, which ensures the implementation of these principles. It includes sources of epidemiological, demographic, geospatial, mobile, and medical data; a centralized data repository; modules for information collection and integration; analytics and modeling; interpretation and visualization of results; monitoring and adaptation; as well as a module for supporting management decision-making. Feedback between the modules ensures the cyclical updating of models, the adaptation of forecasts, and the adjustment of management interventions in accordance with changes in the epidemiological situation.

4.2. Formalization of the management methodology

Management methodology is defined as a system of principles, models, methods, mechanisms, procedures, and tools that determine the sequence and content of actions for the effective achievement of project or organizational goals.

Formally, the management methodology is presented as a tuple:

$$M_{PEM} = \langle P, M, T, R, C \rangle, \quad (1)$$

where M_{PEM} – project-oriented methodology for proactive epidemic management, P – principles and approaches (*Principles*), M – models, methods, and mechanisms (*Methods*), T – tools and techniques (*Tools*), R – roles and responsibilities (*Roles*), C – cycles and processes (*Cycles*).

Key characteristics define the internal structure of the methodology and the specifics of its practical implementation. They are divided into **structural** and **functional**.

Structural characteristics

Structural characteristics include:

Conceptual framework – management philosophy, principles, and strategic guidelines.

Process model – stages, phases, and sequences of management actions.

Toolkit – methods, techniques, and software tools.

Success criteria – metrics and performance evaluation indicators.

The functional characteristics of the methodology include:

Normativity – the definition of formalized rules for performing activities.

Systematic approach – ensuring the integrity and consistency of the approach.

Adaptability – the ability to adjust depending on conditions.

Reproducibility – ensuring the stability of results.

$$M_{\text{feasible}} = \{M \mid C_{\text{resource}}(M) \leq B, C_{\text{time}}(M) \leq T, R_{\text{risk}}(M) \leq R_{\text{max}}\}, \quad (3)$$

where B – budget constraints; T – time constraints; R_{max} – acceptable risk level.

Within the scope of this study, the components of management methodology are viewed not as a general classification of management approaches, but as functionally interrelated elements of proactive epidemic management. The conceptual level is defined by the principles of anticipatory action, adaptability, cyclicity, interpretability, expert involvement in the decision-making loop, and ethicality. The process level is implemented through an iterative cycle of data collection, modeling, interpretation, decision-making, monitoring, and adaptation. The instrumental level encompasses SEIRS models, machine learning methods, geospatial analysis, multi-agent modeling, and digital decision support systems. The organizational level is related to defining the roles of public health specialists, data analysts, managers, and stakeholders in the process of implementing anti-epidemic measures.

4.3. Iterative phases of predictive epidemic management

The conceptual framework of the methodology is implemented through four main iterative phases that form the "Predictive Epidemic Management Cycle".

Project management methodologies are classified as:

- strategic (high-level);
- tactical (operational);
- integrated (comprehensive).

Based on the degree of flexibility, the following are distinguished:

- rigid (waterfall, cascade);
- flexible (agile, adaptive);
- hybrid methodologies.

The effectiveness of a methodology is evaluated using an integral function:

$$E(M) = \sum_{i=1}^n w_i \cdot f_i(P, M, T, R, C, E), \quad (2)$$

where $E(M)$ – methodology effectiveness; w_i – factor weighting coefficients; f_i – influence functions of individual components; E – external environment.

The methodology must satisfy a set of constraints:

Phase 1. Data preparation and integration

Objective. Collection, cleaning, harmonization, and consolidation of heterogeneous data from distributed sources into a centralized repository.

Actions. Identification of sources, automated collection of streaming and static data, NLP text processing, geocoding.

Result. Creation of a consolidated, analysis-ready dataset.

Phase 2. Analysis, modeling, and forecasting

Objective. Identification of patterns, construction and training of predictive models, generation of development scenarios.

Actions. Feature engineering, exploratory data analysis (EDA), training and validation of machine learning model ensembles (time series, graph networks), simulation (e.g., agent-based modeling).

Result. Predictive indicators (R-number, incidence forecast, spread risk assessment), heat maps, "what-if" scenarios.

Phase 3. Interpretation and decision support

Objective. To transform technical forecasts into concrete, interpretable recommendations for management decision-making.

Actions. Visualization of results (dashboards, GIS maps), ranking and explanation of forecasts (SHAP, LIME), development of intervention options with an assessment of their potential impact.

Result. Reports for management, early warning signals, proposals for resource allocation (medicines, doctors, beds).

Phase 4. Implementation, monitoring, and adaptation

Goal. Implementation of selected measures, real-time assessment of their effectiveness, and model adjustment based on new data.

Actions. Integration with supply chain and HR management systems, continuous monitoring of key performance indicators (KPIs), feedback from practitioners, retraining of models. **Result.** Implemented anti-epidemic measures, updated and improved predictive models, closing the cycle.

4.4. Modular architecture and interconnections

The architecture is implemented as five key modules, each of which is responsible for a specific function and interacts with the others through clearly defined APIs and data flows (Table 1).

Table 1. Modular structure of the methodology and interaction between stages

Module	Main functions	Technologies/Approaches	Interaction with other modules
1. Data collection and integration module	Aggregation, cleansing, storage	Apache NiFi/Kafka, Airflow, ETL processes, NLP, APIs	Data source for Module 2. Receives quality feedback from Module 4
2. Analytics and modeling module	Generation of forecasts and analytical insights	Python (Scikit-learn, TensorFlow/PyTorch, Prophet, libraries for geospatial analysis), graph databases	Consumes data from Module 1. Provides forecasts and insights to Module 3. Receives retraining signals from Module 4
3. Interpretation and visualization module		Dash/Streamlit, Tableau, Power BI, GIS platforms (QGIS, ArcGIS), interpretation libraries (SHAP, ELI5)	Receives results from Module 2. Provides an interface and recommendations for users (Module 5)
4. Monitoring and adaptation module	Presenting results, explaining models	MLflow, Evidently AI, automated retraining pipelines, A/B testing	Monitors the productivity of Module 2. Provides quality metrics and data for adaptation. Interacts with Module 5 to assess the impact of decisions
5. Decision-making and project management module	Model validation, performance monitoring, model lifecycle management (MLOps)	Project management systems (Jira, Asana), collaboration and document management platforms	Uses insights from Module 3. Initiates actions, the results of which are returned as data to Module 1

The flow is cyclical. For example, a decision made in Module 5 (e.g., quarantine) generates new data (population mobility, reduced incidence), which enters the system via Module 1, is verified by Module 4, and

may lead to the recalibration of models in Module 2 and the updating of recommendations in Module 3.

Information flows: the data types, sources, and formats used in the methodology are presented in Table 2.

Table 2. Data types, sources, and presentation formats

Data type	Examples of sources	Nature	Output formats	Use in models
Epidemiological	Official reports from the Ministry of Health, global statistics (WHO, ECDC)	Structured time series	CSV, JSON, databases (SQL)	Direct training of forecasting models
Demographic and socioeconomic	Census data, State Statistics Service data, social indices	Structured, semi-structured	CSV, JSON, geo-JSON	Influencing factors, risk segmentation
Mobility data	Anonymous aggregated data from telecommunications providers, Google Mobility, Apple Mobility	Structured time series	CSV, JSON, geospatial formats	
Climatic and environmental	Weather services, satellite data	Structured time series, images	NetCDF, GRIB, CSV	Key predictor of spread dynamics
Text and news		Unstructured texts	Text, HTML, PDF	Accounting for seasonal and weather factors
Healthcare resource data	News, social media, scientific preprints (medRxiv), official announcements	Structured	Databases (SQL), API	Early warning detection, infodemic monitoring, semantic analysis
Genomic data	Registers of hospitals, drug stocks, vaccines, and personnel	Structured biological sequences	FASTA, specific bioinformatics formats	Resource allocation optimization, workload forecasting

4.5. The methodology's adaptability to different types of infectious diseases

The proposed methodology serves as a universal framework whose parameters can be adapted to specific disease categories at the level of data, models, and management processes, enabling the rapid development of an effective forecasting system for new epidemic threats.

1. Adaptation at the data level

For airborne infections (COVID-19, influenza), data on mobility and indoor population density take center stage.

For zoonoses (Ebola, West Nile fever) – climate data and the ranges of animal carriers.

For infections transmitted through water or food – data on water quality and sanitation.

2. Adaptation at the model level

R-number and time series. These are universal, but model parameters (lag, seasonality) are adjusted to the incubation period and dynamics of a specific disease.

Graph neural networks (GNN). The graph topology changes: for airborne epidemics – a graph of transportation links between cities; for STIs – a contact graph.

Agent-based modeling. Agent behavior rules, infection probability, and immunity parameters are customized to the transmission mechanism and pathogenesis of the disease.

3. Adaptation at the project management level

Response speed. For highly contagious diseases with a short incubation period, the "Monitoring-Adaptation" phase cycle is accelerated to a daily basis.

Critical resources. For some diseases, the key resources are respirators and mechanical ventilation; for others, they are antidotes or specific vaccines. The decision support module adapts its recommendations accordingly.

Communication strategy. The selection of channels and content for informing the public depends on the mode of transmission and the target risk groups.

Thus, the proposed methodology is not a rigid framework but a toolkit that allows for the rapid assembly of an effective predictive management system for any new infectious disease threat, minimizing the time required to develop one from scratch.

Integrating the SEIRS epidemiological model with AI and digitalization can enable proactive forecasting of tuberculosis outbreaks, potentially

reducing morbidity by predicting risks in high-burden regions, such as central and southeastern Ukraine.

Principles emphasizing data-driven forecasting, multi-stakeholder collaboration, and the ethical use of AI will help overcome war-related obstacles, although challenges such as access to data in occupied territories may limit effectiveness.

The use of AI for early detection and digital tools for real-time monitoring facilitates sustainable project management in the context of protracted crises, taking into account discussions regarding the accuracy of AI across diverse populations.

The methodology is based on proactive prediction, using AI to forecast and mitigate the spread of tuberculosis before it escalates. Key principles include.

- adaptability to the context of war-torn Ukraine;
- integration of digital health records for a continuous data flow;
- ethical considerations for vulnerable groups, such as displaced persons. Engaging stakeholders – from government health agencies to NGOs – ensures inclusive decision-making with a focus on equity to avoid exacerbating inequalities in regions such as Dnipro and Odesa.

5. An extended mathematical SEIRS–AI model: the case of tuberculosis control in Ukraine

5.1. Prerequisites for applying the model for tuberculosis in Ukraine

The proposed methodology conceptually aligns the compartments of the SEIRS epidemiological model with the phases of project management. This approach allows for the integration of mathematical modeling of the spread of infectious diseases with management decision-making processes in a digital environment.

Specifically, the correspondence between the SEIRS compartments and the project phases is defined as follows:

Susceptible – the planning and risk assessment phase, utilizing artificial intelligence tools for spatial mapping and analysis of risk factors.

Exposed – the prevention phase, implemented through digital tools for contact tracing, vaccination, and preventive interventions.

Infectious – the rapid response phase using intelligent triage, including CAD systems for medical image analysis.

Recovered – the monitoring, reintegration, and relapse prediction phase using AI.

Return to Susceptible – the long-term sustainability phase, which includes educational measures, revaccination, and strategic planning.

Within this framework, artificial intelligence and digitalization serve as the instrumental foundation of the methodology, enabling the analysis of large volumes of data, decision support, and secure information exchange, particularly through the use of blockchain technologies.

Given the current public health challenges in Ukraine – exacerbated by the war since 2022 and prior disruptions to the healthcare system due to the COVID-19 pandemic – the need for proactive management of tuberculosis (TB) control projects is extremely urgent.

The SEIRS epidemiological model is particularly well-suited for analyzing tuberculosis because it accounts for:

- the presence of latent infection;
- temporary immunity;
- the possibility of reinfection after recovery.

Tuberculosis is characterized by a prolonged latent phase, during which individuals may remain infected for years before the disease becomes active. Integrating the SEIRS model into a digital environment supported by artificial intelligence methods enables a shift from reactive to anticipatory management strategies, ensuring early intervention in high-risk scenarios, particularly during mass population displacement and damage to medical infrastructure.

Ukraine remains a country with a high tuberculosis burden, particularly of multidrug-resistant (MDR) forms. Statistical data indicate a complex and unstable epidemiological landscape. During the COVID-19 pandemic, TB incidence temporarily declined (for example, by 31% in 2020), but since 2021, an increase in rates has been observed – by 26% in 2022 and by 7% in 2023.

Central regions of Ukraine, particularly the Dnipropetrovsk and Kirovohrad regions, have experienced a sharp rise in incidence, linked to internal population displacement. Southeastern regions, including the Odesa and Donetsk regions, show consistently high rates, complicated by limited data availability from temporarily occupied territories. A significant increase in cases of MDR-TB has also been recorded among Ukrainian migrants abroad.

Within the proposed methodology, the principles of proactive epidemic management are considered as universal conceptual provisions that can be adapted to a specific infectious disease. For tuberculosis in Ukraine, their application is linked to the disease's prolonged latent phase, the risk of reinfection, regional heterogeneity of the epidemic process, the impact of internal migration, and limitations of the healthcare system under martial law.

In the context of tuberculosis management in Ukraine, these principles are applied as follows:

1. **Anticipation and forecasting** – the use of SEIRS models and artificial intelligence methods to predict the risks of tuberculosis spread before a sharp increase in incidence, particularly in regions with a high disease burden and a significant number of internally displaced persons.

2. **Cyclicity and learning** – organizing management as a continuous cycle of data collection, modeling, interpretation of results, decision-making, monitoring, and model adjustment based on new information.

3. **Adaptability** – the ability to adjust model parameters and management interventions in response to changes in the epidemiological situation, availability of medical resources, migration processes, and crisis conditions, particularly martial law.

4. **Interpretability and transparency** – ensuring that forecasts, scenarios, and recommendations are understandable to physicians, epidemiologists, policymakers, and other stakeholders, which is essential for building trust in decisions made using AI.

5. **Humans in the Decision-Making Loop** – using artificial intelligence as a support tool, not a replacement for expert judgment. Final risk assessments and the selection of interventions must involve public health specialists, clinicians, and policymakers.

6. **Ethics and inclusivity** – ensuring the protection of personal health data, preventing algorithmic bias, and addressing the needs of vulnerable populations, including displaced persons, patients with multidrug-resistant tuberculosis, and populations in regions with limited access to healthcare services. The integration of electronic health records, digital platforms, and geospatial data, as well as collaboration among government agencies, international organizations, NGOs, and local communities, and the monitoring of KPIs, are considered in this work not as separate principles but as organizational

and technological mechanisms for implementing the aforementioned methodological principles.

5.2. The Mathematical Basis of the SEIRS Model

The mathematical basis is the classical SEIRS (Susceptible–Exposed–Infectious–Recovered–Susceptible) model, which is an extension of the basic SIR model and accounts for the latent period and loss of immunity.

The transition diagram between states is as follows:

$$S \rightarrow E \rightarrow I \rightarrow R \rightarrow S. \quad (4)$$

The system of differential equations for the SEIRS model is described as follows:

$$\begin{aligned} \frac{dS}{dt} &= \delta R - \frac{\beta I}{N} S, \\ \frac{dE}{dt} &= \frac{\beta I}{N} S - \sigma E, \\ \frac{dI}{dt} &= \sigma E - \gamma I, \\ \frac{dR}{dt} &= \gamma I - \delta R. \end{aligned} \quad (5)$$

Key model parameters: β – infection rate; σ – rate of transition from the latent state to the infectious state; γ – recovery rate; δ – rate of immunity loss; $R_0 = \beta/\gamma$ – basic reproduction number.

Typical dynamics and possibilities for modeling interventions

The presence of the immunity loss parameter ($\delta > 0$) leads to cyclical waves of morbidity and the potential establishment of an endemic state. The SEIRS model allows for the assessment of the impact of various management interventions, including quarantine measures (reduction of β), vaccination, and social distancing.

5.3. Formalization of the hybrid SEIRS-AI model

Summarizing the results of previous studies presented in publications by the author and co-authors, it is advisable to extend the classical SEIRS compartmental model into a **hybrid multiscale intelligent model** that combines differential equations, geospatial representation, multi-agent systems, and machine learning methods.

Let the epidemiological state of the population at time t be described by the vector:

$$\mathbf{X}(t) = \{S(t), E(t), I(t), R(t)\}, \quad (6)$$

where the components correspond to the classical SEIRS compartments.

Unlike the traditional formulation, **the model parameters are treated as dynamic functions** determined by artificial intelligence methods:

$$\Theta(t, \mathbf{r}) = \{\beta(t, \mathbf{r}), \sigma(t), \gamma(t, \mathbf{r}), \delta(t)\}, \quad (7)$$

where $\mathbf{r}$ are geospatial coordinates.

Then the system of equations takes the form:

$$\begin{aligned} \frac{dS}{dt} &= \delta(t) R - \beta(t, \mathbf{r}) \frac{I}{N} S, \\ \frac{dE}{dt} &= \beta(t, \mathbf{r}) \frac{I}{N} S - \sigma(t) E, \\ \frac{dI}{dt} &= \sigma(t) E - \gamma(t, \mathbf{r}) I, \\ \frac{dR}{dt} &= \gamma(t, \mathbf{r}) I - \delta(t) R. \end{aligned} \quad (8)$$

Intelligent model parameterization

Estimation of parameters based on machine learning

The infection and recovery parameters are determined not by constants, but by the results of ML models:

$$\beta(t, \mathbf{r}) = f_\beta(\mathbf{Z}(t, \mathbf{r})), \quad (9)$$

where $\mathbf{Z}$ is a vector of socioeconomic, demographic, medical, and mobility features, and f_β is an ensemble of models (Random Forest, XGBoost, neural networks) tested in studies analyzing tuberculosis and COVID-19 [17–20, 22, 23].

Similarly, the recovery parameter – or the efficiency with which infected individuals transition to the R compartment – is not defined as a constant value, but rather as the output of a machine learning model:

$$\gamma(t, x) = F_\gamma^{ML}(X(t, x), H(t, x)), \quad (10)$$

where $\gamma(t, x)$ is the dynamic recovery parameter at time t for spatial unit x ; F_γ^{ML} is a machine learning function that estimates the parameter γ ; $X(t, x)$ is a vector of socioeconomic, demographic, medical, and mobility features; $H(t, x)$ – a vector of healthcare system characteristics, including access to medical services, staffing, diagnostic capacity, hospital workload, and access to treatment. This notation allows for formally accounting for the degradation or recovery of the healthcare system under conditions of war or other prolonged crises.

5.4. Geospatial, multi-agent, and RL extensions of the model

The population is divided into a set of agents:

$$A = \{a_1, a_2, \dots, a_n\}, \quad (11)$$

where each agent is described by a state:

$$a_i(t) = \langle x_i(t), y_i(t), c_i(t), h_i(t) \rangle \quad (12)$$

(geolocation, epidemiological status, behavioral characteristics, access to medical care).

Agent interactions occur in a **geospatial environment** such as **GeoCity**, which allows for the modeling of:

- local clusters of infection;
- the migration effect;
- various mobility restriction scenarios;
- as demonstrated in studies using GeoSEIR(D) and multi-agent modeling [21].

Integration of reinforcement learning

Control interventions are formalized as a reinforcement learning problem:

$$F = \langle S, A, P, R \rangle, \quad (13)$$

where S – current epidemiological situation;
 A – actions (screening, mobile clinics, quarantine);
 R – reward function:

$$R = -\alpha I(t) - \beta C(t) \quad (14)$$

(minimizing morbidity and intervention costs).

The optimal strategy π^* is determined via Q-learning or deep RL, which generalizes the approach presented in previous and submitted works

Generalized model structure

Finally, the methodology is described as a tuple:

$$\mathcal{M}_{AI-SEIRS} = \langle SEIRS, GeoCity, MAS, ML, RL \rangle, \quad (15)$$

where $SEIRS$ – basic epidemiological dynamics;
 $GeoCity$ – spatial environment; MAS – multi-agent interaction; ML – estimation of parameters and factors;
 RL – optimization of management decisions.

Table 3. Comparison of traditional (reactive) and proactive approaches to epidemic management

Aspect	Traditional (reactive) approach	A proactive approach based on SEIRS with AI
Focus	Response to detected cases	Prediction and prevention
Tools	Manual monitoring, paper records	AI-based analytics, digital platforms
Effectiveness	Delays in diagnosis and response	Reduction in R_0 through forecasting
Adaptability	Low, especially in crisis situations	High, with real-time updates
Cost	High due to late interventions	Reduced by 19–37% thanks to proactive measures

6. Discussion

6.1. Interpretation of the proposed methodology in the context of contemporary approaches

The proposed project-oriented methodology for proactive epidemic management expands upon traditional epidemiological approaches by integrating mathematical modeling, artificial intelligence, and project management principles. Unlike classic reactive

5.5. A comparison of reactive and proactive approaches

In regions with high epidemiological risk, the methodology calls for the implementation of pilot projects using mobile intelligent diagnostic units. The main challenges remain the limited availability of data from temporarily occupied territories and the need to use proxy sources, particularly satellite data. The expected outcome is a 20–30% reduction in morbidity through early detection and optimization of interventions (Table 3).

The comparative assessment of traditional and proactive approaches presented in Table 3 is of a general nature and is based on a combination of several data sets: results of an analysis of current literature on SEIRS/SEIR modeling, the application of artificial intelligence, geospatial, and multi-agent modeling in epidemiology [7–16, 21–28]; the results of the authors' previous studies on modeling the spread of tuberculosis and COVID-19 in a geospatial environment [17–23]; as well as a conceptual comparison of reactive and proactive management scenarios based on criteria such as management focus, tools, adaptability, effectiveness, and cost. Thus, Table 3 reflects not the results of a single clinical or field experiment, but a systematic comparison of expected management effects derived from literature analysis, preliminary modeling, and the logic of the proposed SEIRS-AI hybrid model.

The quantitative ranges of expected reductions in morbidity and costs are scenario-based and require further empirical validation during the pilot implementation of the methodology in regions with elevated epidemiological risk.

strategies, which focus primarily on recording and processing already identified cases of disease, the proposed approach focuses on risk prediction and the formulation of management decisions before critical scenarios occur.

Integrating the SEIRS model into the project cycle allows the epidemiological process to be viewed not merely as biological dynamics, but as a managed process with clearly defined phases of planning, prevention,

response, monitoring, and adaptation. This interpretation creates a common language among public health specialists, data analysts, and managers.

Advantages of the SEIRS-AI-Project Management Hybrid Approach

One of the key advantages of the proposed methodology is the shift from static parameters of epidemiological models to dynamically controlled parameters estimated using machine learning methods. This allows for the consideration of the influence of socioeconomic, demographic, spatial, and behavioral factors, which are traditionally difficult to formalize in classical models.

Combining SEIRS with geospatial and multi-agent approaches increases the model's sensitivity to local characteristics of infection spread, particularly migration patterns and uneven access to medical resources. The additional integration of reinforcement learning creates a basis for optimizing management interventions, taking into account the trade-off between epidemiological effectiveness and resource constraints.

6.2. Practical implications for public health systems

The study results indicate that the proposed methodology has significant potential for practical implementation in public health systems, particularly during prolonged crises. For Ukraine, this means the possibility of:

- early identification of high-risk regions;
- more rational allocation of limited resources (medical personnel, diagnostic tools, hospital beds);
- reducing delays in diagnosis and response;
- increasing the resilience of the healthcare system to external shocks.

The project-oriented nature of the methodology facilitates its integration into existing management structures, as it allows for the formalization of processes, roles, and responsibilities without the need for a complete overhaul of the organizational architecture.

A comparative analysis of traditional reactive and proactive approaches (table in Section 4) demonstrates a fundamental difference in management focus. Reactive approaches are typically characterized by long time delays, low adaptability, and rising costs during crises. In contrast, a proactive approach based on SEIRS and AI allows for a reduction in the basic reproduction number through early forecasting and preventive interventions, potentially reducing total costs by 19–37%.

These results are consistent with recent studies indicating the effectiveness of forecast-oriented strategies compared to late response, particularly for infections with long latent phases, such as tuberculosis.

6.3. Study limitations

Despite its advantages, the proposed methodology has a number of limitations. First, the quality of predictions depends heavily on the availability and completeness of data, which is problematic in conditions of military operations and the temporary occupation of certain territories. Second, the use of complex machine learning models increases demands on computational resources and staff qualifications.

In addition, there are ethical and legal challenges related to the processing of personal medical data and potential algorithmic biases, which require additional control and validation mechanisms.

Areas for further research

Further research could focus on expanding the methodology by integrating new data sources (e.g., satellite or social media signals), deepening the explainability mechanisms of models, and conducting large-scale pilot implementations in various regions. Particular attention should be paid to validating models in a cross-national context and adapting the methodology to other infectious diseases.

7. Conclusions and Recommendations

This article develops and justifies a project-oriented methodology for proactive epidemic management that combines classical SEIRS epidemiological modeling with methods of artificial intelligence, machine learning, geospatial analysis, and reinforcement learning. The proposed approach allows for a shift from reactive responses to infectious disease outbreaks to anticipatory risk management based on predictive analytical models.

It is shown that integrating the SEIRS model into the management project cycle formalizes the epidemiological process as a controlled system with clearly defined phases of planning, prevention, response, monitoring, and adaptation.

The use of dynamic model parameterization based on artificial intelligence allows for the consideration of the influence of socioeconomic, demographic, and spatial factors, as well as the degradation or recovery of the healthcare system under crisis conditions.

Extending the classical SEIRS model to a hybrid multilevel intelligent structure, which includes geospatial representation, multi-agent systems, and optimization of management interventions, forms a universal methodological framework for forecasting and preventing infectious diseases. The comparative analysis conducted indicates that a proactive approach based on SEIRS and AI has greater adaptability, potentially reduces the basic reproduction number, and allows for cost reduction through early and targeted interventions.

The methodology has particular practical value for Ukraine, where prolonged crisis factors, military operations, and limited data availability significantly complicate traditional epidemiological surveillance. The proposed approach creates a foundation for enhancing the resilience of the public health system and more effective use of limited resources.

Practical Recommendations

Based on the results obtained, the following recommendations are appropriate:

1. **Implementation of a project-oriented approach** to epidemic management at the level of national and regional public health programs, with a clear definition of phases, roles, and performance indicators.
2. **Integration of digital platforms and AI-based analytics** with existing epidemiological surveillance systems to ensure continuous real-time data collection, analysis, and interpretation.
3. **Use of the SEIRS model with dynamic parameterization** to forecast infections with long latent

phases, particularly tuberculosis, taking into account regional characteristics and migration patterns.

4. **Pilot implementation of the methodology** in regions with elevated epidemiological risk, followed by validation of results and scaling up at the national level.

5. **Strengthening cross-sectoral collaboration** between health authorities, think tanks, research institutions, and non-governmental organizations to ensure inclusive and ethical decision-making.

6. **Further research** should focus on the empirical validation of the proposed methodology, expanding datasets, improving the interpretability of AI models, and adapting the methodology to other infectious diseases.

Conflict of interest

The authors declare that they have no conflicts of interest, including financial, personal, authorial, or any other nature, that could influence the study or the results published in this article.

Funding

The study was conducted without financial support.

Data availability

Data will be provided upon reasonable request.

Use of artificial intelligence

The authors confirm that they did not use artificial intelligence technologies to write this paper.

References

1. Tang, L., Zhou, Y., Wang, L., Purkayastha, S., Zhang, L., He, J., Wang, F. and Song, P.X.-K. (2020), "A Review of Multi-Compartment Infectious Disease Models", *International Statistical Review*, Vol. 88, No. 2, pp. 462–513. DOI: <https://doi.org/10.1111/insr.12402>
 2. Brauer, F., Castillo-Chavez, C. and Feng, Z. (2019), *Mathematical Models in Epidemiology*, Springer, Berlin, 619 p. DOI: <https://doi.org/10.1007/978-1-4939-9828-9>
 3. Huang, J. and Morris, J.S. (2025), "Infectious Disease Modeling", *Annual Review of Statistics and Its Application*, Vol. 12, pp. 19–44. DOI: <https://doi.org/10.1146/annurev-statistics-112723-034351>
 4. Bushuyev, S., Bushuyev, D. and Bushuyeva, V. (2020), "Project Management During Infodemic of the COVID-19 Pandemic", *Innovative Technologies and Scientific Solutions for Industries*, Vol. 2, No. 12, pp. 13–21. DOI: <https://doi.org/10.30837/2522-9818.2020.12.013>
 5. Huliiev, N. (2024), "Choice of Machine Learning Models for Predicting the Development of Psychological Disorders in People with Hypothyroidism and Hyperthyroidism", *Innovative Technologies and Scientific Solutions for Industries*, Vol. 2, No. 28, pp. 76–85. DOI: <https://doi.org/10.30837/2522-9818.2024.2.076>
 6. Kiselev, I.N., Akberdin, I.R., Kolpakov, F.A. and Lashin, S.A. (2023), "Delay-Differential SEIR Modeling for Improved Modeling of Epidemic Dynamics", *Scientific Reports*, Vol. 13, Article 11874. DOI: <https://doi.org/10.1038/s41598-023-39052-9>
 7. Fong, S.J., Li, G., Dey, N., Crespo, R.G. and Herrera-Viedma, E. (2020), "Finding an Accurate Early Forecasting Model from Small Dataset: A Case of 2019-nCoV Outbreak", *International Journal of Interactive Multimedia and Artificial Intelligence*, Vol. 6, No. 1, pp. 132–140. DOI: <https://doi.org/10.9781/ijimai.2020.02.002>
-

8. Shinde, G.R., Kalamkar, A.B., Mahalle, P.N., Dey, N., Chaki, J. and Hassanien, A.E. (2020), "Forecasting Models for Coronavirus Disease (COVID-19): A Survey of the State-of-the-Art", *SN Computer Science*, Vol. 1, Article 197. DOI: <https://doi.org/10.1007/s42979-020-00209-9>
9. Venkatramanan, S., Lewis, B., Chen, J., Higdon, D., Vullikanti, A. and Marathe, M. (2018), "Using Data-Driven Agent-Based Models for Forecasting Emerging Infectious Diseases", *Epidemics*, Vol. 22, pp. 43–49. DOI: <https://doi.org/10.1016/j.epidem.2017.02.010>
10. Ye, Y., Pandey, A., Bawden, C., Sumsuzzman, D.M., Rajput, R., Shoukat, A., Singer, B.H., Moghadas, S.M. and Galvani, A.P. (2025), "Integrating Artificial Intelligence with Mechanistic Epidemiological Modeling: A Scoping Review of Opportunities and Challenges", *Nature Communications*, Vol. 16, Article 581. DOI: <https://doi.org/10.1038/s41467-024-55461-x>
11. Wu, J.T., Leung, K. and Leung, G.M. (2020), "Nowcasting and Forecasting the Potential Domestic and International Spread of COVID-19", *The Lancet*, Vol. 395, No. 10225, pp. 689–697. DOI: [https://doi.org/10.1016/S0140-6736\(20\)30260-9](https://doi.org/10.1016/S0140-6736(20)30260-9)
12. Cao, Q., Jiang, R., Yang, C., Fan, Z., Song, X. and Shibasaki, R. (2023), "MepoGNN: Metapopulation Epidemic Forecasting with Graph Neural Networks", In: *Machine Learning and Knowledge Discovery in Databases*, Lecture Notes in Computer Science, Vol. 13718, pp. 453–468. DOI: https://doi.org/10.1007/978-3-031-26422-1_28
13. Riley, S. (2007), "Large-Scale Spatial-Transmission Models of Infectious Disease", *Science*, Vol. 316, No. 5829, pp. 1298–1301. DOI: <https://doi.org/10.1126/science.1134695>
14. Keeling, M.J. and Rohani, P. (2008), *Modeling Infectious Diseases in Humans and Animals*, Princeton University Press, Princeton, NJ, 408 p.
15. Perez, L. and Dragicevic, S. (2009), "An Agent-Based Approach for Modeling Dynamics of Contagious Disease Spread", *International Journal of Health Geographics*, Vol. 8, Article 50. DOI: <https://doi.org/10.1186/1476-072X-8-50>
16. Ajelli, M., Gonçalves, B., Balcan, D., Colizza, V., Hu, H., Ramasco, J.J., Merler, S. and Vespignani, A. (2010), "Comparing Large-Scale Computational Approaches to Epidemic Modeling: Agent-Based Versus Structured Metapopulation Models", *BMC Infectious Diseases*, Vol. 10, Article 190. DOI: <https://doi.org/10.1186/1471-2334-10-190>
17. Vykylyuk, Y., Semianiv, I., Nevinskyi, D., Todoriko, L. and Boyko, N. (2024), "Applying Geospatial Multi-Agent System to Model Various Aspects of Tuberculosis Transmission", *New Microbes and New Infections*, Vol. 59, Article 101417. DOI: <https://doi.org/10.1016/j.nmni.2024.101417>
18. Semianiv, I., Todoriko, L., Vykylyuk, Y. and Nevinskyi, D. (2024), "Application of Geospatial Multi-Agent System for Simulation of Different Aspects of Tuberculosis Transmission", *Infusion & Chemotherapy*, Vol. 7, No. 1, pp. 9–17. DOI: <https://doi.org/10.32902/2663-0338-2024-1-9-17>
19. Vykylyuk, Y., Nevinskyi, D., Chopyak, V., Škoda, M., Golubovska, O. and Hazdiuk, K. (2023), "A Managerial Approach towards Modeling the Different Strains of the COVID-19 Virus Based on the Spatial GeoCity Model", *Viruses*, Vol. 15, No. 12, Article 2299. DOI: <https://doi.org/10.3390/v15122299>
20. Vykylyuk, Y., Levytska, S., Nevinskyi, D., Hazdiuk, K., Škoda, M., Andrushko, S. and Palii, M. (2023), "Decision-Tree and Ensemble-Based Mortality Risk Models for Hospitalized Patients with COVID-19", *System Research and Information Technologies*, Vol. 1, pp. 23–36. DOI: <https://doi.org/10.20535/SRIT.2308-8893.2023.1.02>
21. Vykylyuk, Y., Nevinskyi, D. and Hazdiuk, K. (2024), "Continuous-Discrete GeoSEIR(D) Model for Modelling and Analysis of Geo Spread of COVID-19", *Intelligence-Based Medicine*, Vol. 10, Article 100155. DOI: <https://doi.org/10.1016/j.ibmed.2024.100155>
22. Nevinskyi, D., Martjanov, D., Semianiv, I. and Vykylyuk, Y. (2025), "Studying the Relationship between Tuberculosis and Socioeconomic, Medical, and Demographic Factors in Ukraine", *System Research and Information Technologies*, Vol. 1, pp. 19–31. DOI: <https://doi.org/10.20535/SRIT.2308-8893.2025.1.02>
23. Nevinskyi, D., Semianiv, I., Vykylyuk, Y. and Yakymiv, A. (2025), "Intelligent Agents in the Fight against Global Threats", *Problems of Informatization and Control*, Vol. 81, No. 1, pp. 72–81. DOI: <https://doi.org/10.18372/2073-4751.81.20131>
24. Puterman, M.L. (2014), *Markov Decision Processes: Discrete Stochastic Dynamic Programming*, Wiley, Hoboken, NJ, 684 p.
25. Libin, P.J.K., Moonens, A., Verstraeten, T., Smet, F., Aleman, J., Vandamme, A.-M. and Nowé, A. (2021), "Deep Reinforcement Learning for Large-Scale Epidemic Control", In: *Machine Learning and Knowledge Discovery in Databases. Applied Data Science and Demo Track*, Lecture Notes in Computer Science, Vol. 12461, pp. 155–170. DOI: https://doi.org/10.1007/978-3-030-67670-4_10
26. Zino, L. and Cao, M. (2021), "Analysis, Prediction, and Control of Epidemics: A Survey from Scalar to Dynamic Network Models", *IEEE Circuits and Systems Magazine*, Vol. 21, No. 4, pp. 4–23. DOI: <https://doi.org/10.1109/MCAS.2021.3118100>
27. Kraemer, M.U.G., Tsui, J.L.-H., Chang, S.Y., Lytras, S., Khurana, M.P. et al. (2025), "Artificial Intelligence for Modelling Infectious Disease Epidemics", *Nature*, Vol. 638, No. 8051, pp. 623–635. DOI: <https://doi.org/10.1038/s41586-024-08564-w>
28. World Health Organization (2021), *Global Strategy on Digital Health 2020–2025*, World Health Organization, Geneva. ISBN: 978-92-4-002092-4, available at: <https://www.who.int/publications/i/item/9789240020924>

Received (Надійшла) 24.02.2026

Accepted for publication (Прийнята до друку) 21.05.2026

Publication date (Дата публікації) 27.06.2026

Відомості про авторів / About the Authors

Невінський Денис Володимирович – кандидат технічних наук, доцент, Національний університет "Львівська політехніка", Львів, Україна;

Denys Nevinskyi – Candidate of Technical Sciences, Associate Professor, Lviv Polytechnic National University, Lviv, Ukraine;

e-mail: nevinskyi90@gmail.com

ORCID ID: <https://orcid.org/0000-0002-0962-072X>

Виклюк Ярослав Ігорович – доктор технічних наук, професор, Національний університет "Львівська політехніка", Львів, Україна;

Yaroslav Vyklyuk – Doctor of Technical Sciences, Professor, Lviv Polytechnic National University, Lviv, Ukraine;

e-mail: vyklyuk@ukr.net

ORCID ID: <https://orcid.org/0000-0003-4766-4659>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=6504272188>

ПРОЄКТНО-ОРІЄНТОВАНА МЕТОДОЛОГІЯ ПРОАКТИВНОГО УПРАВЛІННЯ ЕПІДЕМІЯМИ НА ОСНОВІ ШТУЧНОГО ІНТЕЛЕКТУ Й МОДЕЛЕЙ SEIRS

Предметом дослідження в статті є процеси проактивного управління епідеміями як складними соціально-медициніми системами в умовах високої невизначеності, обмеженості ресурсів і динамічних зовнішніх впливів, зокрема із застосуванням математичного епідеміологічного моделювання та інтелектуальних цифрових технологій. **Мета роботи** полягає в створенні проєктно-орієнтованої методології управління епідеміями на основі інтеграції класичної моделі SEIRS зі штучним інтелектом, машинним навчанням і цифровими системами підтримки прийняття управлінських рішень, що забезпечує перехід від реактивного реагування до проактивного прогнозування й запобігання спалахам інфекційних захворювань. У статті необхідно виконати такі **завдання**: сформулювати концептуальну й процесну модель проактивного управління епідеміями в межах проєктного підходу; розробити інтегровану архітектуру збору, аналізу та інтерпретації епідеміологічних показників; адаптувати компартментну модель SEIRS до умов динамічної параметризації з використанням методів машинного навчання; обґрунтувати застосування геопросторових і мультиагентних підходів для врахування неоднорідності поширення інфекцій; визначити роль інтелектуальних методів оптимізації управлінських втручань у проєктному циклі протидії епідеміям. Упроваджено такі **методи**: математичне моделювання епідеміологічних процесів на основі систем диференціальних рівнянь SEIRS; машинне навчання й ансамблеве прогнозування для динамічного оцінювання параметрів моделі; геопросторовий аналіз і мультиагентне моделювання для відтворення мобільності населення й просторових ризиків; навчання з підкріпленням для оптимізації управлінських рішень; проєктно-орієнтовані підходи до організації та адаптації управлінських процесів. Досягнуто таких **результатів**: сформульовано принципи проактивного управління епідеміями на основі інтеграції епідеміологічного моделювання й штучного інтелекту; запропоновано проєктно-орієнтовану методологію, що поєднує цикл управління проєктом із фазами моделі SEIRS; розроблено узагальнену структуру гібридної інтелектуальної моделі, здатної адаптуватися до різних типів інфекційних захворювань, зокрема туберкульозу та COVID-19; обґрунтовано можливість підвищення ефективності раннього виявлення ризиків і раціонального розподілу ресурсів у кризових умовах. **Висновки**: застосування запропонованої проєктно-орієнтованої методології та інтегрованої моделі SEIRS зі штучним інтелектом створює підґрунтя для переходу до проактивного управління епідеміями, підвищує адаптивність систем громадського здоров'я та може бути використане як універсальний методологічний каркас для розроблення інтелектуальних систем прогнозування й запобігання епідеміям у складних і нестабільних середовищах.

Ключові слова: проактивне управління епідеміями; проєктно-орієнтована методологія; штучний інтелект; модель SEIRS; машинне навчання; цифрове здоров'я; епідеміологічне прогнозування; системи підтримки прийняття рішень.

Бібліографічні описи / Bibliographic descriptions

Невінський Д. В., Виклюк Я. І. Проєктно-орієнтована методологія проактивного управління епідеміями на основі штучного інтелекту й моделей SEIRS. *Сучасний стан наукових досліджень та технологій в промисловості*. 2026. № 2 (36). С. 95–108. DOI: <https://doi.org/10.30837/2522-9818.2026.2.095>

Nevinskyi, D., Vyklyuk, Y. (2026), "Project-oriented methodology for proactive epidemic management based on artificial intelligence and SEIRS models", *Innovative Technologies and Scientific Solutions for Industries*, No. 2 (36), P. 95–108. DOI: <https://doi.org/10.30837/2522-9818.2026.2.095>

Vladimir Pevnev, Oles Yudin

CLASSIFICATION AND EXPERIMENTAL STUDIES OF THE FACTORIZATION ALGORITHMS EFFICIENCY

The subject of the study is factorization algorithms, namely experimental testing of the efficiency of modern algorithms for integer factorization, identifying patterns between factorization algorithms and the size of the numbers being factorized in terms of the time required to perform the factorization operation. **The purpose of the article is to** analyze the performance of factorization algorithms using various composite numbers, in particular medium-sized numbers (~ 20 digits), which allows evaluating their efficiency and execution time. In the course of the research, the following **tasks** must be performed: classify modern factorization algorithms, experimentally verify the efficiency of modern Pollard factorization algorithms (p and $p-1$), the elliptic curve method, the Dixon algorithm, and the quadratic sieve. To achieve the goal, general scientific **methods** were used: analysis of the subject area and mathematical apparatus, set theory, numbers and fields, planning, and experimental research. **Results achieved.** Experimental studies have shown that Pollard's algorithms are effective for numbers with small divisors, but lose performance as the size of the numbers increases. The elliptic curve method has proven its usefulness in finding medium-sized divisors and has shown better scalability compared to classical stochastic methods. The Dixon algorithm, despite its relative simplicity of implementation, has demonstrated stochastic fluctuations in execution time, which limits its practical value in scenarios with strict time constraints. The most stable and predictable results were achieved for the quadratic sieve, which confirmed its suitability for factoring medium-sized numbers and provided the smallest spread of time values under conditions of multiple runs on identical data sets. **Conclusions.** The experiments fully confirm the theoretical expectations regarding the performance of the methods under study. The results indicate that simple algorithms (Pollard, elliptic curve method) are appropriate for preprocessing and detecting weak keys, while the quadratic sieve is the optimal choice for factoring medium-sized numbers. For large RSA modules, it is practical to use more complex algorithms, such as the general number field sieve. Further development of factorization algorithms involves parallelizing the factorization process and developing algorithms that use screening of impossible solutions.

Keywords: integer factorization; cryptography; Pollard method; elliptic curve method; quadratic sieve; Dixon's algorithm; algorithm efficiency.

Introduction

The qualitative leap in information technology that we have witnessed over the past decade has significantly changed approaches to the organization of technical and technological processes in industry. First of all, it is worth noting the use of artificial intelligence in the organization of production processes. This includes the creation of enterprises in which most of the processes are performed by robotic platforms, research and creation of new materials with specified properties, product control, ranging from food and medicines to spacecraft and nuclear power plants, etc.

Along with the introduction of new technologies, the number of accompanying documents is also growing, including new technologies, technological maps, reports on research work, accounting documentation, and information about individuals involved in various projects. All this data is stored both on hard media and in electronic form. The electronic form allows for the

rapid transfer of documents between executors and real-time coordination, but at the same time, the confidentiality of the information being transferred and stored must be ensured. The most common way to ensure confidentiality is encryption. This method allows information to be transmitted over open communication channels, stored in cloud storage, and read only with a key.

The encryption process itself is based on the use of cryptographic systems. There are two large groups of these systems: symmetric and asymmetric. While the cryptographic strength of symmetric encryption systems is based on key secrecy, asymmetric systems rely on solving problems that are difficult to solve. In the first case, a single key is used for encryption and decryption, while in the second, two keys are used, one for encryption and the other for decryption. One such problem is the factorization of large numbers.

Factorization, or the decomposition of integers into prime factors, is one of the fundamental problems of

number theory and cryptography. It underlies the security of modern cryptographic protocols, in particular the RSA system. The security of RSA is based on the computational complexity of factoring large composite numbers, which are the product of two large prime numbers.

In modern research, factoring is considered a key element of cryptanalysis, as well as a training and research platform for testing the algorithmic efficiency of various calculation methods. There is a wide range of factorization algorithms. These include classical methods such as Fermat's method, Pollard's method, and Dixon's algorithm, as well as modern lattice methods and algorithms based on elliptic curves. Each of these algorithms has its own limitations and areas of application, which determine their practical suitability for factoring numbers of different sizes.

The aim of this study is to investigate algorithms in terms of their suitability for RSA systems. For each algorithm, the execution time is evaluated and a comparison of efficiency is performed on a set of composite numbers of different sizes.

Relevance of the study

Number factorization, i.e., the decomposition of integers into prime factors, is one of the oldest mathematical problems that has attracted the attention of many mathematicians over the centuries. Eratosthenes (276–194 BC) was the first known mathematician to develop a simple algorithm for finding prime factors. It is called Eratosthenes' sieve and enumerates all prime numbers smaller than a given integer N . Other prominent mathematicians who have made various contributions to factorization are Fermat (1607–1665) and Euler (1707–1783) [1].

The development of computer technology in the mid-20th century allowed mathematicians to create new and effective factorization algorithms. Significant improvements in the theory and practice of factorization were driven by the development of the well-known RSA public-key encryption algorithm.

The basic concept of asymmetric cryptographic algorithms is the use of two different keys: one for encrypting a message and the other for decrypting it [2, 3].

A key feature of such algorithms is the irreversibility of the encryption procedure, even with the public key available [4]. That is, knowing the

public encryption key and the encrypted text, it is impossible to recover the original message. Decryption is only possible with the use of the corresponding private key. Thus, the public key can be publicly available – its disclosure does not pose a security threat, since it does not allow the encrypted information to be read. The security of the algorithm is ensured by the complexity of the inverse operation – determining the original prime numbers by their product [5]. As a result, factorization of numbers is the mathematical basis of modern asymmetric cryptography.

As information security becomes increasingly important, cryptographic mechanisms have become an integral component of every information system. The use of cryptographic mechanisms and protocols can solve many security problems (confidentiality, integrity control, authenticity, and non-repudiation of information). The development of cryptographic technologies has a direct impact on the economic, sociological, and political aspects of society as a whole. The development and use of factorization algorithms, analysis, and evaluation of their effectiveness allow new requirements to be developed for future versions of the RSA system, in particular for the size of the system module and the keys involved in the process of encrypting and decrypting information.

Thus, the widespread use of cryptography today increases the importance of developing cryptographic algorithms, and research into the features of applying number factorization algorithms for RSA is relevant.

Literature review

The operation of factorization algorithms is considered in the works of contemporary Ukrainian and foreign authors [6–8].

In [9], the dependence of the factorization of composite numbers on the time required to decompose this composite number using trial division, Fermat's algorithm, Pollard's p -method, Pollard's $(p-1)$ -method, and Brent's method is compared. However, the authors do not take into account the conditions for using algorithms to solve cryptography problems.

In [10], a modification of Fermat's algorithm using Euler's function is considered. However, the author does not investigate the effectiveness of applying the proposed algorithm specifically for the implementation of encryption algorithms.

In [11], the authors investigated three factorization algorithms: the divisor search algorithm, Fermat's algorithm, Pollard's p -method, and Shor's algorithm. The research is interesting in that emulating a quantum algorithm on a classical computer does not provide significant advantages. At the same time, the authors did not consider other algorithms used for the RSA system.

In [7, 12], the authors investigated the effect of the number of threads on the speed of execution of the division verification algorithm and their energy efficiency. However, the authors also did not take into account the peculiarities of the RSA system in their research.

Thus, modern research considers various factorization algorithms. However, in most cases, such studies do not analyze the possibility of using them for the RSA system.

The purpose of this study is to analyze the performance of factorization algorithms on a home computer using various composite numbers, including medium-sized numbers (~ 20 digits), to evaluate their efficiency and execution time. Particular attention is paid to algorithms that can be used to analyze the security of RSA systems with weak keys, as well as their educational value for demonstrating the principles of factorization.

Presentation of the main material

The factorization of natural numbers is a basic operation in arithmetic, which consists of representing a number $n \in \mathbb{N}$ as a product of prime factors. According to the fundamental theorem of arithmetic, every natural number $n > 1$ can be represented uniquely (up to the order of the factors) as a product of prime numbers.

The formalized formulation of the problem is as follows.

Let there be a given natural number
 $n \in \mathbb{N}, n > 1$

The general problem of factorization is to find a set of prime numbers $\{p_1, p_2, \dots, p_k\} \subset P$ and corresponding exponents $\{a_1, a_2, \dots, a_k\} \subset P$ such that the following equality holds

$$n = p_1^{a_1} \cdot p_2^{a_2} \cdot \dots \cdot p_k^{a_k}, \text{ where } p_i < p_{i+1}, \forall i$$

p_i are prime numbers, $a_i \geq 1$, $k \geq 1$ and the layout is unique with precision down to the rearrangement of factors [6, 7, 9].

The development of effective factorization algorithms allows for the improvement of both protection and cryptanalysis tools.

There are a large number of factorization algorithms.

1. Classical methods (based on quadratic congruences) [1, 9, 11].

These methods are based on finding numbers a, b , such that $a^2 \equiv b^2 \pmod{n}$.

a is not equal $\pm b$ by the module n .

Then the expression can be rewritten as

$$(a^2 - b^2) \equiv 0 \pmod{n}$$

or

$$(a - b)(a + b) \equiv 0 \pmod{n}.$$

It is possible to find a non-trivial divisor of the number n by calculating the greatest common divisor (GCD): $GCD(a - b, n)$ and $GCD(a + b, n)$. At least one of these GCDs will be a non-trivial divisor of n , which is the goal of factorization.

2. Stochastic and probabilistic methods [13].

Stochastic, or probabilistic, factorization methods do not guarantee finding a divisor within a specified time, but they have a high probability of success. Their effectiveness depends on the properties of the number itself or its unknown divisors. The basic idea is to use randomness to "stumble upon" a divisor. Unlike deterministic methods, they do not try all possible options, but make "lucky guesses".

3. Sieve methods [1, 6, 10, 11, 13].

Sieve methods are the most powerful known algorithms for factoring large integers. Their name comes from the main stage, which resembles the famous Sieve of Eratosthenes for finding prime numbers.

The general principle of these methods is to reduce the factorization problem to two basic steps:

– Sieving stage – efficient search for special "smooth" numbers (i.e., numbers that can be factored into small prime numbers).

– Linear algebra stage – combining the numbers found to construct a congruence of squares.

The ultimate goal, as in many other algebraic methods, is to find two different numbers a and b such that:

$$a^2 \equiv b^2 \pmod{n}$$

where n is a number that needs to be factored. If such a pair is found, the divisors can be easily calculated using $GCD(a - b, n)$ and $GCD(a + b, n)$.

The key idea that makes these methods so fast is that instead of checking millions of numbers one by one for "smoothness", the sieve allows you to do this en masse.

4. Algebraic methods [4, 6, 7, 13].

Algebraic factorization methods are based on using the properties of numerical structures, equations, and forms to find non-trivial divisors of composite numbers. Unlike classical or stochastic algorithms, which work mainly through brute force or random iterations, algebraic methods apply strict mathematical patterns that reduce the search space for factors.

The basic idea of such methods is to represent a given number nmn in the form of algebraic relations, such as equations or quadratic forms, and find solutions that lead to a non-degenerate divisor of the number.

5. Quantum algorithms [10].

Shore's algorithm is the most efficient theoretically known method for factorization on quantum computers. It performs decomposition in polynomial time, which threatens modern cryptography.

The most well-known algorithms of different classes are shown in Fig. 1.

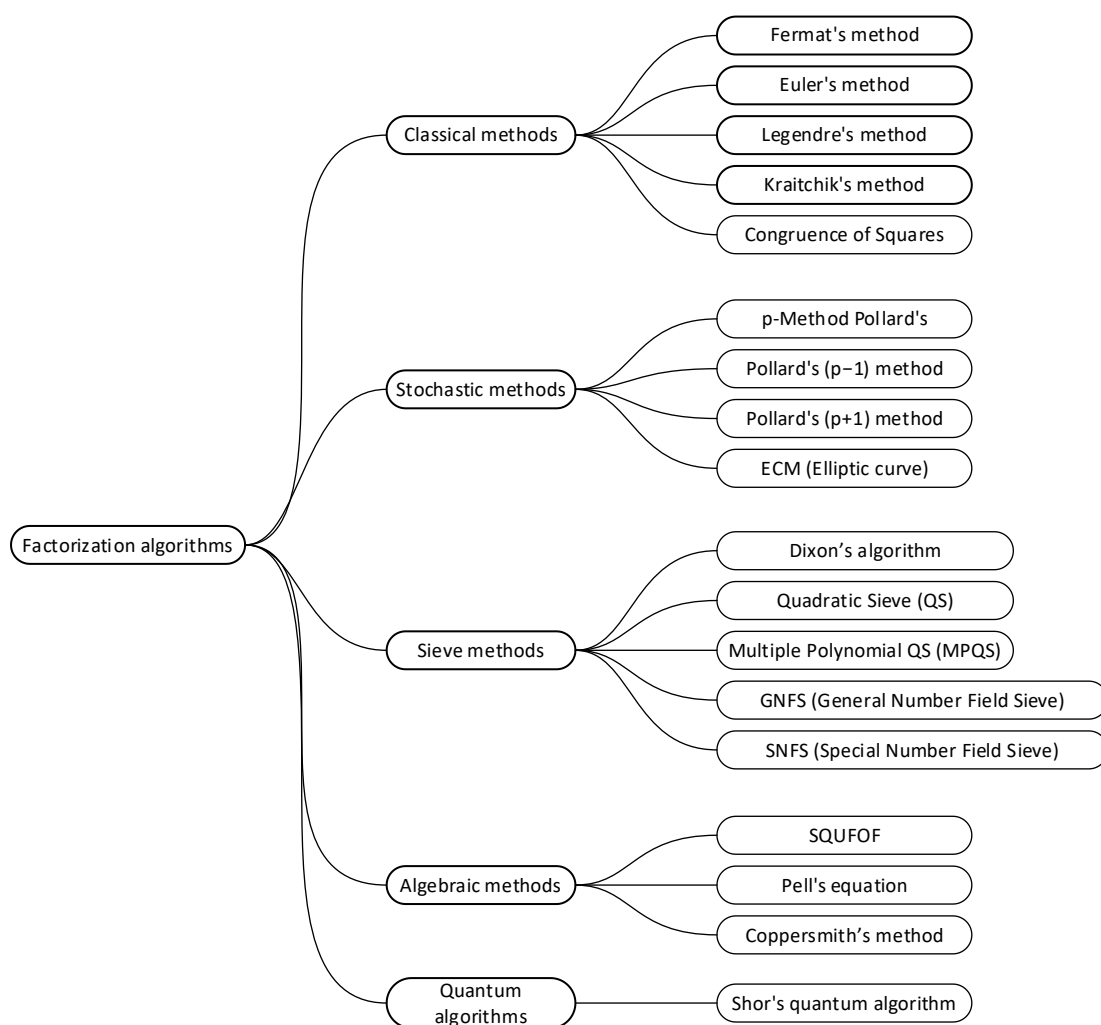


Fig. 1. Classification of factorization algorithms

Despite the significant variety of factorization algorithms, their use in RSA systems is somewhat limited.

Algorithms of a special structure, which include Fermat's method, Euler's method, Pollard's $(p-1)$ method, Williams' $(p+1)$ method, Bent's method,

Pell's equation, SQUFOF, Legendre's method, Legendre-Gauss's method, and Kraitchik's method, are only effective in cases where the divisors have a special form (for example, $p-1$ or $p+1$ are easily factorized, or the numbers are close to each other). As a result, they are not applicable to RSA systems, since prime

numbers for RSA are chosen randomly and specifically checked for the absence of weak structures.

Stochastic and probabilistic methods are suitable for finding "small" prime divisors. ECM, in particular, is the most effective method for finding small prime factors in large numbers. For RSA systems, such methods can only be effective in the case of incorrect key generation (for example, when one of the prime factors is significantly smaller than the other). In correctly generated RSA keys, prime numbers have the same length, so these methods are ineffective.

Lattice-type algorithms are among the most powerful classical algorithms. QS is effective for medium-sized numbers (up to ~ 110 digits). For very large numbers, such as RSA-1024, the only known practical method is GNFS. All successful attacks on RSA keys with a length of 768 bits and less are based on GNFS. Algebraic methods are used to construct a factor base at an auxiliary stage in the implementation of lattice methods (QS, GNFS).

The only known algorithm that makes RSA fundamentally vulnerable is Shor's quantum algorithm. However, its practical application requires a high-powered quantum computer, while emulation of this algorithm on a classical computer does not demonstrate high efficiency.

Therefore, lattice methods remain the most relevant methods for systems [14].

To compare factorization algorithms, this study chose methods that can find RSA module divisors in cases of poor key generation or demonstrate the principle of RSA key construction. All algorithms were implemented in Python, and Google Colab cloud platform was chosen as the execution environment. The experimental study was conducted on a sequence of numbers in the range from 4 to 1,000,000 with an interval of 10,000. This approach ensured equal execution conditions for each of the algorithms and enabled a correct comparison of their effectiveness.

Pollard's Rho method (Pollard's p -method) was chosen as the first algorithm for the study. It is usually used to find small divisors. If the RSA module was built with a "bad" prime number (for example, not a random number, but a number that has a certain structure), this method can quickly find it. Therefore, it is used to check RSA keys for weakness.

The idea of the algorithm is based on the birthday paradox and the idea of cycles in pseudo-random number sequences. First, a sequence of numbers is generated.

To do this, an arbitrary initial number x_0 is selected. The sequence of numbers $x_1, x_2, x_3, \dots$ is generated using a pseudorandom function $x_1, x_2, x_3, \dots$, where N is the number to be decomposed, and c is a small constant (for example, $c = 1$) [8, 15].

Since the sequence is generated modulo N , it is cyclic. The algorithm searches for repetitions not in the sequence modulo N itself, but in the sequence modulo p , where p is a prime factor of N . Due to the birthday paradox, the probability that two numbers in the sequence will be equal modulo p is significantly higher than modulo N . For an efficient search for a match, Floyd's cycle detection algorithm, also known as the "tortoise and hare method", is used.

The formalized algorithm can be written as follows:

1. Select random initial values x_0 and $y_0 = x_0$.

2. Select constant $c \neq 0$.

3. Repeat:

– $x_i = (x_{i-1}^2 + c) \bmod N$

– $y_i = (y_{i-1}^2 + c) \bmod N$, and then again

$y_i = (y_i^2 + c) \bmod N$ (i.e. $y_i = x_{2i}$)

– Calculate $d = \text{GCD}(|y_i - x_i|, N)$.

– If $d > 1$ and $d \neq N$, then d is a simple multiplier of the number N . The algorithm is complete.

– If $d = N$, this means that we have "jumped" over the multiplier. In this case, we need to start over with different initial values x_0 or c .

– If $d = 1$, continue iteration.

The complexity of this algorithm is approximately $O(N^{1/4})$.

The results of applying the algorithm are shown in Fig. 2.

Analysis of the results shows that the execution time of Pollard's p -method is on the order of microseconds, varying from approximately 5.8×10^{-7} to 4.9×10^{-5} seconds. Most measurements are concentrated near the lower limit, indicating the stability and relatively high efficiency of the method for the selected input data. At the same time, there are isolated peak values that exceed the average level by more than an order of magnitude. This may be due to the specifics of the random component of the algorithm, which sometimes leads to longer searches for non-trivial divisors.

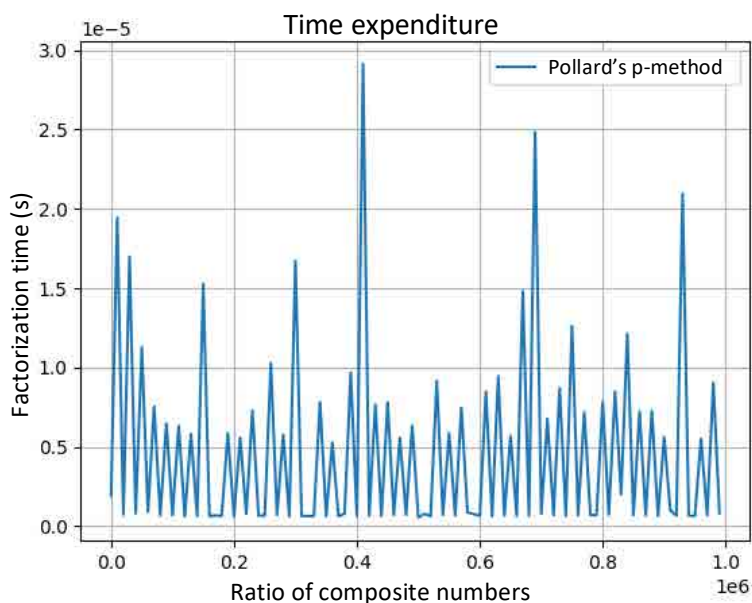


Fig. 2. Results of testing the Pollard's p -method algorithm

Thus, Pollard's p -method demonstrates high performance for factorization problems on small and medium-sized numbers, but is characterized by uneven execution time due to its probabilistic nature. This makes it suitable for practical application in conditions where execution time fluctuation is acceptable, but requires caution when estimating guaranteed computational costs.

The second algorithm considered is Pollard's $(p-1)$ method. This is a specialized factorization method that is effective for decomposing a number N if N has a prime factor p for which the number $p-1$ is smooth. A smooth number is a number whose prime factors are all less than a certain specified limit.

The algorithm is based on Fermat's little theorem, which states that if p is a prime number and a is an integer not divisible by p , then $a^{p-1} \equiv 1 \pmod{p}$. [8, 16]

If $p-1$ is a smooth number, it can be decomposed into small prime factors. The algorithm uses this fact to find a multiple of $p-1$ (or $k(p-1)$).

The formalized algorithm can be written as follows:

1. Choose any initial number a , for example, $a = 2$.
2. Select the smoothness threshold B . This number determines how "smooth" $p-1$ should be.
3. Calculate K as the product of all powers of prime numbers less than or equal to B .

4. Calculate $a^K \pmod{N}$ using the algorithm for rapid exponentiation by modulus.

5. Calculate $d = \text{GCD}(a^K - 1, N)$.

6. If $d > 1$ and $d \neq N$, then d is a simple factor of the number N . The algorithm has been successfully completed.

7. If $d = 1$, then the chosen boundary B is too small, and $p-1$ is not smooth with respect to B . It is necessary to increase B and repeat.

8. If $d = N$, then the chosen number a is not suitable. This can also happen if N has several factors p_i with smooth $p_i - 1$. You need to start again with another a .

The time complexity of this algorithm for factoring a sequence of numbers is shown in Fig. 3.

The results show that the execution time of Pollard's $(p-1)$ method varies widely: from values of the order of $10^{-6} - 10^{-5}$ seconds to individual peaks reaching 0.015–0.020 seconds. Most of the results are concentrated in the microsecond range, confirming the high efficiency of the algorithm for most of the numbers processed. At the same time, the presence of significant deviations from the average level indicates a substantial dependence of the execution time on the structure of the factorized number, in particular on the properties of its prime divisors and the order of elements in the multiplicative group modulo.

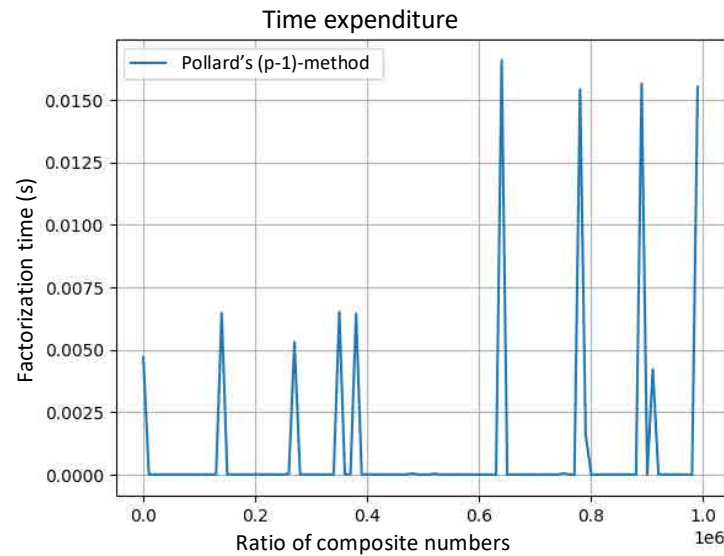


Fig. 3. Results of algorithm $(p-1)$ – Pollard's method testing

The Pollard $(p-1)$ method demonstrates asymmetric performance. In most cases, it works very quickly, but for some numbers it may require significantly more time.

The next algorithm under consideration is the Elliptic Curve Method (ECM). This algorithm is used to find non-trivial factors of integers. It is particularly effective for finding relatively small prime factors of large numbers, and its performance depends on the size of the smallest prime factor, not the size of the number itself, making it one of the most powerful modern methods for finding medium-sized factors.

ECM combines the ideas of Pollard's $(p-1)$ method and the group structure of points on an elliptic curve. Instead of performing calculations with numbers modulo N , ECM performs addition operations on points on an elliptic curve modulo N [11].

The algorithm selects a random elliptic curve E and a point P on it. An elliptic curve is generally described by the equation $y^2 = x^3 + ax + b \pmod{N}$.

The operation of "addition" can be defined on the points of an elliptic curve. This operation is associative, and the points form an Abelian group. Similar to the $(p-1)$ method in ECM, point P is "added" to itself K times, and $K \cdot P$ is calculated. The calculation of $K \cdot P$ is performed modulo N .

The main advantage of ECM is that the order of the group of points on an elliptic curve modulo p (where p is a prime factor of N) can be different for different curves.

The formalized algorithm can be written as follows:

1. Select the smoothness limit B_1 .
2. Select a random elliptical curve E and a point P on it.
3. Calculate $Q = K \cdot P \pmod{N}$, where K – the product of all prime numbers less than or equal to B_1 , in the appropriate degrees.
4. When calculating Q , each step checks whether an error occurs when dividing by $0 \pmod{N}$. If the GCD of the denominator and N is greater than 1, the factor is found.
5. If no multiplier is found, you can increase the limit B_1 or try another curve and point P .
6. The computational complexity of ECM depends on the size of the smallest prime factor p of the number N and is approximately equal to $O\left(e^{\sqrt{2 \log p \log \log p}} \cdot \log^2 N\right)$. The complexity is sub-exponential with respect to p , not N . This means that the smaller the prime factor, the faster ECM will find it.

The probability of successfully finding a factor increases with the number of attempts (i.e., with the number of elliptic curves used).

The time required for factorization using this algorithm is shown in Fig. 4.

The results obtained demonstrate significant variability in the execution time of the ECM method. The values range from about 10^{-5} seconds to tens of seconds, with both relatively fast factorizations (e.g., in the millisecond range) and significant peak

loads recorded – in particular, 6.58 seconds and even 14.26 seconds. This uneven performance is explained by the mathematical nature of ECM. Its effectiveness depends significantly on the choice of random elliptic curves and the properties of the desired prime divisor.

If the parameters are "favorable", the algorithm quickly finds the factor, but in unfavorable conditions, a large number of iterations are required, resulting in exponential peaks in execution time.

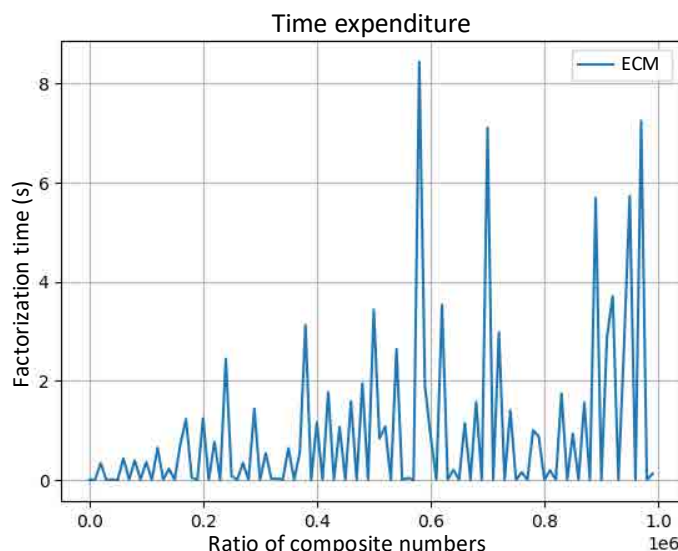


Fig. 4. Results of testing the ECM algorithm

Quadratic Sieve (QS) is one of the most efficient classical algorithms for factoring large integers. It belongs to the class of algorithms based on the idea of finding congruence of squares and is the best choice for factoring numbers up to 100–120 digits.

The algorithm is based on the idea of finding two numbers x and y such that $x \not\equiv y \pmod{N}$ and $x^2 \equiv y^2 \pmod{N}$. If such a pair is found so $x^2 - y^2 \equiv 0 \pmod{N}$, then $x^2 - y^2 \equiv 0 \pmod{N}$ is divisible by N . In this case, with high probability $\text{GCD}(x - y, N)$ or $\text{GCD}(x + y, N)$ will give a non-trivial factor of the number N [16].

The method for finding such a congruence consists in searching for a set of numbers z_i such that $z_i^2 \pmod{N}$ is a smooth number (i.e., it can be factorized from a small "factor base").

The computational complexity of the quadratic sieve is sub-exponential and is expressed by the formula $L_N[1/2, c] = e^{(c+o(1))\sqrt{\ln N \ln \ln N}}$, where $c=1$ for the basic algorithm.

The time required for factorization using this algorithm is shown in Fig.5.

The results of experiments with a quadratic lattice (QS) demonstrate high stability of execution

time. Most values are concentrated in the range of $5.0 \times 10^{-5} - 7.0 \times 10^{-5}$ seconds, which indicates relatively predictable computational behavior of the algorithm. Only a few experiments showed deviations, in particular peak values of about 1.0×10^{-4} seconds, but such cases are rare. Unlike stochastic methods, such as ECM, the quadratic sieve manifests itself as a more deterministic algorithm with uniform load. This is due to its structured construction based on the factor base and the method of linear algebra. It is thanks to this that QS is considered one of the most effective universal methods for factoring medium-sized numbers.

Thus, the experimental results confirm the theoretical estimates – the quadratic sieve ensures consistency and stability of execution time when calculating the factors of numbers in the studied range. This makes it practically suitable for factorization attacks on small and medium-sized RSA modules.

The last algorithm studied is the Dixon algorithm. It is a probabilistic method for factoring integers. The algorithm belongs to a class of algorithms based on finding congruences of squares and is the theoretical basis for more practical and faster algorithms, such as the Quadratic Sieve (QS).

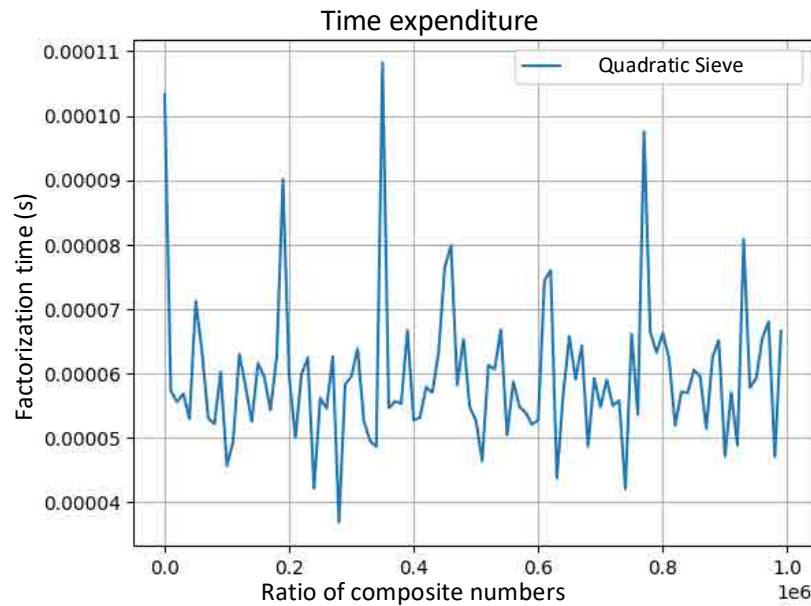


Fig. 5. Results of testing the Quadratic Sieve method

The basic principle of Dixon's algorithm is to find two numbers x and y such that $x \not\equiv \pm y \pmod{N}$ and $x^2 \equiv y^2 \pmod{N}$. If such a pair is found, then $x^2 - y^2 \equiv 0 \pmod{N}$, which means is divisible by N [17, 18].

The method for finding such a pair is based on the system of smooth numbers.

The time required for factorization using this algorithm is shown in Fig.6.

For small composite numbers, the algorithm takes about 3.5×10^{-5} seconds to run. These values are quite

close to QS, which shows the effectiveness of the method for small numbers.

With an increase in composite numbers, there is a gradual increase in the duration of calculations, in some cases up to 0.0015–0.0023 s. This is because Dixon requires more iterations to find dependencies in the factor base as the size of the number increases. There is a strong variability of values between neighboring numbers (values can jump from 0.0006 to 0.0017 s). This indicates the stochasticity of the algorithm, i.e., the efficiency depends on the random selection of squares and the solution of linear systems.

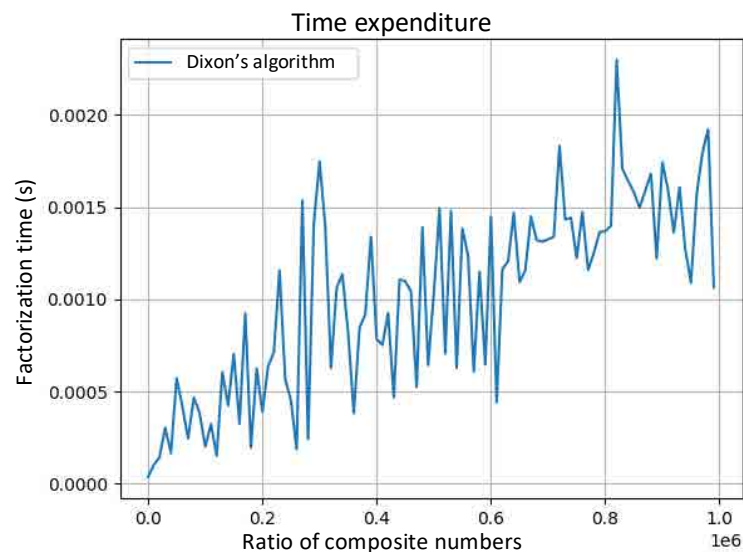


Fig. 6. Results of testing the Dixon algorithm

The upper peaks during the experiment are isolated, but significantly exceed the average values, which confirms low stability compared to QS. The Dixon algorithm demonstrates average efficiency for the numbers studied. It is faster than Pollard's methods $(p, p-1)$ on large numbers, but inferior to the quadratic sieve in terms of stability and predictability of running time. This confirms its status as a transitional method between simpler factorization algorithms and more advanced ones (QS, NFS).

Conclusions

The experimental study allows us to draw the following conclusions regarding the effectiveness of factorization methods.

Pollard's algorithms showed some variability in their performance, which makes them suitable for practical application in conditions where execution time fluctuations are acceptable, but requires caution when estimating guaranteed computational costs. As a result, it is advisable to use these methods in combination with other factorization algorithms, which allows compensating for their uneven performance.

The elliptic curve method (ECM) showed higher efficiency compared to Pollard's algorithms when searching for divisors of average size. The results of the experiments confirm the theoretical position on the advisability of its use in the initial stages of factorization of large numbers.

The Dixon algorithm demonstrated stochastic behavior and significant fluctuations in execution time, which corresponds to its probabilistic nature. The results confirmed its limited practical value compared to more modern methods.

The quadratic sieve (QS) showed the best results among all implemented algorithms, ensuring stability and relatively predictable execution time. This confirmed the theoretical complexity estimates and proved the feasibility of using this method for factoring medium-sized numbers.

Thus, the experiments fully confirm the theoretical expectations regarding the performance of the methods studied. The results obtained indicate that simple algorithms (Pollard, ECM) are appropriate for preprocessing and detecting weak keys, while the quadratic sieve is the optimal choice for factoring medium-sized numbers. For large RSA modules, it is practical to use more complex algorithms, such as the general number field sieve (GNFS).

It should be noted that when solving the problem of compromising the RSA system, factorization algorithms have some differences compared to the mathematical problem of factorization. This difference is due to the number of factors used to construct the number to be factored. Based on the research conducted, the following directions for the further development of factorization algorithms can be formulated. The first is the parallelization of the factorization process, and the second is the screening of impossible solutions.

Conflict of interest

The author declares that there is no conflict of interest, particularly of a financial, personal, authorial, or any other nature, that could influence the research or the results published in this article.

Funding

The research was conducted without financial support.

Data availability

The manuscript has no associated data.

Use of artificial intelligence

The authors confirm that they did not use artificial intelligence technologies to write this paper.

References

1. Mahato, P. and Shah, A. (2023), "A review of prime numbers, squaring prime pattern, different types of primes and prime factorization analysis", *International Journal for Research in Applied Science and Engineering Technology*, 11, pp. 2036–2043. DOI: <https://doi.org/10.22214/ijraset.2023.54904>
2. Klesov, O.I. (2016), *Elementary Number Theory and Elements of Cryptography*, TViMS, Kyiv, 412 p., available at: <https://ela.kpi.ua/handle/123456789/30046> (accessed 11 September 2025).

3. Pieprzyk, J. (2019), "Integer Factorization – Cryptology Meets Number Theory", *Scientific Journal of Gdynia Maritime University*, 1(109), pp. 7–20. DOI: <https://doi.org/10.26408/109.01>
4. Pevnev, V., Yudin, O., Sedlaček, P. and Kuchuk, N. (2024), "Method of testing large numbers for primality", *Advanced Information Systems*, 8(2), pp. 99–106. DOI: <https://doi.org/10.20998/2522-9052.2024.2.11>
5. Fedorchenko, V., Yeroshenko, O., Shmatko, O., Kolomiitsev, O. and Omarov, M. (2024), "Methods of information systems protection", *Advanced Information Systems*, 8(4), pp. 82–92. DOI: <https://doi.org/10.20998/2522-9052.2024.4.11>
6. Kudinov, M. and Muntean, P. (2025), "Modern Number Factorization Algorithms: Efficiency Analysis and Applications", *Collection of Abstracts of Scientific Reports by Higher Education Applicants of Berdiansk State Pedagogical University*, 208 p. DOI: <https://doi.org/10.5281/zenodo.15521215>
7. Prots'ko, I.O. and Gryshuk, O.V. (2019), "Computation factorization of number at chip multithreading mode", *Radio Electronics, Computer Science, Control*, 3, pp. 117–122. DOI: <https://doi.org/10.15588/1607-3274-2019-3-13>
8. Montgomery, P.L. (1994), "A Survey of Modern Integer Factorization Algorithms", available at: <https://ir.cwi.nl/pub/18252/18252B.pdf> (accessed 11 September 2025).
9. Detto, S. (2025), "The New Fermat-Type Factorization Algorithm", *arXiv preprint*, arXiv:2503.07151, pp. 1–33. DOI: <https://doi.org/10.48550/arXiv.2503.07151>
10. Kozukalov, M. and Boiko, M. (2020), "Research and analysis of number factorization algorithms", *ΛΟΓΟΣ. The Art of Scientific Thought*, pp. 64–68. DOI: <https://doi.org/10.36074/2617-7064.10.012>
11. Boudot, F., Gaudry, P., Guillevis, A., Heninger, N., Thomé, E. and Zimmermann, P. (2022), "The state of the art in integer factoring and breaking public-key cryptography", *IEEE Security & Privacy*, 20(2), pp. 80–86. DOI: <https://doi.org/10.1109/MSEC.2022.3141918>
12. Yareschenko, V. and Kosenko, V. (2024), "Low-energy coding method in data transmission systems", *Innovative Technologies and Scientific Solutions for Industries*, 3(29), pp. 121–129. DOI: <https://doi.org/10.30837/2522-9818.2024.3.121>
13. Rabah, K. (2006), "Review of methods for integer factorization applied to cryptography", *Journal of Applied Sciences*, 6(2), pp. 458–481. DOI: <https://doi.org/10.3923/jas.2006.458.481>
14. Lteif, G. (n.d.), "Integer Factorization Algorithms: A Comparative Analysis", available at: <https://softwaredominos.com/home/science-technology-and-other-fascinating-topics/integer-factorization-algorithms-a-comparative-analysis/> (accessed 11 September 2025).
15. Barnes, C. (n.d.), "Integer Factorization Algorithms", available at: <https://connellybarnes.com/documents/factoring.pdf> (accessed 11 September 2025).
16. Wang, B., Hu, F., Yao, H. and Wang, C. (2020), "Prime factorization algorithm based on parameter optimization of Ising model", *Scientific Reports*, 10(1), 7106. DOI: <https://doi.org/10.1038/s41598-020-62802-5>
17. Somsuk, K. (2020), "The new integer factorization algorithm based on Fermat's Factorization Algorithm and Euler's theorem", *International Journal of Electrical and Computer Engineering (IJECE)*, 10(2), pp. 1469–1476. DOI: <https://doi.org/10.11591/ijece.v10i2.pp1469-1476>
18. Kendre, S. (n.d.), "Integer Factorization Algorithms", available at: <https://sauravkendre.medium.com/integer-factorization-algorithms-8f3937502bcc> (accessed 11 September 2025).

Received (Надійшло) 01.11.2025

Accepted for publication (Прийнята до друку) 21.05.2026

Publication date (Дата публікації) 27.06.2026

Відомості про авторів / About the Authors

Певнев Володимир Яковлевич – доктор технічних наук, доцент, Національний аерокосмічний університет "Харківський авіаційний інститут", професор кафедри комп'ютерних систем, мереж і кібербезпеки, Харків, Україна;

Vladimir Pevnev – Doctor of Technical Sciences, Associate Professor, National Aerospace University "Kharkiv Aviation Institute", Professor of the Computer Systems, Networks and Cyber Security Department, Kharkiv, Ukraine;

e-mail: v.pevnev@csn.khai.edu

ORCID ID: <https://orcid.org/0000-0002-3949-3514>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=57194525720>

Юдін Олександр Вікторович – Національний аерокосмічний університет "Харківський авіаційний інститут", здобувач вищої освіти ступеня доктора філософії кафедри комп'ютерних систем, мереж і кібербезпеки, Харків, Україна;

Oles Yudin – National Aerospace University "Kharkiv Aviation Institute", Postgraduate Student of the Computer Systems, Networks and Cyber Security Department, Kharkiv, Ukraine;

e-mail: o.yudin@csn.khai.edu

ORCID ID: <https://orcid.org/0009-0006-8079-0488>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=58882520600>

КЛАСИФІКАЦІЯ ТА ЕКСПЕРИМЕНТАЛЬНІ ДОСЛІДЖЕННЯ ЕФЕКТИВНОСТІ АЛГОРИТМІВ ФАКТОРИЗАЦІЇ

Предметом дослідження є алгоритми факторизації, а саме експериментальна перевірка ефективності сучасних алгоритмів цілочисельної факторизації, виявлення закономірностей між алгоритмами факторизації та розміром чисел, що факторизуються, з погляду часу, необхідного для виконання факторизації. **Мета роботи** – аналіз продуктивності алгоритмів факторизації з використанням різних складених чисел, зокрема чисел середнього розміру (~20 розрядів), що дає змогу оцінити їх ефективність і час виконання. У процесі дослідження необхідно виконати такі **завдання**: систематизувати сучасні алгоритми факторизації, експериментально перевірити ефективність сучасних алгоритмів факторизації Полларда (p та $p-1$), методу еліптичної кривої, алгоритму Діксона й квадратичного решета. Для досягнення мети використовувалися загальнонаукові **методи**: аналіз предметної галузі та математичного апарату, теорія множин, числа й поля, планування та експериментальне дослідження. **Досягнуті результати**. Експериментальні дослідження продемонстрували, що алгоритми Полларда ефективні для чисел з малими дільниками, але втрачають продуктивність зі збільшенням розміру чисел. Метод еліптичної кривої довів свою доцільність у пошуку дільників середнього розміру та кращу масштабованість порівняно з класичними стохастичними методами. Алгоритм Діксона, незважаючи на відносну простоту реалізації, продемонстрував стохастичні коливання часу виконання, що обмежує його практичну цінність у сценаріях зі строгими часовими обмеженнями. Найбільш стабільних і передбачуваних результатів було досягнуто для квадратичного решета, яке довело його придатність для факторизації чисел середнього розміру й забезпечило найменший розкид значень часу за умов багаторазового виконання на ідентичних наборах даних. **Висновки**. Експерименти повністю підтверджують теоретичні очікування щодо продуктивності досліджуваних методів. Результати свідчать про те, що прості алгоритми (метод Полларда й метод еліптичної кривої) ефективні для попереднього оброблення та виявлення слабких ключів, тоді як квадратичне решето є оптимальним вибором для факторизації чисел середнього розміру. Для великих модулів RSA практично використовувати більш складні алгоритми, наприклад загальне решето числового поля. Подальший розвиток алгоритмів факторизації передбачає паралелізацію процесу факторизації та розроблення алгоритмів, що застосовують скринінг неможливих розв'язків.

Ключові слова: цілочисельна факторизація; криптографія; метод Полларда; метод еліптичної кривої; квадратичне решето; алгоритм Діксона; ефективність алгоритму.

Бібліографічні описи / Bibliographic descriptions

Певнев В. Я., Юдін О. В. Класифікація та експериментальні дослідження ефективності алгоритмів факторизації. *Сучасний стан наукових досліджень та технологій в промисловості*. 2026. № 2 (36). С. 109–120. DOI: <https://doi.org/10.30837/2522-9818.2026.2.109>

Pevnev, V., Yudin, O. (2026), "Classification and experimental studies of the factorization algorithms efficiency", *Innovative Technologies and Scientific Solutions for Industries*, No. 2 (36), P. 109–120. DOI: <https://doi.org/10.30837/2522-9818.2026.2.109>

Natalia Pyrih, Serhii Udovenko, Olena Grinyova, Larysa Chala

AUTOMATED SYSTEM FOR DAMAGE ANALYSIS OF URBAN INFRASTRUCTURE OBJECTS USING COMPUTER VISION MODELS

The current problem is to ensure the integration of multi-level data for effective monitoring and analysis of damage to buildings resulting from military operations in Ukraine, and to determine restoration priorities. The purpose of this work is to develop an intelligent system for automated analysis of damage to urban infrastructure objects, enabling effective processing of large volumes of images obtained using computer vision tools for further assessment of the level of destruction and decision-making regarding the priority of their restoration. The detailed reports generated by the system on each building's condition significantly simplify the expert evaluation process, eliminating the need for specialists to be physically present on-site. Ranking damaged objects by their impact on the country's vital activities ensures a more rational, well-balanced allocation of the resources required for infrastructure restoration. In addition, the system actively interacts with relevant state structures, local governments and specialists responsible for planning and organizing restoration work. Within the framework of such interaction, it is envisaged to obtain access to comprehensive databases with information on the architectural features of buildings, their historical and cultural value, as well as the use of existing technical documentation, which makes it possible to take into account all aspects when assessing the severity of damage and forming building restoration plans. Analysis of data obtained from specialised engineering surveys and studies using high-precision instrumental methods enables additional refinement of damage parameters and the creation of extended datasets for further training and validation of computer vision models. The paper investigates and justifies the feasibility of using instance segmentation with YOLOv11 and Mask R-CNN models, enabling not only damage detection but also accurate boundary and quantity determination, which is important for a comprehensive assessment of the scale of destruction and planning restoration measures. The results show that the proposed system works effectively when approaches to combat class imbalance are used. An important stage of the system is the automated formation of optimized restoration plans based on a comprehensive analysis of the identified damage. Such plans should take into account not only the technical parameters of the destruction, but also financial, material and personnel limitations. The proposed system is an effective tool for supporting decision-making at the state and regional levels on the restoration of damaged structures, contributing to the acceleration of the restoration processes of critical infrastructure and minimizing socio-economic losses associated with military destruction.

Keywords: damage analysis; urban infrastructure; deep learning; computer vision models; instance segmentation; image set; intelligent system.

1. Introduction

As a result of the armed aggression launched by the armed forces of the Russian Federation in 2022, the infrastructure of many settlements in Ukraine has suffered significant damage. This has disrupted the normal functioning of residential infrastructure, reduced the quality of life for residents, and led to the loss of cultural and historical heritage. There is an urgent need for the effective restoration of damaged structures, and the issue of prioritizing such restoration becomes particularly important given the limited availability of necessary resources.

To determine the order of restoration, it is necessary to consider the extent of physical damage to the facilities, structural features, social role, and cultural and historical value of the destroyed buildings. Traditional methods of damage assessment, based on expert inspections, are

labor-intensive and typically subjective. Given the significant number of damaged structures across the country, it is advisable to improve approaches to analyzing the extent of infrastructure damage and planning the corresponding restoration work. In doing so, the task of restoring destroyed and damaged structures must be considered both during the war and in the post-conflict period. It should be noted that, given the ongoing hostilities, there is a threat to workers performing restoration work, as well as the risk of new destruction.

Currently, Ukraine lacks information systems capable of automating the processes of obtaining the information on the extent of destruction necessary for making operational decisions and of proposing a comprehensive approach to prioritizing the restoration of damaged facilities, taking into account their physical condition, social significance, and historical and cultural value. Most existing solutions are limited to assessing only

external signs of damage and are unable to provide recommendations on the order of restoration based on a multifactorial analysis.

In light of this, it is particularly important to create an intelligent system that, based on architectural and technical documentation data, building metadata, and using computer vision methods, is capable of determining priorities for the restoration of urban infrastructure facilities. The development and implementation of such a system will not only improve the efficiency of existing resources but also ensure a socially equitable restoration strategy that takes into account community needs, the importance of the facility's functions, and the preservation of national memory.

This work explores the subject area related to the analysis of damage to urban infrastructure and identifies the main problems and challenges that arise during damage assessment and restoration planning. Particular attention is paid to the innovative aspects of the system under development, which is based on the use of computer vision models and deep machine learning for the automated detection of building damage using image data.

Deep learning can, in particular, be effectively used to solve complex tasks: classification, segmentation, and identification of damaged fragments of destroyed building infrastructure. Modern deep learning models are based primarily on the use of convolutional neural networks (CNNs), which allow for improved quality of damage analysis based on the results of computer vision methods and tools. Detecting building damage during military operations and assessing the extent of such damage using intelligent models can ensure effective analysis of the condition of damaged structures and enable prioritized planning for the restoration of these buildings.

The principle of traditional computer vision methods lies in extracting object features from images for subsequent classification. Convolutional neural networks complement the capabilities of computer vision methods through deep learning principles, which ensure the flexibility of neural network models that can be retrained for any possible dataset in tasks involving the identification of defects in damaged buildings [1].

The objective of this work is to develop an intelligent system for the automated detection of damage to buildings based on images, the assessment of the extent of damage, and the prioritization of their restoration, taking into account their cultural and social significance.

2. Literature Review and Problem Statement

The development of systems for the automated detection of building damage in the context of military conflicts, as well as the prioritization of their restoration based on the extent of damage and the cultural and historical significance of the sites, is accompanied by a number of significant challenges. One of the key limitations is the shortage of high-quality annotated data required for effective training of machine learning models. Currently, there is a limited number of publicly available datasets illustrating building damage, particularly as a result of natural disasters. Among these, the xBD dataset [1] is worth highlighting; it contains over 850,000 annotations of buildings on satellite images acquired before and after events (such as hurricanes or earthquakes), with a detailed classification of damage levels. However, the lack of specialized datasets focused on war damage significantly complicates the adaptation of existing methods to the specifics of combat operations and their consequences.

Fig. 1 shows examples of image pairs from the aforementioned dataset.

Models trained on such datasets require paired satellite images (before and after destruction) to effectively analyze damaged objects. This results in significant costs for collecting such data, as well as a reliance on external imaging sources. Moreover, the application of these models may be limited due to the lack of pre-damage historical imagery, which is particularly relevant in the context of military conflicts. The accuracy of models trained on xBD-type datasets is often insufficient when analyzing buildings damaged by combat operations. This is because satellite images do not always capture minor defects (e.g., cracks, localized potholes), and the nature of damage from shelling differs significantly from the effects of natural disasters, which are predominantly represented in existing datasets. The latter focus on global landscape changes, whereas in wartime conditions such changes may be absent or may not correlate with the actual condition of buildings [2]. The overall accuracy of damage assessment depends largely on the quality of the input data. The complexity of analyzing satellite imagery lies in the presence of obstacles, such as neighboring buildings, the surrounding environment, and clouds, which partially or completely obscure the image [3]. Additional factors – shadows, vegetation, adverse weather conditions, smoke, and suboptimal shooting

angles – also distort visual information, making it difficult to identify key features of objects. Furthermore, the inability to safely collect additional data on-site

(due to the threat to experts' lives) limits the ability to expand training datasets to the sizes required for effective model training.



Fig. 1. Landscape samples before (top row) and after (bottom row) the disaster from the xBD dataset

One of the key challenges in developing building damage analysis systems is the need to account for the variety of building materials, such as concrete, brick, or precast concrete, which significantly affect the stability of structures and the feasibility of their restoration [4]. For example, concrete structures are capable of withstanding significant dynamic loads, whereas brick walls demonstrate greater vulnerability to blast waves. Furthermore, the nature of damage in brick and concrete buildings can differ significantly: cracks in brick masonry are often less critical for operational safety than similar defects in concrete structures.

This heterogeneity complicates the creation of unified damage assessment models, as it requires a differentiated approach to analyzing each type of material. An important aspect is the assessment of the condition of load-bearing floors, since damage to them directly affects the overall stability and safety of the building. To detect hidden defects and accurately determine the extent of damage, specialized equipment – such as ultrasonic scanners, laser rangefinders, or ground-penetrating radar – is often required [5]. However, their use is limited due to high costs and the complexity of deployment in wartime conditions. Economic and organizational constraints also play a significant role. Countries affected by armed conflicts often face a shortage of the financial and human resources needed for large-scale infrastructure restoration. In Ukraine, in particular, investment in reconstruction is limited due to the prolonged war [6], which complicates the

implementation of high-tech damage assessment methods. Another challenge is the lack of systematic criteria for assessing the cultural and historical value of sites. Uncertainty regarding restoration priorities, stemming from a lack of standardized methodologies and objective indicators, significantly complicates the decision-making process regarding the order of building reconstruction.

Below, we examine existing platforms and software solutions currently used for the automated detection and classification of damage to civil infrastructure objects.

At the start of the full-scale war in 2022, the UN funded the "Infrastructure Semantic Damage Detector" (ISDD) project for Ukraine [7], under which a model was created that uses machine learning algorithms and big data processing methods to analyze infrastructure damage by processing reports. The model uses the open-source ACLED database, which contains global real-world data on the consequences of destruction, as its data source. The system based on this model was developed to determine the location of infrastructure damage and the scale of the resulting losses to enable rapid response. It was also intended to use this system to generate ready-made reports for the purpose of submitting requests for assistance to international humanitarian organizations and calculating the amount of aid required. By analyzing the text of the reports, data on the date, time, location, cause, and type of damage are recorded. One of the challenges addressed by the system's developers is determining the type of damaged infrastructure in the required format. Nine categories were proposed for

classification: industrial facilities, logistics, energy/power supply, telecommunications, agriculture, healthcare, education, housing, and business. By identifying keywords for each of these categories and using the cosine similarity measure to establish a semantic link between the text of the reports and the infrastructure categories, the system can automatically classify the types of damaged infrastructure.

The final stage of the system's operation is the visualization of the results on an interactive map of Ukraine. The ability to filter the buildings identified and displayed on the map by region and settlement, as well as by their function, allows users to examine the geographic distribution of damage across various types of infrastructure. The ability to generate bar and pie charts allows users to examine the results in greater detail. Figure 2 shows an example of an interactive dashboard built on the ISDD platform. The advantages of the described system include: speed and efficiency; scalability; and the ability to analyze various types of infrastructure.

The disadvantages of this system include: limitations caused by processing only text data without analyzing visual materials; dependence on the quality of text data without the ability to verify the accuracy of the obtained data against data from other sources; the system's focus solely on determining the type of damaged infrastructure and its location, while ignoring the severity of the damage and the assessment of losses.

In 2022, the Kyiv School of Economics, in collaboration with Ukrainian government agencies,

initiated the "Russia Will Pay" project, aimed at a comprehensive analysis and assessment of the damage inflicted on the country's infrastructure as a result of military aggression [8]. The project involves collecting and systematizing data from various sources, including official reports from government agencies, business associations, verified media materials, as well as satellite imagery and the results of drone surveys of liberated territories. This approach allows for the assessment of damage not only to buildings but also to critical infrastructure, major roadways, and agricultural land. A key partner in the project is the Polish-American company Tensorflight, which provides satellite imagery of sites free of charge and uses machine learning and artificial intelligence tools to analyze them. The innovation of the method lies in the use of imagery captured at different angles (0° and $30\text{--}50^\circ$ relative to the axis), which makes it possible to assess damage to both roofs and building facades – even if the roof remains intact. After preprocessing the images, they are segmented into undamaged and damaged objects, and the degree of damage is classified – from partial (repairable) to total. Based on historical data on damaged buildings, Tensorflight also provides a preliminary estimate of reconstruction costs. To date, the system has been used to analyze damage in cities such as Bucha, Irpin, and Mariupol. Additionally, the project helps prioritize sites for restoration by providing analytical support for decision-making.

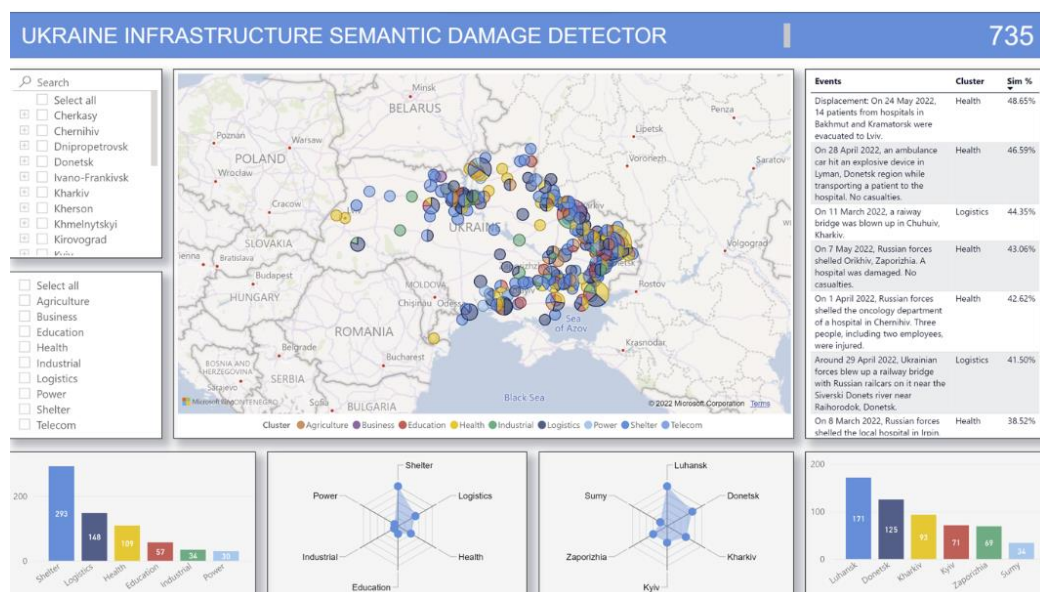


Fig. 2. Example of an interactive board built using the "Infrastructure Semantic Damage Detector" platform

The system offers the following advantages: a multimodal approach to gathering information about damage; detection of damage to building facades with intact roofs using the proposed satellite imaging approach. The disadvantages of the system under consideration include: the dependence of users (government agencies) on a commercial company from another country; the possibility of confidential information about infrastructure facilities being leaked; and the high cost of collecting and processing satellite data.

One of the leading commercial companies specializing in assessing damage to buildings caused by external factors is FlyPix AI [9]. The company has developed a geospatial analytics platform (Geospatial AI platform) that integrates pre-trained machine learning and deep learning models, as well as artificial intelligence tools for processing drone imagery, aerial photography, and high-quality satellite images. The platform provides automated detection, segmentation, and classification of building damage in real time. Initially, FlyPix AI's models were focused on analyzing damage caused by natural disasters and were successfully used to assess damage to settlements following earthquakes in Turkey and Japan. In 2025, the company adapted its solutions to analyze building damage resulting from military operations in Ukraine. The official website notes that previously developed models, trained on natural disaster data, did not provide sufficient accuracy for military damage. Therefore, new models were created, trained on specialized custom datasets containing information about destruction in combat conditions. To improve accuracy, multimodal data fusion was also implemented.

The Geospatial AI Platform features machine learning and deep learning models with various architectures for the segmentation and classification of images of damaged buildings (specifically, U-Net and Mask R-CNN deep learning models combined with trained ResNet34, SeResNext50, Inception v3, and EfficientNet B4 models). The highest accuracy was achieved by using an ensemble of Mask R-CNN models (for detecting damaged structures) and Inception v3 models (for assessing the degree of damage). After the selected deep learning models have completed their tasks, the system allows users to work with dashboards and interactive maps to examine the results in greater detail.

The advantages of the described system include: a wide selection of models with different architectural types, with the ability to fine-tune them on user-provided datasets to solve a wide range of diverse tasks; convenient tools for automated image annotation in the user dataset; multifunctionality; no requirement for users to possess knowledge or skills in machine learning and artificial intelligence to use the system; integration of the configured model into cloud environments; real-time analysis of satellite and drone imagery. The system's drawbacks include: limited capacity for building monitoring and free usage based on the number of requests; the inability for users to view detailed information about detected damage; lack of information regarding the type of structure being inspected; and the inability to prioritize building restoration based on their cultural and social value.

Another platform for analyzing building damage is the "AI Damage Detector" platform from T2D2 [10]. This cloud-based platform allows for the detection and classification of damage to buildings using computer vision algorithms (Fig. 3). It should be noted that T2D2 did not aim to assess damage to structures caused by military actions, but rather focused on detecting various types of damage on facades during building condition monitoring, which may partially relate to the topic of this study.

Platform users can use images and videos captured from mobile devices, cameras, or drones as input data for the models they build. The system built on this platform allows users to track the deterioration of buildings over time by comparing images and videos taken at different intervals. The "Inspection Cloud" cloud environment developed by T2D2 allows users to store necessary materials and integrates with Google Drive, OneDrive, and Box. Advantages of this system: the ability to generate detailed reports describing various types of damage found; high classification accuracy due to large proprietary training datasets used to train the proposed models. The system's drawbacks include: the lack of a free plan (with the exception of a two-week trial period); a lack of focus on identifying specific damage sustained by buildings during military operations; lack of assessment of the socio-cultural significance of buildings; lack of assistance in making decisions regarding the restoration of structures as well as determining the order of building restoration.

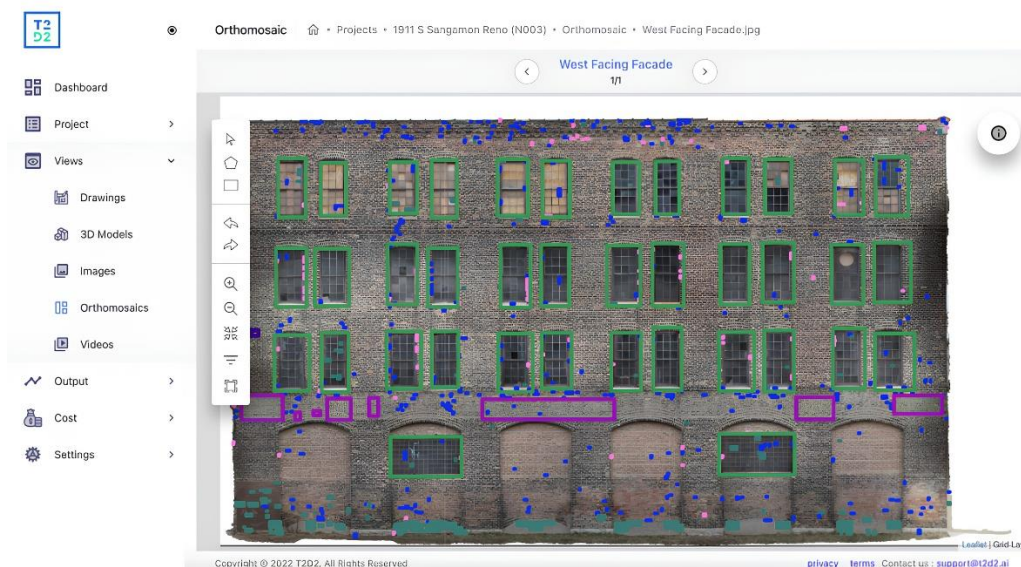


Fig. 3. Interface of the "AI Damage Detector" platform

At the start of the war, various volunteer organizations launched several projects aimed at documenting building damage by creating 3D models. These projects cover both simple structures and those of cultural or historical value. Drones and aerial photography are typically used to survey structures, and 3D models are created using photogrammetry. The results obtained allow not only for documenting damage but also for

specialists to use them for further analysis and planning of restoration work. Among the most well-known projects in this field is the "War up Close" project, which uses 360° panoramic photographs, drone footage, and 3D modeling to document the criminal destruction of Ukraine's infrastructure for subsequent presentation to the global community [11]. Fig. 4 shows one of the pages of the War in 3D project featuring constructed 3D models.

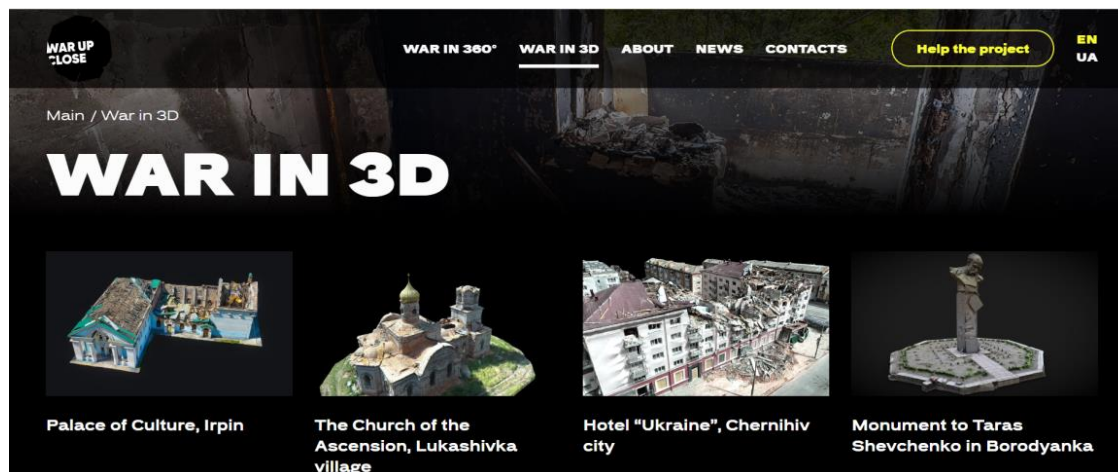


Fig. 4. Example of a page from the "War in 3D" section of the "War up Close" project

The project team is actively collaborating with various government agencies and helping rescue workers clear debris and assess the true extent of the damage.

The global initiative for Ukraine's recovery, UNITED24, has launched a joint project with the IT company LUN, using an app with patented technology that builds 3D models of structures based on drone

imagery [12]. This initiative allows for tracking the progress of the reconstruction of damaged buildings and raises funds for infrastructure restoration. Currently, the program is limited to Kyiv and covers 17 sites. Figure 5 shows an interactive map of Kyiv with the locations of damaged buildings.

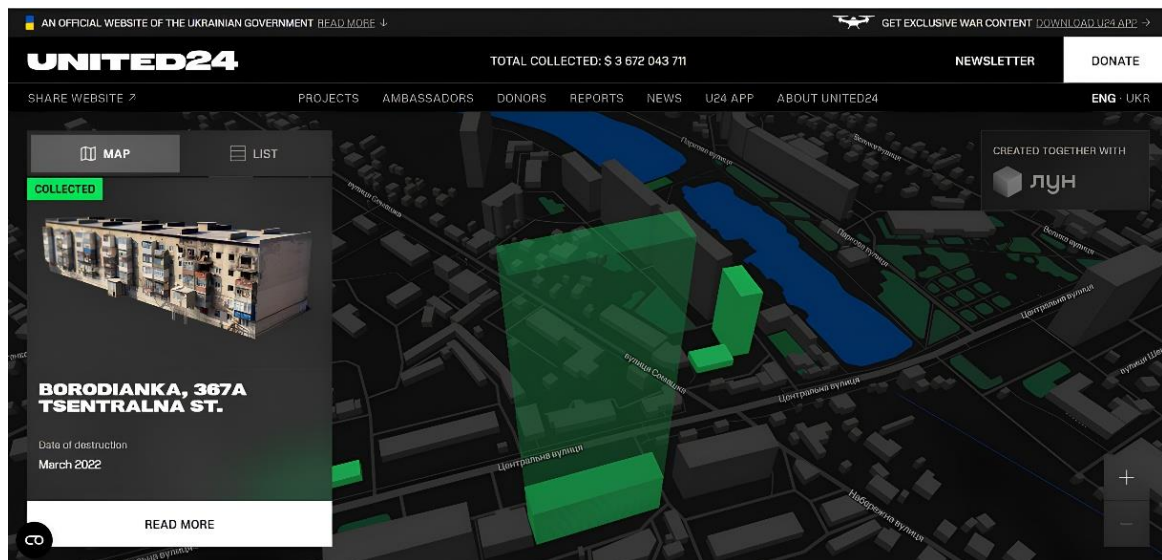


Fig. 5. Example of an interactive map showing the locations of damaged buildings, provided by UNITED24

Unlike previous initiatives, which have focused more on ordinary residential buildings, the volunteer initiative Scan UA works with Ukrainian heritage sites that have been damaged or are at risk of disappearing [13]. The goal of this initiative is to create a digital archive of Ukrainian heritage to preserve the memory of cultural heritage, as well as to provide information about the damage inflicted on the country's architecture. Figure 6 shows a 3D model of the damaged Fire Station No. 4 in Kharkiv.

The advantages of all the initiatives described are that the 3D models produced as a result of their work allow specialists to examine the damage present on the buildings in greater detail and plan restoration work more thoroughly. At the same time, a significant drawback is the lack of an automated mechanism for detecting damage and determining the level of destruction of buildings. These initiatives do not propose an objective mechanism for selecting buildings for structural damage assessment or for determining the order of restoration work.

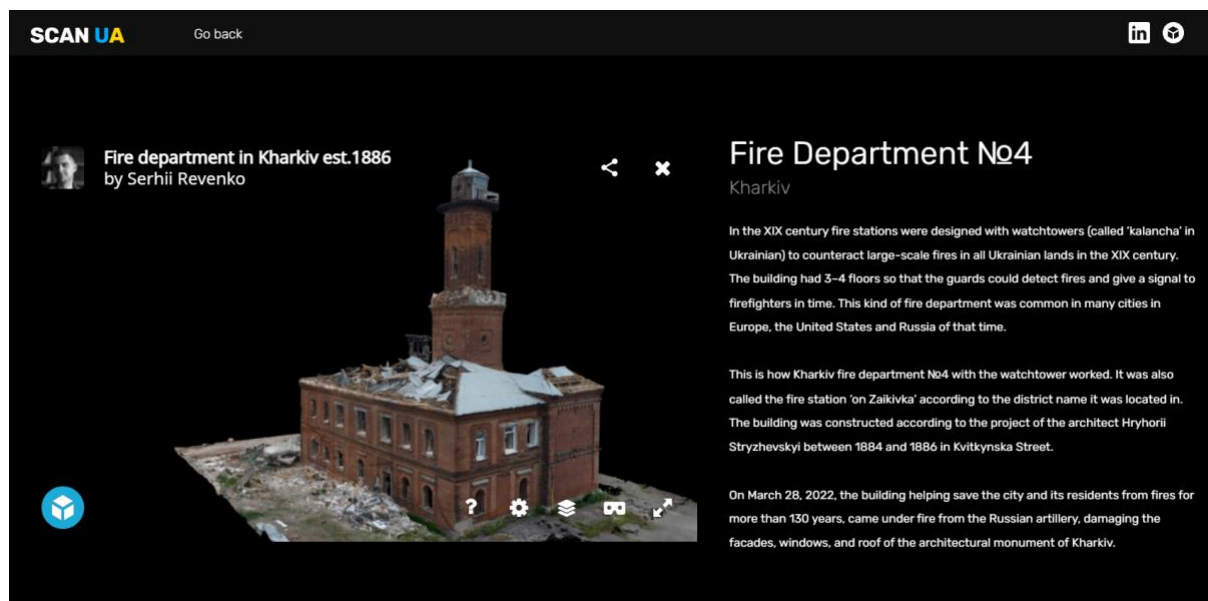


Fig. 6. Interface page showing a 3D model of the damaged Fire Station No. 4 in Kharkiv

Below we present the results of scientific studies most relevant to the topic of this paper.

In [14], the authors proposed an approach to the automated detection of damaged buildings in satellite and

aerial imagery, based on unsupervised machine learning methods. Images of Ukrainian cities from Google Earth, obtained before and after the start of the war, were used as input data. Before applying machine learning algorithms, the images underwent preprocessing, which included georeferencing, orthorectification, contrast enhancement, and brightness equalization. The algorithm consisted of three main stages. In the first stage, Principal Component Analysis (PCA) was applied to determine the principal components describing the directions of maximum data variance. In the second stage, the K-means algorithm was used to partition the image into clusters by grouping similar areas of change. In the final stage, the DBSCAN (Density-Based Spatial Clustering of Applications with Noise) clustering method was applied to identify dense regions corresponding to damage. The result of the algorithm's operation was a binary damage mask, where damaged areas were highlighted for further analysis. The advantages of the proposed method lie in its high damage detection accuracy (Accuracy = 98.13%, F1 Score = 70.90%), as well as its ability to work with data of varying quality and spatial resolution without the need for pre-training complex deep learning models. However, the algorithm has certain limitations. First, it does not accurately define the boundaries of different types of damage. Second, the presence of visual artifacts, such as shadows, glare, or uneven lighting, can significantly reduce the accuracy of the results. Furthermore, the effective application of the method requires the availability of satellite or aerial images taken before and after the damage, which are not always obtainable. Additionally, not all types of damage are clearly visible from above, which limits its scope of application.

In [15], an innovative approach to determining the degree of building damage was proposed, based on the construction of a Damage Detection Map (DDM) model. Photographs of damaged buildings obtained after the earthquakes in Nepal and Ecuador, Typhoon Rui, and Hurricane Matthew were used to train the model. A distinctive feature of this approach is that the images were captured at ground level, which allows the model to be used in the future to analyze photographs posted on social media by ordinary users without the need for specialized equipment.

The model's architecture is based on a convolutional neural network (CNN), specifically a modified VGG19 model, whose outer layers were fine-tuned on collected and labeled datasets. To visualize and interpret the results, the Grad-CAM method is applied in the final

convolutional layer, which calculates weights based on gradients for each object map, highlighting the most significant areas. This allows for the generation of a DDM heatmap, where areas of high intensity correspond to detected damage. To classify the severity of damage (minor, moderate, severe), the authors proposed the DAV (Damage Assessment Value) metric, which is calculated as the average of the intensities on the heatmap. The higher the DAV value, the more severe the damage. Using DAV thresholds allows the model to be adapted to different types of damage, providing flexibility in determining their severity. The model's output is a thermal map highlighting the damage, as well as a classification of the building by damage severity. This approach opens up possibilities for automated analysis of photographs taken by non-specialists, significantly expanding the capabilities for documenting and assessing damage in real time.

The main advantages of the proposed approach [15] are its cost-effectiveness and ability to be deployed immediately following the occurrence of emergencies. This is particularly relevant in situations where large volumes of data can be rapidly collected from social media, enabling a prompt response to damage without the need for specialized equipment. Another important advantage is the clarity and comprehensibility of the results, which helps build trust in automated systems for detecting and classifying the extent of damage to buildings.

In a subsequent study [16], the authors expand on this approach, emphasizing the advantages of the Grad-CAM and Grad-CAM++ methods, which provide improved interpretability of results. Additionally, the application of the model for classifying damage to various types of infrastructure – including roads, bridges, and buildings – is demonstrated, broadening its scope of practical use. However, the method has certain limitations. First, small-scale damage may not be visible on heatmaps, creating a risk that it will be overlooked during the planning of restoration work. Second, although the method effectively visualizes large-scale damage, its capabilities are limited when it comes to detecting minor defects. At the same time, the algorithm for automatically assessing the severity of damage allows buildings to be classified by degree of damage, which is an important advantage for prioritizing restoration efforts.

In [17], a method was proposed that allows for the assessment of the degree of building damage on a scale from 0 to 1 based solely on aerial imagery, without the need for before-and-after photographs of the structures.

To train the model, a custom dataset was created, comprising 250 images of buildings damaged by natural disasters, collected from open sources and pre-annotated by experts. The method's architecture involves three parallel feature processing streams, whose results are combined for the final damage severity prediction. The first stream processes unprocessed color images using a VGG-based convolutional neural network to extract deep features. The second stream first applies semantic segmentation using DilatedNet to remove extraneous objects and extract only relevant elements (e.g., building walls), and then processes the resulting color mask in a manner similar to the first stream. The third pipeline converts the segmented image into a binary mask representing the building's contours with damaged areas and analyzes it using a simplified LeNet architecture. The features obtained from all three streams are combined into a single vector, which is processed by a fully connected layer with a sigmoid activation function to determine the building's damage level as a continuous value. The advantages of the proposed approach lie in the absence of a need for data on the condition of objects prior to destruction, as well as in the ability to obtain a continuous assessment of the degree of damage, which allows for the flexible prioritization of the restoration of structures with different types of damage. However, the method also has significant limitations. First, the high computational complexity requires substantial resources for model training. Second, the quality of the results depends on the training dataset, which must be pre-annotated by experts. Discrepancies in assessments by different groups of experts or changes in the criteria for determining the severity of damage may necessitate the preparation of a new dataset and retraining of the model, which complicates its practical application.

In [18], an approach was proposed aimed at reducing the computational complexity of models by utilizing computer vision methods that account for the building materials used during construction. The data source for the experimental studies was a custom Google My Maps map created in April 2022, which contains images of damaged buildings in Mariupol taken from both ground-level and drone perspectives. Image annotation was performed on the Roboflow platform using four damage severity classes adapted from the EMS-98 classification. Various augmentation methods were applied to expand the training dataset and improve model robustness. A pre-trained YOLOv8 model was used to automatically determine the damage level; its

training lasted 50 epochs and yielded the following results: mean average precision (mAP) was 58.4%, precision was 64.1%, and recall was 52.5%. The same architecture was used to determine the type of development project. The main advantage of this approach is the use of open data collected by eyewitnesses to the events. The images used to train the model vary in quality, distance to the object, and shooting angles, which contributes to the model's stability when working with new data. In addition, the use of the modern universal YOLOv8 model allows for a balance between prediction accuracy and productivity. However, the method has certain limitations. First, the classification of damage severity is based on a five-point scale for assessing post-earthquake damage, which may not fully correspond to the nature of damage caused by military actions, potentially affecting the model's accuracy. Second, the subjectivity of assigning structures to different classes can lead to ambiguities in classification results, complicating the interpretation of the obtained data.

In [19], an approach was proposed that addresses these issues using a simplified version of the YOLOv5 model, optimized for real-time operation with minimal computational cost while maintaining high accuracy. To train and test the model, a specialized dataset titled "Ground-level Detection in Building Damage Assessment" was created, comprising ground-level photographs of damaged buildings with detailed annotations of various types of damage. The model was trained to recognize damage types such as collapses, cracks, material delamination, and debris accumulation. The modification of the base YOLOv5 architecture, named LA-YOLOv5, involved the introduction of Ghost Bottleneck elements and the CBAM (Convolutional Block Attention Module), as well as improvements in handling objects of different scales. These changes allowed the model to maintain high accuracy while significantly reducing the memory requirements for its deployment. The main advantage of the proposed approach is its high processing speed, making the model suitable for use in situations that require rapid damage detection and quick response. In addition, the model's ability to recognize different types of damage allows for a more in-depth analysis of the damage, which facilitates informed decision-making regarding the priorities of restoration work.

However, the method has a significant limitation – the lack of a formalized algorithm for assessing the degree of building damage, which complicates the automated determination of the level of damage and

may require additional expert assessment for accurate classification of objects.

The main drawback of the methods discussed above is the lack of an objective approach to assessing the degree of building damage and the ability to adapt this method to various situations requiring the application of the developed models. Recently, a number of studies have appeared dedicated to solving the problem of determining the priority of restoring one structure over another. For example, Article [20] describes the application of the Multi-Criteria Decision Analysis (MCDA) method to determine priority infrastructure objects for construction or reconstruction under limited financial resources, consisting of a combination of the two-phase AHP method and the TOPSIS method. To determine the weights of the selected criteria, the AHP method was used, which in the first phase assigns weights to the criteria taking into account the opinions of all stakeholders, and in the second phase adjusts the determined weights based on the decision of the lead expert (key person). Next, the priority for each object is calculated using the TOPSIS method by ranking the objects according to their proximity to the ideal object. The advantages of this approach include the transparency of the ranking process and the ability to account for different perspectives on the task at hand. Its disadvantages include both the time and computational complexity of configuring the AHP method when considering multiple criteria with the participation of a large number of experts.

In [21], a comprehensive approach was implemented to determine the order of implementation for proposed disaster recovery projects, where a set of comparison criteria was selected for each project. The Best Worst Method (BWM) was used to determine the importance of each criterion, allowing experts to select the most and least important criteria for the projects and, based on their comparisons, calculate weights for the entire evaluation system. The TOPSIS, ELECTRE III, VIKOR, and PROMETHEE methods were applied for ranking. The final result (a ranked list of projects) was obtained by linearly combining the results provided by each of the ranking methods into a single priority vector while minimizing discrepancies. To verify the reliability of the results, an ANN neural network trained on historical data about completed projects was used. The results obtained using the MCDA method coincide with the neural network's predictions. The advantages of the described approach include the hybrid use of four ranking methods,

which allowed for more accurate results, and the BWM method for determining criterion weights, which is less labor-intensive compared to the AHP method but also ensures high consistency of decisions.

Let us now consider the problem of applying computer vision tools, which are to be used in conjunction with deep learning methods to implement the tasks set in this study. In the context of the task of detecting damage to buildings during wartime and assessing the extent of such damage, it is necessary to select a computer vision method capable of effectively detecting and analyzing damage in input images. To solve this problem, approaches such as image classification, object detection, semantic segmentation, and instance segmentation can be applied [22, 23]. One of the fundamental areas of computer vision is image classification, which involves assigning an image to one or more categories based on its visual characteristics [24]. Classification models are trained on a dataset of pre-labeled images, where each image must correspond to a specific class. In the context of this work, these may be images with text labels ("damaged building", "partially damaged", or "undamaged") or numerical values (degrees of building damage). There are a number of pre-trained classification models (e.g., ResNet or VGG) that require minimal fine-tuning to achieve high accuracy and image processing speed [25]. At the same time, these models are typically unable to determine the location or boundaries of damage, which significantly limits their usefulness for tasks requiring the localization of urban infrastructure objects. Furthermore, they require the involvement of experts for data annotation, which can lead to subjectivity in assessments, as different specialists may interpret the extent of building damage differently during visual analysis. This creates a risk of label inconsistency, which negatively impacts the model's reliability. It is advisable to focus on computer vision methods that enable the localization and detailed analysis of damage, allowing the process of assessing the extent of destruction to be separated from the model and carried over to the next stage.

An important area of computer vision is object detection, which focuses on identifying and localizing objects in an image by creating bounding boxes around them, followed by classification. During the detection process, the model learns to recognize objects based on their visual characteristics, such as shape, texture, or color, and to create rectangular bounding boxes that enclose these objects, along with labels indicating their

class [26]. In this process, detection identifies objects as separate units but does not provide information about their exact boundaries. The main advantage of this approach is that it allows for the precise determination of an object's location within an image, which represents significant progress compared to image classification. However, its drawbacks are that the bounding boxes do not reflect the exact shape or boundaries of the objects, which reduces the accuracy of analysis in tasks requiring detailed information about the structure of damage. In cases where objects overlap, have blurred boundaries, or are closely spaced, this approach may lose accuracy, as the bounding boxes may incorrectly enclose objects or miss them [27].

To achieve the highest efficiency in detecting building damage, it is worth considering the field of image segmentation, which allows not only determining the location of objects but also their exact contours. There are two main types of image segmentation: semantic segmentation and instance segmentation [28].

Semantic segmentation involves classifying each pixel in an image, determining its membership in a specific class. A distinctive feature of this approach is that individual instances of the same class are not distinguished but are treated as a single category. Because the method provides a detailed analysis of the image at the pixel level, it becomes possible to accurately estimate the area and distribution of objects in the image. Additionally, semantic segmentation performs well when analyzing images with a large number of objects or complex structures, where segmentation of the entire scene is required. At the same time, the fact that semantic segmentation yields detailed information about objects and their clear boundaries is a drawback of this approach, as it does not allow for distinguishing individual instances of the same class, which limits its use in tasks requiring the identification of individual objects. For the task at hand, using semantic segmentation allows us to estimate the total area of damage, but it cannot provide information about the number or shape of individual instances of a single type of damage.

This limitation can be addressed by instance segmentation, which combines the capabilities of object detection and semantic segmentation, enabling the identification, localization, and pixel-level classification of individual object instances in an image. This method not only determines the category of objects but also creates separate pixel masks for each instance, allowing for the clear separation of different objects within the

same class. The advantages of this approach include the ability to perform a comprehensive, detailed analysis of images containing a large number of objects or complex structures, providing information on all detected elements. A drawback of this approach is the need for significant computational resources.

Therefore, for the task of detecting damage to buildings, the best approach is to use instance segmentation, which allows not only for the detection of damage but also for the precise determination of its boundaries and quantity, which is important for a comprehensive assessment of the scale of destruction and the planning of restoration measures.

3. Purpose and Objectives of the Study

Based on the results of the analysis presented above, we formulate the purpose and objectives of this work.

The objective is to develop an intelligent system for the automated analysis of damage to urban infrastructure objects, which would enable the effective processing of large volumes of images obtained using computer vision techniques, for the subsequent assessment of the level of damage and decision-making regarding the priority of their restoration.

To achieve this goal, the following must be carried out:

- selection of computer vision models for segmenting instances of the analyzed image sets of damaged objects;
- balanced formation of datasets containing images of buildings with various types of damage;
- ranking of damaged urban infrastructure objects to determine the priority of their restoration;
- definition of the structure and functions of the proposed intelligent system for automated damage analysis;
- software implementation of the main modules of the proposed system;
- implementation of the server-side and system interface;
- experimental research and testing of the system.

4. Materials and Methods

4.1. Selection of computer vision models for segmenting instances in the analyzed image sets

Instance segmentation is an effective approach for solving the problem of detecting war-induced damage to buildings, where a single image may contain 10 to 100 small instances of damage with a significant

class imbalance. After determining the field of computer vision, it becomes important to correctly select a model from this field that is capable of accurately identifying all damage and operating at high speed. This section is devoted to a review of modern computer vision models for instance segmentation, with an explanation of their suitability for solving the given task.

Currently, there is no consensus on the best object segmentation model, given the variety of comparison criteria: segmentation accuracy, processing speed, computational efficiency, and adaptability to class imbalance [29]. It is advisable to compare segmentation models based on mAP scores for bounding boxes, masks, and ease of configuration. In terms of speed, models from the YOLO family have a significant advantage; their latest versions have the following mAP values: 54.7 for bounding boxes and 43.8 for masks [30]. Among the models that are slower but provide high object detection accuracy and attention to small details, the modern architectures of BEiT3, Mask2Former, MaskDINO, and OneFormer are worth noting. However, a significant drawback of these models, which complicates their practical use, is their high computational complexity. In addition, these models require large amounts of finely annotated data for fine-tuning.

Other relevant segmentation models include Mask R-CNN, Panoptic FPN, and YOLACT++. Mask R-CNN provides accurate pixel-level segmentation and is suitable for tasks requiring maximum accuracy, such as visualizing sections of damaged buildings. Panoptic FPN combines instance segmentation and semantic segmentation, providing a comprehensive understanding of the scene. YOLACT++ is designed for fast, real-time object segmentation, which is useful for video analysis, action recognition, and object counting. The advantages of each model should be considered based on the challenges encountered in each specific application of computer vision for segmentation.

In particular, for the task of detecting building damage among the models considered, special attention should be paid to the latest models in the YOLO and Mask R-CNN families, as the former are capable of combining high image processing speed with good accuracy, while the latter remain leaders in tasks where maximum pixel segmentation accuracy and detail are critical.

The YOLOv11 and Mask R-CNN models represent two different approaches to solving the object segmentation problem, each with its own architectural features, advantages, and limitations.

The YOLOv11 model implements a modern single-stage algorithm for object segmentation tasks, built on the YOLO family architecture, which was originally developed for fast and efficient object detection in images and later adapted for segmentation tasks [31]. The basic principle of such models is to divide the image into a grid and simultaneously predict bounding boxes, object classes, and segmentation masks for each region within a single pass through the image [31]. The YOLOv11 architecture features an improved backbone based on Darknet-53, which ensures effective feature extraction using convolutional layers. The model employs a modified C2F module, which improves the quality of fine-scale feature extraction and reduces the size of convolutional filters from 6×6 to 3×3 , thereby lowering computational costs. A significant feature is the use of a decoupled head (decoupled head), which separates the tasks of coordinate regression and classification, improving segmentation productivity and accuracy.

To improve the stability and quality of predictions during segmentation, the Pseudo Ensemble method is often used, allowing the results of several models to be combined for a more stable outcome. These architectural solutions make YOLOv11 a very fast and efficient model capable of processing up to 128 frames per second, which is particularly important for real-time applications and systems with limited resources.

The main advantages of YOLOv11 are its high data processing speed and balanced accuracy. However, despite its improved architecture, this model may exhibit reduced accuracy compared to two-stage methods, especially when dealing with complex or small objects.

In contrast to YOLOv11, Mask R-CNN is a two-stage model that has gained widespread recognition due to its high object segmentation accuracy [32]. The Mask R-CNN architecture includes a backbone network (typically ResNet-50 or ResNet-101 with FPN), which extracts multi-level features from the image. Next, the Region Proposal Network (RPN) generates region proposals, which are precisely aligned using the RoI Align layer to extract features of specific regions of interest.

The scaled features are fed into the model's head, which concurrently performs object classification, bounding box regression, and the generation of pixel-level segmentation masks for each object. Model training is based on a combined loss function that accounts for all three of these tasks. A key feature of Mask R-CNN is its high-quality pixel-level segmentation and robustness to challenging image conditions, making it highly sought

after for applications that require precise object detection. The main drawback of Mask R-CNN is its relatively high computational complexity and slow inference time, which limits the model's use under strict time constraints.

Thus, the YOLOv11 and Mask R-CNN models implement different approaches to image segmentation and are effective for use in automated defect analysis systems for damaged urban infrastructure objects. YOLOv11 provides higher processing speed and generally better accuracy compared to Mask R-CNN. At the same time, Mask R-CNN remains the leader in tasks requiring the most accurate and detailed object segmentation, despite its significant computational load. The choice of model should be based on the requirements of the specific task and the available computational resources.

4.2. Addressing the Issue of Class Imbalance in the Training Image Dataset

In the process of developing an automated system for detecting damage to buildings during military operations in Ukraine, a significant problem arises due to the uneven distribution of classes in the training dataset. The class imbalance is caused by the natural heterogeneity of the types of damage observed in the images. In particular, some categories of damage (e.g., broken windows) are significantly more common due to their high frequency of occurrence. For example, even in the absence of a direct hit on a building, the blast wave often causes damage to windows, which significantly increases the proportion of such objects in the dataset. At the same time, buildings with critical damage, such as collapsed structures, usually also have broken windows, which further increases the number of instances in this class. Since there are numerous windows in every building, the number of instances in this class significantly exceeds other damage categories, such as cracks or collapsed walls.

This uneven distribution of data creates conditions for the model to be biased toward the dominant class. When using standard training approaches, the model will demonstrate high overall accuracy, since most annotations in the dataset belong to the dominant class, and their correct classification ensures a high accuracy score. However, when the model is used in practice, its effectiveness will be significantly reduced, as such a model will be unable to detect underrepresented but critically important classes that are crucial for assessing the extent of damage and planning reconstruction.

Thus, in this case, high overall accuracy does not reflect the model's actual ability to solve the given task, which reduces its practical value. Class imbalance issues should be taken into account both for preprocessing the dataset and for tuning computer vision methods. One of the basic and effective methods for addressing class imbalance through preprocessing of image datasets in computer vision tasks is data augmentation, which allows not only for increasing the total size of the training sample but also for balancing the representation of minority classes by varying the structure and appearance of images.

All image augmentation methods can be broadly divided into basic and advanced [33]. The basic approach combines methods of geometric and non-geometric image manipulation, as well as image erasure techniques. Geometric transformations, whose main operations are rotation, translation, scaling, axis translation, and angle translation, change the spatial position of objects, allowing for increased model invariance to the orientation and position of objects. On the one hand, such transformations allow for increased model robustness, given that in reality, damage may be captured from different camera angles or in different positions. On the other hand, transformations that radically alter the spatial arrangement of objects in the frame can lead to the creation of training images that do not occur in reality and may negatively impact the model, resulting not merely in a lack of improvement in accuracy but in a significant deterioration in quality. In general, the selection and use of each geometric transformation method requires a comprehensive preliminary check to ensure that the resulting images correspond to reality. Non-geometric transformation methods include changes in color space, brightness, and contrast, as well as noise injection and filtering, which improve the model's robustness to variations in lighting and textural noise. Unlike geometric transformations, this approach is safer to use and does not have such a strong impact on reducing the model's prediction accuracy. Using this approach to data augmentation makes the model more robust to cases where images were captured under different shooting conditions (using different devices and with varying image quality). However, it is important not to overuse these methods, as this can lead to a significant deterioration in image quality and the model's inability to distinguish between the desired classes during training. Erasing methods (specifically CutOut, Random Erasing, and GridMask) involve the deliberate removal of random or regular sections

of the image, which forces the model to pay more attention to different parts of the object rather than just the most prominent features, thereby facilitating better recognition of rare or inconspicuous objects.

Advanced image augmentation techniques include image mixing methods (such as Mixup, CutMix, SaliencyMix, and PuzzleMix), which create synthetic images by combining regions from two or more images, while proportionally combining labels, as well as self-learning-based methods (in particular, AutoAugment and Fast AutoAugment), which are capable of automatically selecting optimal augmentation strategies.

In the analyzed images of buildings with severe damage sustained during combat operations, significant destruction of building parts and minor damage (e.g., broken or destroyed window frames) are typically present. It turns out that minor classes (collapsed structures, cracks, plaster debris) and major classes (broken windows) are neighbors in the same frames, and simply multiplying the number of input images only proportionally increases the proportion of both types of objects without changing their relative balance. With limited tools, for example, when using the basic Roboflow plan, this problem can be solved through region-selective transformations [33]. First, targeted cropping allows us to extract individual fragments from the input images that are concentrated on minor classes. In this case, the cropped areas containing cracks or debris become independent instances, which increases their frequency in the training sample and enhances the contribution of the corresponding gradients during optimization. Given the presence of both minor and major classes in a single image, as well as certain limitations of augmentation approaches on Roboflow, it is advisable to consider addressing the class imbalance problem using other approaches. To this end, one can use a Weighted Dataloader [34], which works by adjusting the probability of selecting samples when forming training batches: elements of rare classes are assigned higher weights, increasing their frequency of appearance during training. The key advantages of this approach are its ease of use and versatility, as it requires no modification of the model architecture and affects only the data selection strategy. A more complex but more effective approach to addressing class imbalance is the two-phase learning approach, in which the model is first trained on a balanced subsample and then fine-tuned on the full, unbalanced dataset. In the first stage, methods such as Random Under Sampling and minority class

augmentation are used to create a more balanced sample, allowing the model to effectively learn to recognize rare objects. Afterward, the model is fine-tuned on the full dataset with the natural class distribution to adjust predictions based on real-world statistics and improve generalization ability. The advantage of this approach is the ability to focus on rare classes without losing information about the actual distribution. The method is versatile and easily combines with techniques for addressing class imbalance that use specialized loss functions. The main drawbacks of this approach are its increased computational time and resource requirements, which stem from training the model twice on different subsets. Overall, despite this limitation, two-phase training remains an effective and widely used approach for improving model quality in tasks with significant class imbalance.

4.3. Ranking Damaged Urban Infrastructure Objects to Determine the Priority of Their Restoration

Determining the priority of restoring damaged urban infrastructure objects requires the use of ranking methods that ensure an objective ordering of structures based on various multiple criteria. In modern research, MCDM and LTR methods are used to solve such problems, as they allow for the integration of heterogeneous data and the optimization of the decision-making process. Both approaches are applied in practical planning tasks but differ in terms of input data, implementation complexity, scalability, and the level of transparency of the results. To analyze them, we will consider the most commonly used MCDM and LTR algorithms.

MCDM methods are widely used in urban infrastructure planning, environmental science, and strategic analysis. Their main advantage is the ability to formally account for heterogeneous evaluations (quantitative and qualitative), as well as ease of implementation without the need for prior model training [35]. One of the most common methods in this group of algorithms is the TOPSIS algorithm [36]. This method involves selecting the alternative that is closest to the hypothetically best solution and, at the same time, farthest from the hypothetically worst solution. If the input data consists of m alternatives and n criteria, then the first step is to construct a decision matrix containing the value of each criterion j for each alternative i .

$$X = [x_{ij}] , \quad (1)$$

where $i=1,\dots,m$; $j=1,\dots,n$.

The next step in the algorithm involves creating a matrix of normalized values $R = [r_{ij}]$ by normalizing the decision matrix constructed in the previous step, in order to eliminate the influence of different units of measurement on the prioritization process:

$$r_{ij} = x_{ij} / \sqrt{\sum_{i=1}^m x_{ij}^2}. \quad (2)$$

To account for the relative importance of each criterion for the elements of the normalized matrix obtained in the previous step, criterion weights w_j are applied. In the current step, a weighted normalized matrix is formed:

$$V = [v_{ij}], \quad (3)$$

where $v_{ij} = w_j \times r_{ij}$; w_j – the weight value of the j -th criterion; r_{ij} – the normalized value of the j -th criterion for the i -th alternative.

Next, the best and worst decisions are identified by determining the corresponding value for each criterion.

$$A^+ = \{v_1^+, \dots, v_n^+\}, \quad (4)$$

$$A^- = \{v_1^-, \dots, v_n^-\}, \quad (5)$$

where $v_j^+ = \max_i v_{ij}$ – the best value of the criterion;

$v_j^- = \min_i v_{ij}$ – the worst value of the criterion.

At this stage, the preparatory steps are completed and the ranking process itself begins. For each alternative, the Euclidean distances to the predefined ideal and anti-ideal points are calculated.

$$D_i^+ = \sqrt{\sum_{j=1}^n (v_{ij} - v_j^+)^2}, \quad (6)$$

$$D_i^- = \sqrt{\sum_{j=1}^n (v_{ij} - v_j^-)^2}, \quad (7)$$

where D_i^+ – the distance from alternative i to the optimal solution; D_i^- – the distance from alternative i to the anti-ideal solution.

The algorithm concludes by calculating, for each alternative, the relative proximity coefficient to the optimal solution C_i :

$$C_i = \frac{D_i^-}{D_i^+ + D_i^-}, \quad 0 \leq C_i \leq 1. \quad (8)$$

It is generally accepted that the higher the value obtained, the closer the alternative is to the ideal. Therefore, the best alternative is the one with the highest relative proximity coefficient. The result of this algorithm

is a priority list containing alternatives ranked in descending order of their relative proximity coefficients.

Another widely used MCDM method is the AHP algorithm, which is used to structure complex problems and rank alternatives based on multiple criteria. The AHP method is based on decomposing the decision-making problem into a hierarchical structure consisting of an objective, criteria, subcriteria, and alternatives [35]. The basic algorithm of this method is implemented sequentially in several steps. In the first step, a pairwise comparison matrix is formed based on the results of expert surveys $A = [a_{ij}]$, in which each element reflects the relative importance of criterion i compared to criterion j . When constructing such a matrix, the principles of reciprocity and self-comparison are taken into account. The next step involves forming a vector of normalized weights as the principal eigenvector of matrix A .

$$Aw = \lambda_{\max} w, \quad (9)$$

$$\lambda_{\max} \geq n, \quad (10)$$

where $\lambda_{\max}$ – is the largest eigenvalue of the matrix A ; w – normalized main eigenvector; n – number of criteria.

To determine the normalized weight vector, it is important to first calculate the largest eigenvalue $\lambda_{\max}$:

$$\lambda_{\max} = \frac{\sum a_j w_j - n}{w_j}. \quad (11)$$

To calculate the vector w it is also possible to use the method of power products:

$$w = \lim_{k \rightarrow \infty} \frac{A^k e}{e^T A^k e}, \quad \text{where } e = (1, 1, \dots, 1)^T. \quad (12)$$

To verify the algorithm's results and the consistency of pairwise comparisons, the consistency index CI and the consistency ratio CR are calculated:

$$CI = \frac{\lambda_{\max} - n}{n - 1}, \quad (13)$$

$$CR = CI / RI, \quad (14)$$

where RI – the average consistency index of a random matrix of size $n \times n$.

The AHP method yields a vector of criterion weights that reflects the relative importance of each criterion in the decision-making process, the values of the alternatives' global priorities, and a ranked list of alternatives in order of priority from highest to lowest, based on the calculated global priorities.

The AHP method has both advantages and disadvantages [35]. Its advantages include: flexibility, stemming from the AHP method's ability to easily adapt to multi-criteria analysis tasks, taking into account both quantitative and qualitative criteria; support for collective expert participation in decision-making by calculating the geometric mean of individual pairwise comparisons; resilience to uncertainty and risk, due to the ability to form a rating scale in cases where standard measurements are insufficient or incomplete and scattered data are present; transparency of the process of calculating ranks and obtaining results. The disadvantages of this method include: the need to perform numerous pairwise comparisons in cases with a large number of criteria or alternatives; the problem of disrupting the priority of objects when adding very similar alternatives or copies of their existing variants to the set of alternatives. Due to the additive method of aggregating scores, high scores on some criteria can offset low scores on others, which can lead to excessive smoothing of results and the loss of important details.

In cases where historical information on object rankings is available or expert evaluations exist, it is advisable to use LTR approaches. These methods are widely used in information retrieval, recommendation systems, and application sorting, and can be adapted to the task of planning the priority restoration of damaged urban infrastructure objects [37]. The RankNet, LambdaRank, and LambdaMART algorithms are important representatives of this field, actively used to build models that sort or rank documents or alternatives according to their relevance or significance [38]. These methods have evolved from simple pairwise document comparisons to more complex approaches that optimize list metrics such as NDCG. An important aspect of the development of these algorithms is their ability to optimize ranking via gradient descent using various loss functions that account for the specifics of ranking.

The RankNet algorithm, proposed by Microsoft Research, is the first method in the LTR family and is based on a pairwise approach. The basic idea is that the model compares pairs of documents and determines which one is more relevant. For each pair of documents (d_i, d_j) , where d_i is supposed to rank higher than d_j , RankNet uses a sigmoid function to predict the probability that d_i will be better than d_j

$$P_{ij} = \frac{1}{1 + e^{-\sigma(s_i - s_j)}}, \quad (15)$$

where s_i and s_j are the scores generated by the model for documents d_i and d_j , respectively; σ is a model parameter.

The loss function is based on cross-entropy, which is minimized to train the model to make accurate predictions. However, this approach focuses on minimizing the error for each document pair rather than directly on the ranking metric, which is the main limitation of this method.

LambdaRank is an improvement over RankNet – it introduces lambda gradients that allow for the optimization not just of pairwise accuracy, but specifically of list-based quality metrics such as NDCG. The idea behind LambdaRank is that, instead of adjusting the model solely based on errors in pairwise comparisons, the gradients are adjusted to minimize changes in the list quality metric caused by reordering document pairs. This allows for direct optimization of the ranking, reducing the likelihood that relevant documents will end up at the bottom of the list. LambdaMART is a further development of LambdaRank and combines the ideas of λ -gradients with gradient boosting on MART trees. This algorithm uses an ensemble of decision trees, where each tree corrects the errors of the previous one. Each tree is trained on the residual errors made by the previous trees, thereby optimizing the ranking metric while accounting for the values of the λ -gradients. As a result, LambdaMART optimizes list metrics such as NDCG and provides more accurate results with large amounts of data (both categorical and numerical). However, this algorithm is computationally complex, requiring extensive parameter tuning and retraining.

The RankNet, LambdaRank, and LambdaMART algorithms successively improve the approach to Learning to Rank: from simple models that work with document pairs to more complex methods that can optimize ranking quality using boosting techniques. Compared to previous methods, LambdaMART is the most powerful, as it combines the strength of λ -gradients and gradient boosting on trees, which allows for a significant improvement in ranking accuracy with large datasets and high variability. However, all these methods require significant computational resources and careful tuning to avoid overfitting.

For prioritizing building restoration, MCDM methods (TOPSIS, AHP) are a good choice when there are multiple criteria but no labeled data, as they offer simplicity and transparency. RankNet, LambdaRank,

and LambdaMART are suitable for large datasets with prior estimates, as they are focused on optimizing ranking metrics such as NDCG. For tasks with limited data or expert estimates, MCDM methods will be more convenient. The TOPSIS method is simple to implement and provides a clear and intuitive ranking of alternatives based on their distance from ideal and anti-ideal solutions. It also works well when there is a limited amount of data and no need for in-depth modeling of dependencies between criteria.

4.4. Structure of the Proposed Intelligent System for Automated Analysis of Damage to Urban Infrastructure Objects

The intelligent system for automated detection of damage to buildings based on images is primarily designed to assess the extent of damage and prioritize their restoration, taking into account cultural and social significance. This system must be effective, scalable, and transparent, and be based on the integration of artificial intelligence principles, image analysis methods using computer vision, and contextual information to support decision-making regarding the order of reconstruction of damaged infrastructure objects in Ukraine.

The system under development aims to address the aforementioned shortcomings of existing solutions by proposing an innovative approach to assessing damage and prioritizing the restoration of buildings in Ukraine.

The system's operation involves the sequential implementation of the following key stages: identification of damaged buildings; analysis of damage and assessment of its severity; and determination of restoration priorities, taking into account the social and cultural value of each building.

Among the advantages of the system under development is, first and foremost, the ability to obtain data on building damage from multiple sources, including satellite imagery, reports from municipal authorities and emergency services, and information from volunteers. This approach allows for gathering more information about the presence of structural damage and accounting for damage to structures that may have been overlooked during the analysis of satellite imagery. Although working with satellite imagery plays an important role in identifying damaged structures – as it allows for the rapid detection of a larger number of damaged structures – it is not mandatory.

Another advantage of this system is the ability to conduct on-site surveys, which allows for a more detailed

examination of existing damage. Importantly, the system will provide not only a general assessment of the extent of damage but also a list of identified damages, enabling preliminary cost estimates for restoration work and planning for partial reconstruction of buildings in the event of insufficient funds. The system uses 3D models of buildings, which allow reconstruction specialists to examine the scope of work and the extent of damage in greater detail.

When calculating restoration priorities, the system balances a building's functionality, the extent of its damage, and its cultural value, using UNESCO standards and local registries. This ensures objectivity in the rating process, unlike similar systems from FlyPix AI and T2D2, which do not account for cultural and social aspects.

Finally, the system is transparent, giving ordinary citizens the ability to track where allocated funds are directed, thereby eliminating corruption risks. In the context of the ongoing martial law in Ukraine, where resources are limited, this advantage will ensure support from the community and donors, facilitating the rapid restoration of critical infrastructure and cultural sites.

Figure 7 shows a general diagram of the proposed system's functionality. According to the diagram in Figure 7, the process of analyzing building damage and determining the priority of their restoration proceeds as follows: first, data sources are analyzed to identify all damaged buildings and compile a corresponding list of their addresses; next, detailed information on each structure from the list is collected and analyzed; and in the final stage, a list of all surveyed buildings is created, sorted by the priority of their planned restoration, and a report is generated that includes the system's results.

To obtain information about all damaged buildings, the system is designed to interact with three different, independent primary sources. Rapid monitoring and subsequent analysis of potential damage across large areas is carried out (where possible) using satellite imagery, which enables the task of identifying and classifying damaged objects by assigning them a damage class (level). After the computer vision method completes its work, the system proceeds to the next step, which involves linking the identified building to its address and generating a list of addresses marked with the building's level of damage. Given that, in such aerial surveys, buildings are typically surveyed only from above, there is a possibility of missing buildings that have an intact roof but damaged walls and facades. To account for this type of building, the system is designed to work with two additional data sources.

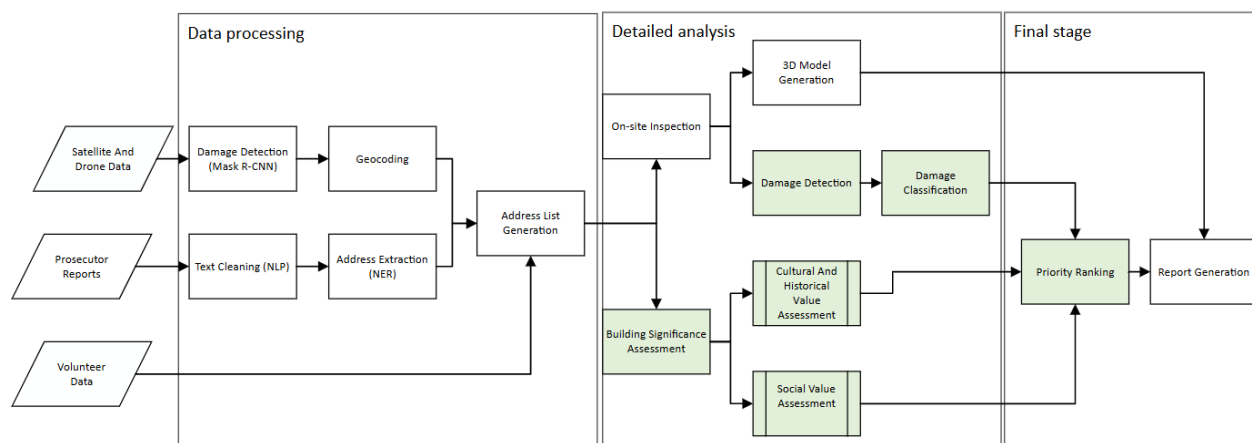


Fig. 7. General block diagram of the proposed system's functions

Typically, emergency response teams are immediately deployed to the sites of attacks to document damage to buildings. By collaborating with government agencies – such as the prosecutor's office – based on reports from these teams, it is possible to obtain addresses, photos of the damage taken from various angles, or even an approximate assessment of the extent of damage to a structure. Using these reports within the system allows for a significant refinement of the list of destroyed buildings identified through satellite imagery analysis. However, for de-occupied territories that are extensive in area and have a high level of destruction, there is often an insufficient number of such reports available. In such cases, the final information resource plays a crucial role: collaboration with volunteers and concerned citizens who can report, through established communication channels, the location of a damaged building at a specific address, subjectively assess the severity of the damage, or send photo reports. Thus, the result of the first module's work is a list of addresses generated after analyzing satellite imagery, sorted by damage severity ratings, and supplemented with information received from operational services and volunteers.

After generating this list, the system proceeds to the second stage. At this stage, all available information is searched for and collected for each building on the generated list. This stage can be divided into two independent blocks, which can therefore be executed in parallel. The first block involves working with archives and databases to obtain complete information about the building's construction history, its architectural style, the architect, significant historical events that may be associated with it, information about the building's importance to society, and its current functional role. The final step is to search for data on the type of

structure. Most buildings are constructed according to standard designs, and typically, the construction information characteristic of a standard type and series of buildings applies to the structure itself. Next, following specific guidelines, the building's socio-cultural value and its social significance are determined. Both assessments are rated on a scale from 1 to 10.

The second part of this stage involves deploying mobile teams that visit all addresses and conduct on-site inspections, during which they photograph the building from all angles to obtain the most comprehensive information about its condition. It is also permissible at this stage to record a video of a full walk-through of the building, if possible. After the collected information is entered into the system, the stage of automated visual analysis begins using computer vision methods to detect various types of damage. The next step involves determining the extent of the building's damage by correlating the identified types of damage with a ten-point scale, ranging from minor to critical damage. The video footage obtained during the survey is processed separately using algorithms that implement Structure-from-Motion (SfM) reconstruction, resulting in a detailed 3D model of the structure. This allows specialists to conduct a more detailed preliminary assessment of the scope of repair work without the need for a physical site visit.

In the final stage, the ranking method implemented in the system integrates the previously obtained indicators (level of damage, cultural and historical value, and social significance) into a single priority metric.

The result of the entire system is a priority-sorted list of damaged sites and detailed reports on each of them, containing information on defect types, a 3D model, and an assessment of sociocultural value.

4.5. Software Implementation of Project Modules

The system for prioritizing building restoration, taking into account automatically detected damage, the calculated damage severity score, its social significance, and cultural value, was created as a client-server web application, with the server-side implemented using Flask. When building this project, the principle of separation of concerns was followed; that is, the project structure is organized in such a way as to ensure its modularity, ease of maintenance, and the ability to expand functionality in the future. Figure 8 shows the structure of the project's software components in the development environment.

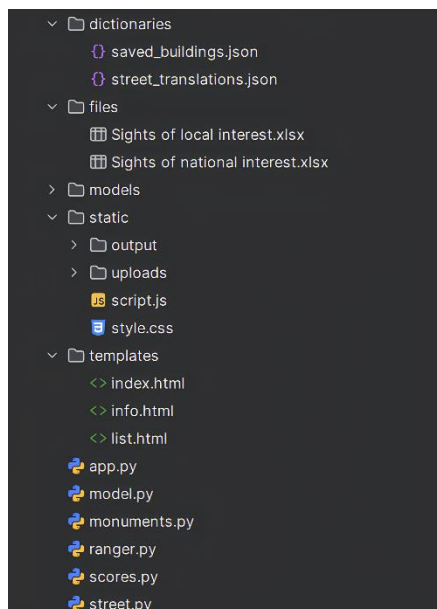


Fig. 8. Structure of the project's software components

The project consists of the following components:

- the "dictionaries" directory, which contains files with reference information required for data processing. The `saved_buildings.json` file stores records with data on damaged buildings for ranking purposes. This file serves as a database that is accessed when there is a need to display a list of all damaged addresses on a web page and to add a record after inspecting a new damaged building. The `street_translations.json` file is a user dictionary and contains the correspondence between old and new street names in the city, enabling the correct translation of street names from various languages into Ukrainian. This file is used in the application to retrieve all necessary information about a building, which comes from various sources where it is stored in different formats and has different street name conventions;

- the files directory, which contains `xlsx` files downloaded from the official website of the Ministry of Culture and Strategic Communications of Ukraine, as well as from the State Register of Immovable Monuments of Ukraine. This directory contains the files `Sights_of_local_interest.xlsx` and `Sights_of_national_interest.xlsx`, which contain information about sites of local and national cultural significance, respectively. These files are used as a data source for assessing the cultural and historical significance of buildings;

- the models directory, which contains the weights of pre-trained deep learning models that can be used to detect damage in images of buildings. This directory is critical for automating the image analysis process and generating a vocabulary that reflects damage classes and the number of instances of each class detected in an image;

- the static directory, which stores the web application's static resources. It contains the output and uploads subdirectories, used to store the results of image processing by the computer vision model and user-uploaded images, respectively. In addition, the `script.js` and `style.css` files implement the dynamic behavior and styling of the client-side application interface;

- the templates directory, which contains HTML templates for rendering the web interface. In particular, the files `index.html`, `info.html`, and `list.html` display the home page, detailed information about buildings, and a list of results, respectively;

- the `venv` library root directory, which represents a Python virtual environment used to isolate project dependencies;

The project's root directory contains scripts that form the core of the project. All project functionalities are implemented in various Python files, namely:

- `app.py`, which is the main application file that initializes the Flask server and contains defined routes for server-client interaction as well as the general logic of the web interface;

- `model.py`, which is the module responsible for loading and using a pre-trained deep learning model to detect damage in images;

- `monuments.py`, which contains the logic for working with data on historical and cultural sites and the results of processing the corresponding Excel files;

- `ranger.py`, which implements a mechanism for multi-criteria ranking of objects using the TOPSIS method;

– scores.py, which is a module for calculating three key building assessments: damage level, functional significance, and cultural value;

– street.py, which contains functions for processing addresses, including street name translations, dictionary lookups, and queries to various databases and archives to obtain all necessary information about the building.

The performance of computer vision models used to detect damage in building images depends significantly on the quality of the training dataset. Therefore, the task is to prepare such a dataset with a sufficient number of images that are correctly annotated. Given that there is a limited number of high-quality images available for the tasks of implementing the proposed system, it is advisable to combine several relevant datasets. However, this raises issues regarding the quality and detail of annotations in such image sets. In the selected combined dataset for analyzing the condition of damaged buildings, a general "damage" class was defined, which is used to assess damage of all levels. To provide the model with greater flexibility in detecting different types of damage, the system performs autonomous image annotation.

Work with the dataset was conducted on the Roboflow platform using semi-automated annotation tools. Figure 9 shows the interface window for image annotation, where the annotated image is displayed on the right and information about the labels of the

defined classes is shown on the left. Image annotation requires the development of a comprehensive list of damage types to most accurately identify different classes of damage, which vary in severity. After conducting a detailed analysis of various survey reports on damaged buildings, as well as a visual inspection of the available dataset, it was decided to define the following basic classes: Big-Crack, Medium-Crack, Small-Crack, Collapsed wall, collapsed building part, damaged roof, burn marks, damaged facade, broken balcony, broken door, and broken window. Note that the datasets subsequently used for training and testing the system's computer vision models contain a greater number of different damage classes.

The defined list of classes allows for the coverage of a wide range of damages, thereby eliminating the issue of subjectivity in determining which class a particular damage belongs to. To definitively avoid errors during dataset annotation, it was appropriate to define a damaged facade as a breach in the building's outer layer (e.g., cladding), a destroyed wall as a breach in the wall's integrity, and a destroyed part of the building as a breach involving two or more adjacent walls. A burned facade is a type of damaged facade defined to distinguish damage caused specifically by fire, which has no direct impact on the structure's stability but may indicate the presence of internal damage.

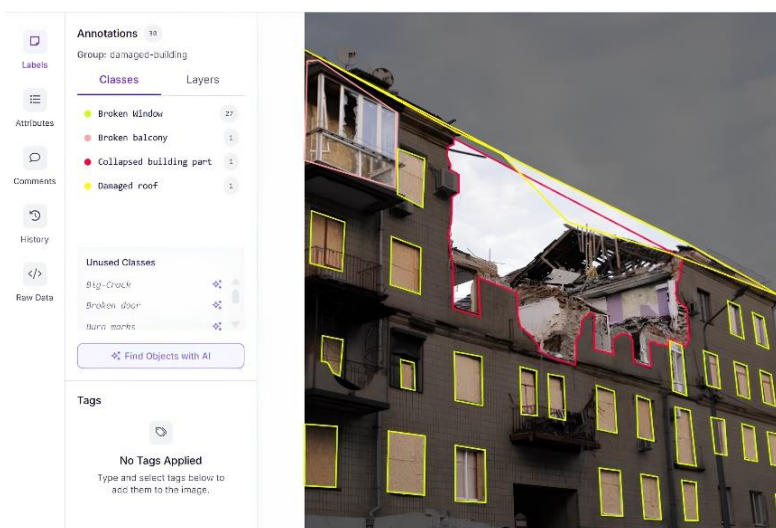


Fig. 9. The interface window of the Roboflow visualization environment

Figure 10 shows, for comparison, images with the initial annotations (left) and those resulting from the automated annotation process (right). After annotating the base test dataset, the final version contained 383 images across 11 annotation classes. A total of 12,973 annotation

masks were applied. Most of the annotations were created for the broken windows class (10,234) and the damaged facades class (1,629). All other classes have a total of between 100 and 600 annotations. The class with the fewest annotations is the class of small cracks – 114 annotations.



Fig. 10. Comparison of the original and annotated images of damaged buildings

The significant class imbalance is due to the varying nature of the damage and the different frequencies at which it occurs. This imbalance can significantly affect the accuracy of damage detection models, causing them to analyze the most prevalent classes more thoroughly than other classes. This issue was addressed later, during the model training phase. After annotating the image set is complete, the analytical information automatically generated on the Roboflow platform is analyzed. Selecting the optimal features for the dataset allows us to subsequently determine which model to prioritize and plan the preparation of the dataset for its training. To address the issue of varying image dimensions during the dataset creation phase, it is necessary to standardize the size for all images. Typically, images are resized to 640×640 , but given the large number of small-sized defects, it was decided to resize the images to 1024×1024 , which will slow down model training but improve the quality of defect detection. Figure 11 shows an Annotation Heat Map, which visualizes the distribution of annotations across the entire image space in the dataset.

As can be seen, the central region (red and orange zones) contains the highest concentration of annotations, while the periphery (blue zone) contains virtually no objects. This pattern indicates a spatial bias – most detected objects are located in the center of the images. This can lead to overfitting, where the model expects the objects needed for detection to be in the center of the image and may fail to detect them in other areas. To address this spatial bias issue, geometric transformations are typically used. However, since all images contain buildings with a fixed spatial location, the most effective approaches that can be used in our

case during the augmentation stage remain the Crop method, which will force the model to see objects in different parts of the image, and the Shear method, which allows the data to be made more geometrically diverse without distorting orientation. Using only these augmentation processes to address the previously described dataset issues, a version of the image set comprising 696 photographs was created, which can now be used to train computer vision models. To further increase the number of images, augmentation procedures such as brightness adjustment, image blurring, and noise addition can be used. This not only increases the number of images in the dataset to 4,048 but also makes the model more robust to varying shooting conditions and the quality of the input images. Thus, at this stage, image datasets were created for training models to recognize building damage.

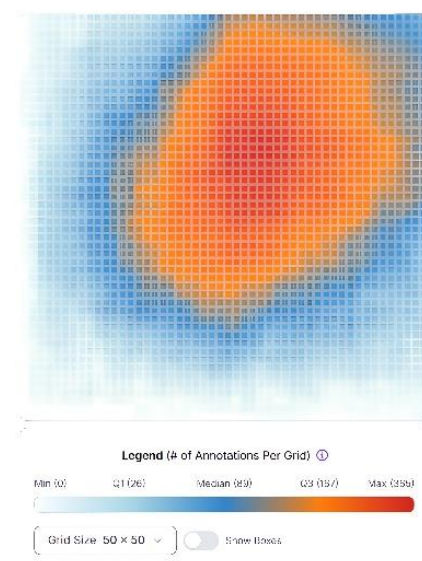


Fig. 11. Annotation density map

Let's take a closer look at the implementation details of the main modules of the proposed system: the damage detection module, the building assessment calculation module, and the object ranking module.

The main function of the *damage detection module* is the `predict_and_visualize_return_data` function, which is responsible for applying the weights of a pre-trained damage detection model to the required images and preparing the output data for further analysis and visualization. At the initial stage of the function's operation, the model weights are loaded as a dictionary, where each key corresponds to a specific model from several available options. Next, the original image in BGR format is read using the OpenCV library. The image is processed by the selected model with a specified confidence threshold of 0.5 on a CPU device. The result of this processing is an object containing information about detected bounding boxes, masks, and damage classes. In the next step, the number of objects in each class and the total area of the segmented regions are calculated. For this, a counter is used, which is incremented for each detected frame, and the pixel areas of the masks are accumulated in the variable "total_area_pixels". After obtaining all necessary statistics, the "plot" method is called to generate an annotated image with overlaid masks, which is returned in RGB format by default. Since the subsequent saving is performed using OpenCV in BGR format, a function is used that reverses the color order (RGB→BGR) when the input array has three channels. After completing all previous steps, the resulting image is saved. The save function creates and returns a dictionary with three keys: the path to the saved image, statistics on counted objects by class, and the total area of the segments in pixels. Thus, the module provides a convenient way to obtain an image with analysis results and detailed information about the number and size of detected damages.

The *building assessment calculation module* consists of three interconnected submodules that sequentially determine the cultural and functional values and the degree of damage to the building. Each submodule is implemented as a separate function that takes the object's metadata as input and returns a numerical score ranging from 1 to 10. As a result, these ratings are stored in a single scores dictionary, which is subsequently used by the ranking module. The cultural value assessment function analyzes information about the protected site's status, date of construction, architectural style, and the building's architect. First, the system checks whether the

building is listed on the UNESCO, national, or local heritage lists, with the highest score given to World Heritage sites, the middle score to sites of national significance, and the lowest to those with local or undetermined status. Next, the year of construction is analyzed, allowing for additional points to be awarded based on the structure's age: the oldest sites are given higher priority in terms of cultural significance. The architectural style is checked for inclusion in the list of rare schools, which allows for the identification of structures representing architectural movements important for preservation. After that, information about the architect is processed: mention of the participation of famous architects in the building's construction adds extra points to the score. After analyzing all available information, a final result is generated, which is adjusted depending on the type of building – cultural sites receive additional points, while ordinary residential buildings constructed according to standard designs, conversely, lower the final score. The resulting value is normalized relative to the maximum possible score and rounded to a whole number, with a range limited to between one and ten. The functional assessment is based on an analysis of the building's purpose and its current condition. First, the key "purpose/name" field of the building is analyzed to determine its type: residential, hospital, educational, commercial, or industrial. Each type corresponds to a base number of points, and the presence of a secondary commercial function (e.g., a residence with a store) provides a slight advantage over a standard residential building. Next, the current technical condition of the structure is taken into account – a building in use receives additional points, while neglected or abandoned structures, on the contrary, lower the score because they are not a priority for restoration. In the case of residential complexes and private homes, if data on the number of apartments is available, the scale of the property is taken into account: large complexes receive additional points, given that they are used by a significant number of people. Finally, a final adjustment based on the type of structure is added to the total score, after which the result is normalized to a scale of 1 to 10. The damage assessment function relies on the results of the computer vision module, which determines the number of objects in each damage class, and on construction information obtained from open sources. When calculating the assessment for each type of damage, a specific weight is assigned, reflecting the danger such damage poses to the building's stability. Next, a correction is added to the

base contribution for each type of damage, taking into account the number of objects of that class found. For large numbers, multipliers are applied, even if the damage is not severe. If additional data on the building's structural features is available (e.g., the material of exterior or load-bearing walls, the presence of a technical floor, etc.), these parameters adjust the weight of each type of damage according to its impact on the structural integrity. After accumulating the total damage score, the entire indicator is normalized and scaled to a general scale, and the resulting value is limited to a range from 1 to 10. Thus, the assessment calculation module provides a multidimensional qualitative analysis of each building, taking into account historical-cultural, socio-functional, and technical characteristics, and generates input parameters for the multi-criteria ranking model.

The *object ranking module* is critical for determining the priority of building restoration. For ranking in this module, the TOPSIS algorithm is used, with the implementation of the main function `topsis_rank` and the auxiliary function `_to_number`. The auxiliary function `_to_number` is needed only to correctly convert the input rating value to an integer, preventing the program from crashing due to incorrect or missing data. The main ranking function `topsis_rank` takes as input a list of buildings, each of which is represented by a dictionary with a nested sub-dictionary scores containing the values of three criteria: degree of damage, social significance, and cultural-historical value. Additionally, two parameters are passed – vectors of criteria for the ideal and anti-ideal examples, containing the values of the criteria considered the best and worst in terms of the priority of the object's restoration. At the start of the system's use, the criterion values for the best and worst cases were defined as follows: a building that should be restored in the near future must have the highest social significance (10 out of 10 possible), while its cultural significance must be average (5 out of 10), because a low value would not reflect its cultural value, while a high value would likely impose restrictions on restoration and require special techniques for rebuilding the building, which is costly and not always feasible in practice. The most acceptable level of building damage should be 3, because a value that is too low may indicate only damage to windows and minor damage to the facade, which does not threaten the building's stability and therefore the building can still be used for some time without restrictions; A high damage rating indicates a possible need to demolish the building,

and therefore does not require restoration work. A building that can be restored last is a socially insignificant structure with no cultural significance. If the damage level of such a building is 9, this indicates significant damage to the building. If such a building is not a cultural monument, it is advisable to demolish it rather than restore it. During the execution of the `topsis_rank` function, the following steps are performed for each object individually: retrieving the values of the three scores from the scores sub-dictionary; processing each criterion value using the `_to_number` helper function to ensure algorithm stability and prevent errors caused by missing or incorrectly formatted data; forming a score vector. Next, the distance between the current building's vector and the ideal vector is calculated, as well as between that vector and the anti-ideal vector. This determines how close or far the object under consideration is from the ideal or, conversely, from the worst possible option. Based on the values of these two distances, an aggregated proximity index to the ideal solution is calculated. This index takes a value between 0 and 1, where a value close to 1 means that the building is closer to the ideal scenario and, accordingly, has a higher priority in the ranking. After calculating the index for each building, the list of all building records is sorted in descending order of this index, meaning that the highest-priority buildings appear at the top of the list. Additionally, each object is assigned a `priority_rank` field, which determines its order number in the sorted list, starting with 1 for the object with the highest index. Thus, as a result of executing the specified function, two additional values are added for each building in the input list: `topsis_score`, which contains a numerical indicator of proximity to the optimal solution, and `priority_rank`, which indicates the building's position in the overall priority ranking. The function returns an updated list of buildings, sorted by importance, which is then displayed to the user. This approach, using the `topsis_rank` function, allows for more informed decisions regarding the order of building restoration, taking into account not only their technical condition but also their social and cultural significance.

4.6. Server-side component of the system

The server-side component of the application is implemented as a Flask application and serves as the single entry point for all client requests. It handles HTTP request routing, coordinates calls to all necessary functions for retrieving and preparing data from external services, calculating scores, and saving them.

The main route at the path "/" is required to process the GET request and read the saved_buildings.json file, which contains saved records of buildings with preliminary analysis results. If such a file exists, its contents are deserialized into a list of objects; otherwise, an empty array is created. In the next step, this list is passed to the list.html template, which is responsible for displaying a tabular list of buildings with their addresses and previously calculated ratings. A GET request to the /add route returns an empty form implemented in the index.html template, where the user can upload a photo and enter the building's address.

After submitting the form, a POST request is sent to the /upload route. The image file and a string containing the building's full address are sent to the server, after which the function handling this route checks for the image's existence and saves it to the static/uploads folder. Next, the predict_and_visualize_return_data function is called, which is responsible for running the pre-trained damage detection model on the user-uploaded image, and then returns the path to the saved annotated image, the count of detected damage classes, and the total area of the masks.

In the next step, the house number and street name are extracted from the received string containing the building address, after which a request to an external source (web scraping) is initiated via the get_building_info function, which returns a dictionary of keys containing basic and additional information about the building obtained from open sources. Fields for both types of information are filtered according to a predefined list, retaining only relevant attributes. If the external search yields information about the building's series and design, the get_project_info function is called, which searches for the necessary information about standard building series and designs and returns a dictionary with the retrieved data, which is also stored in the session. The find_monument_info function is used to search for data regarding the status of the historical and cultural monument. After all necessary data has been prepared, the user is redirected to the /info route, where the info.html template displays: an annotated photo, a damage counter, and basic and additional information about the building.

The /scores route is implemented as an API method that returns the results of three evaluation functions – estimate_damage_severity, estimate_functional_value, and estimate_cultural_value. It passes the necessary data previously stored in the session to the corresponding submodules, returning a JSON object with three numeric fields: damage, functional, and cultural. This route

implements asynchronous updates of building assessment values on the client side without reloading the web page. To save the retrieved building information and calculated assessments, the user makes a POST request to the /save route. After that, the server reads the address and damage statistics from the session, and the final values of the three ratings from the request body, then sends a new record to the saved_buildings.json file and confirms successful saving. The GET route /list, which replicates the functionality of the main page, is designed to display only records that have already been saved. The /topsis route triggers a re-ranking of all records when a new building record is received from the user. After the re-ranking is complete, the session rewrites the saved_buildings.json file according to the new restoration order and redirects the user back to /list, where the updated ranking is displayed.

Thus, the server-side component of the Flask-based application provides a complete data lifecycle: from image loading and analysis (through external data processing and score calculation) to their storage and multi-criteria ranking. This architecture allows for scaling functionality and integrating new services or modules without significant code refactoring.

4.7. System User Interface

Upon launching the application, the user immediately receives a list of buildings that were damaged during military operations and require restoration, presented as a pre-sorted list (from buildings requiring the most urgent intervention to those that can wait for new investments). Figure 12 shows the appearance of the application's main page.

After reviewing the list of damaged buildings, the user can add a new entry – register a new damaged building – by clicking the "Add Entry" button. After clicking this button, the user is redirected to a page for adding a record to enter the building's address and upload one or more photos of the building from different angles. After the user clicks the "Submit" button, the uploaded information is processed and sent to the server. A search is conducted at the specified address for general information about the building, information about the standard design according to which it was built, and a check for the address's presence in the registers of local and national heritage sites. At the same time, the uploaded images are processed for automatic damage detection. The status of these processes is displayed on a separate page to allow users to track the progress.

List of saved buildings

Priority	Address	Number of damages	Destruction	Func. value	Hist.-cult. value
1	Constitutsii Square, 1	Broken Window: 71 Damaged roof: 1	4	7	6
2	Rymarska Street, 20	Broken Window: 33	2	7	5
3	Rymarska Street, 19	Broken Window: 17	1	7	6
5	Lesia Serdiuka Street, 54	Broken Window: 72 Broken balcony: 2	3	7	1
6	Heroiv Kharkova Avenue, 94	Broken Window: 7 Damaged facade: 1	2	5	5
7	Sumska Street, 25	Broken Window: 14	1	5	3
8	Yaroslava Mudroho Street, 15	Broken Door: 1 Broken Window: 12 Damaged facade: 1	2	5	1
10	Yaroslava Mudroho Street, 13	Broken Window: 18 Damaged facade: 1	3	5	1
11	Sumska Street, 61	Broken Window: 11	1	1	7

Add entry
Rank

Fig. 12. The app's home page

Figures 13 and 14 show what this page looks like after the user has uploaded all the necessary data and sent it to the server. After reviewing the building information, the user can calculate the ratings needed for ranking by clicking the "Get Ratings" button. A pop-up

window (Fig. 15) displays the automatically calculated ratings and comments, which the user can save.

The pop-up window (Fig. 14) displays automatically calculated estimates and comments, which the user can save.



Building information

Balcony/loggia	balconies
Ceiling height	2.48 m
Internal walls	concrete multi-hollow panels (0.22 m), gypsum-concrete partitions (0.08 m)
Additionally	—
External walls	brick, silicate brick (thickness 0.35–0.45 m)
Facade exterior	sometimes ceramic tiles
Apartments in bldg.	1-, 2-, 3-room apartments
Number of entrances	2-8
Floor slabs material	oval-hollow or hipped reinforced concrete panels
Disadvantages	small kitchens, moral and physical aging of the building series, cracking of the external brick load-bearing walls due to poorly burnt brick

Fig. 13. The first section of the page containing information about the building

Detected damages

- Broken Window: 16

Basic information

Name/purpose:: residential building

Current status:: in use

Series:: 163

Cultural and historical value

Monument of local importance: —

Monument of national importance: —

UNESCO heritage: —

Get estimates

List of damages

Fig. 14. The second part of the page containing information about the building

Detected damages

- Broken Window: 16

Building estimates

Degree of destruction: 1/10 (minor)

Functional value: 6/10 (medium)

Historical and cultural value: 1/10 (minor)

✓ Saved

Cultural and historical value

Monument of local importance: —

Monument of national importance: —

UNESCO heritage: —

List of damages

Fig. 15. Pop-up window displaying the estimates calculated for the building after the information has been saved

After saving the building record – including its address and ratings – the user can navigate to the page listing all buildings to verify that the building has been added to the list of buildings registered for restoration,

but that a restoration priority has not yet been assigned to it. An example of the page listing buildings in this state is shown in Fig. 16.

5	Lesia Serdiuka Street, 54	Broken Window: 72 Broken balcony: 2	3	7	1
6	Heroiv Kharkova Avenue, 94	Broken Window: 7 Damaged facade: 1	2	5	5
7	Sumska Street, 25	Broken Window: 14	1	5	3
8	Yaroslava Mudroho Street, 15	Broken Door: 1 Broken Window: 17 Damaged facade: 1	2	5	1
10	Yaroslava Mudroho Street, 13	Broken Window: 10 Damaged facade: 1	3	5	1
11	Sumska Street, 61	Broken Window: 11	1	1	7
12	Metrobudivnykiv Street, 6	Broken Window: 25 Broken balcony: 2 Collapsed building part: 1	8	7	1
13	Rymarska Street, 21	Broken Window: 19	1	1	5
14	Pushkinska Street, 46	Broken Window: 6	1	1	2
	Metrobudivnykiv Street, 29	Broken Window: 16			

Add entryRank

Fig. 16. Example of a page displaying a list of buildings after creating a record for the added building

To include the added building in the restoration priority list, the user must click the "Rank" button, after which the list of buildings will be updated to reflect the newly added entry.

5. Research Results

To evaluate the effectiveness of object detection and segmentation models under conditions of severe class imbalance in the dataset, two architectures – YOLOv11 and Mask R-CNN (implemented via the Detectron2 framework) – were trained and investigated. Model training was conducted using several strategies aimed at improving the recognition quality of objects from rare classes, namely: the baseline approach, data augmentation, the use of a weighted dataloader, and a combination of augmentation with a weighted dataloader. We will now examine the specifics of training and tuning the selected instance segmentation models, as well as perform a comparative analysis of the detection quality for different types of damage using the applied architectures, taking into account the class imbalance.

Images were loaded via the Roboflow API in the format required by the models, and model training was conducted in the Colab Notebook cloud environment.

The first model considered in the experiment was one of the latest models in the YOLO family – YOLOv11 with segmentation weights. This model was trained using a standard configuration adapted to the tasks at hand, employing loss functions characteristic of this architecture (including box loss, classification loss, objectness loss, and dfl loss), which are already built into the model for faster and more accurate training on image datasets with complex scenes.

To implement the second model selected for the study, Mask R-CNN, the Detectron2 library was used, which provides a convenient interface for training and evaluating various computer vision models. This model was also loaded and configured with default parameters.

Various approaches were considered to address class imbalance. The first approach was data augmentation, which included both transformations aimed at increasing data diversity (adding noise, altering brightness and contrast, random scaling) and random cropping of images. This allowed us to diversify the input images, increasing their number in the dataset and forcing the models to adapt to varying image quality and different shooting formats. The second approach allowed the model to account for small objects and minority objects, which occupied too little area in the original large images for effective recognition. This approach did not require additional tuning of the model training processes but was implemented by loading a different version of the dataset with the changes incorporated into it.

The third approach to solving this problem involves using a weighted dataloader, which selects training examples based on class weights and increases the probability of images containing rare objects being included in a batch. This approach was implemented on the first version of the dataset, which was not augmented, requiring the addition of the YOLOWeightedDataset and CustomTrainer datasets to the main training process for the YOLOv11 and Mask R-CNN models, respectively [39].

The final approach used a combination of data augmentation and the weighted dataloader. To implement it, the models were trained on an augmented dataset using the aforementioned classes to load images with weights. The resulting mAP@0.5, and mAP@0.5:0.95 scores for each combination of model and class imbalance mitigation method are shown in Table 1.

Table 1. Evaluations of the effectiveness of damage detection in models using different approaches to reducing class imbalance

	Model	An approach to addressing the issue of class imbalance			
		No configuration	Data augmentation	Weighted dataloader	Combination of approaches
mAP@0.5	YOLOv11	0.65	0.74	0.89	0.89
mAP@0.5:0.95		0.31	0.37	0.43	0.43
mAP@0.5	Mask-R-CNN	0.61	0.73	0.90	0.91
mAP@0.5:0.95		0.29	0.35	0.44	0.45

An analysis of the results for the metrics mAP@0.5 and mAP@0.5:0.95 shows that the use of methods to address class imbalance significantly improves the quality of predictions for both models. Data augmentation provides a noticeable improvement compared to baseline training,

which is explained by the improvement in the models' generalization ability and the increased diversity of training examples. The use of a weighted dataloader demonstrates further improvement in metrics, driven by more balanced training, where models pay more attention to rare classes.

For Mask R-CNN, the combination of augmentation and a weighted dataloader yielded the best results, both for this model and among all experiments conducted. In the case of YOLOv11, combining these methods did not lead to a significant improvement in the mAP metric compared to using only the weighted dataloader.

Overall, the results presented in Table 1 confirm that the proposed solutions to the class imbalance problem are effective tools and help significantly improve the damage detection accuracy of both models.

To illustrate the damage detection results of the models, collages with four images are provided below, where the top section shows the results for the YOLOv11 model, and the bottom section shows the results for the Mask-R-CNN model. The images on the left represent the worst results obtained (for the approach without class imbalance mitigation methods), while those on the right show the best model performance (for the combination of data augmentation approaches and the use of a weighted dataloader). Images that were problematic to process were selected for the experiment (Fig. 17).

According to Fig. 17, the YOLO v11 model was able to detect a crack even without the necessary configuration, while the Mask-R-CNN model did not detect it. After completing training using a combination of approaches to address the class imbalance issue, the accuracy of damage detection by the models changed. Both models were able to detect all damage, but Mask-R-CNN showed better results (it detected the boundaries of the damage more clearly and with greater probability).

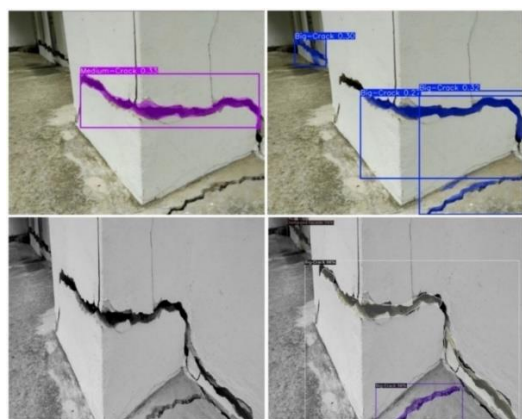


Fig. 17. Results of crack detection by the YOLOv11 (top row) and Mask-R-CNN (bottom row) models before and after applying a combined approach to address class imbalance

Figure 18 shows the results of the models on an image of a distant building containing a large number of major-class instances and a single minor-class

instance. After training with default settings, the YOLOv11 model detected only the damaged windows in the building, i.e., it detected major-class instances. The Mask-R-CNN model correctly detected the boundaries for the minor class but misclassified the damage type and failed to detect instances of the minor class that had fewer pixels in the image. After training the models to account for class imbalance, both models significantly improved their performance. The YOLOv11 model was able to detect the minor class and correctly define its boundaries, while the Mask-R-CNN model correctly identified the minor class and began detecting the most obvious window damage.

Although the YOLO model was able to recognize smaller damage instances and detect a total of more instances in this case, it is difficult to declare it the clear winner. It is worth noting that failing to detect a damaged window is less serious than identifying, with lower confidence, a more severe class of damage that is actually minor. The Mask-R-CNN model identified the damage with greater confidence in both cases.



Fig. 18. Damage detection results using the YOLOv11 and Mask-R-CNN models on complex scenes before and after applying a combined approach to address class imbalance

6. Discussion of the results

After conducting a visual comparison of the models' results and analyzing the metrics obtained, we can conclude that the use of methods to address class imbalance has a positive impact, particularly on images of buildings with damage that is difficult for models to

detect. Overall, both the YOLOv11 and Mask R-CNN models demonstrated high performance metrics, confirming their suitability for practical application in damage detection tasks provided that class imbalance mitigation methods are employed. Given the differences in architecture and the specifics of image processing, it is advisable to allow the user to select a model depending on the type of damage and analysis conditions. This approach will enable the system to be adapted to specific scenarios and improve the accuracy of results under real-world operating conditions.

7. Conclusions

As part of this work, a concept was developed and a prototype system was implemented for the automated detection of damage to buildings in urban infrastructure based on computer vision methods to determine the priority of building restoration, taking into account social significance, the degree of damage, and cultural and historical value. The system's performance was tested on real images of damaged buildings in the city of Kharkiv, resulting in a prioritized list of buildings registered in the system for restoration. These results can be utilized within the university's current projects, as well as in collaboration with government agencies and research institutions working on the restoration of damaged areas.

The proposed system is a development and adaptation of modern international solutions in the field of automatic analysis of infrastructure damage and multi-criteria decision-making, taking into account the specifics of the military conflict in Ukraine. Unlike existing analogues, the proposed system allows for the integration of processes for collecting available information on damaged buildings, automatically detecting damage, and prioritizing their restoration. The uniqueness of this system lies in the proposed approach to determining the priority of a building's restoration, taking into account criteria such as its social significance, historical and cultural value, and degree of damage. Experimental studies of the YOLOv11 and Mask R-CNN computer

vision models and analysis of the obtained metrics allow us to conclude that the use of methods to address class imbalance has a positive effect, especially for images of buildings with damage that is difficult to detect. Overall, both models demonstrated high quality metrics, confirming their suitability for practical application in damage detection tasks. The results show that the proposed system works effectively when class imbalance mitigation approaches are applied. Given the differences in architecture and the specifics of image processing, it is advisable to allow the user to select a model depending on the type of damage and analysis conditions.

Future research is planned to further improve the functionality and performance of the proposed system. In the long term, the system could become an effective tool for supporting decision-making at the national and regional levels regarding the restoration of damaged structures, helping to accelerate the restoration of critical infrastructure and minimize the socioeconomic losses associated with significant destruction resulting from hostilities.

Conflict of interest

The authors declare that they have no conflicts of interest, including financial, personal, authorial, or any other nature, that could influence the research or the results published in this article.

Funding

The study was conducted without financial support.

Data availability

The manuscript has no associated data.

Use of artificial intelligence

The Grammarly service (v1.2.254.1880) was used for spell-checking.

References

1. Gupta, R., Goodman, B., Patel, N., Hosfelt, R., Sajeew, S., Heim, E., Doshi, J., Curtis, K., Seferbekov, S., Alker, A. and Lee, S. (2019), "xBD: A Dataset for Assessing Building Damage from Satellite Imagery", *arXiv*, pp. 1–9. DOI: <https://doi.org/10.48550/arXiv.1911.09296>
2. Risso, M., Goffi, A., Motetti, B.A., Burrello, A., Bove, J.B., Macii, E., Poncino, M., Pagliari, D.J. and Maffei, G. (2025), "Building Damage Assessment in Conflict Zones: A Deep Learning Approach Using Geospatial Sub-Meter Resolution Data", in *Proceedings of the 2025 IEEE 6th International Conference on Image Processing, Applications and Systems (IPAS)*, pp. 1–6. DOI: <https://doi.org/10.1109/IPAS63548.2025.10924534>

3. Yuan, Q., Shen, H., Li, T., Li, Z., Li, S., Jiang, Y., Xu, H., Tan, W., Yang, Q., Wang, J., Gao, J. and Zhang, L. (2020), "Deep learning in environmental remote sensing: Achievements and challenges", *Remote Sensing of Environment*, 241, 111716. DOI: <https://doi.org/10.1016/j.rse.2020.111716>
4. Shevchenko, L. (2022), "Mass Housing in Ukraine in the Second Half of the 20th Century", *Docomomo Journal*, 67, pp. 72–78. DOI: <https://doi.org/10.52200/docomomo.67.08>
5. Teng, S., Chen, G., Liu, Z., Cheng, L. and Sun, X. (2021), "Multi-sensor and decision-level fusion-based structural damage detection using a one-dimensional convolutional neural network", *Sensors*, 21(12), 3950. DOI: <https://doi.org/10.3390/s21123950>
6. UN News (2025), "Ukraine: Post-War Reconstruction Set to Cost \$524 Billion" [online]. Available from: <https://news.un.org/en/story/2025/02/1160466> (Accessed: 2 April 2025).
7. UNDP Ukraine (n.d.) "In Ukraine, Machine-Learning Algorithms and Big Data Scans Used to Identify War-Damaged Infrastructure" [online]. Available from: <https://www.undp.org/ukraine/blog/ukraine-machine-learning-algorithms-and-big-data-scans-used-identify-war-damaged-infrastructure> (Accessed: 11 March 2025).
8. Geospatial World (n.d.) "Assessing Infra Damage in War-Torn Ukraine" [online]. Available from: <https://geospatialworld.net/prime/technology-and-innovation/assessing-infra-damage-war-torn-ukraine/> (Accessed: 11 March 2025).
9. FlyPix (n.d.) "AI-Driven Building Damage Assessment: Innovations and Applications" [online]. Available from: <https://flypix.ai/blog/building-damage-assessment/> (Accessed: 11 April 2025).
10. T2D2 (n.d.) "AI Damage Detection" [online]. Available from: <https://t2d2.ai/damage-detector> (Accessed: 11 March 2025).
11. WarUpClose (n.d.) "Memorialization of Destructions" [online]. Available from: <https://war.city/about-us/> (Accessed: 23 April 2025).
12. Ukrainska Pravda (2023), "Rebuilding of Housing Destroyed by Russians Can Now Be Seen Online in 3D" [online]. Available from: <https://www.pravda.com.ua/eng/news/2023/10/3/7422472/> (Accessed: 23 April 2025).
13. SCAN UA (n.d.) "Save Ukrainian Heritage" [online]. Available from: <https://scanua.com/> (Accessed: 24 April 2025).
14. Kashtan, V. and Hnatushenko, V. (2023), "Automated building damage detection on digital imagery using machine learning", *Naukovyi Visnyk Natsionalnoho Hirnychoho Universytetu*, 6, pp. 134–140. DOI: <https://doi.org/10.33271/nvngu/2023-6/134>
15. Li, X., Caragea, D., Zhang, H. and Imran, M. (2018), "Localizing and quantifying damage in social media images", in *Proceedings of the IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining (ASONAM)*, Barcelona, 28–31 August 2018. DOI: <https://doi.org/10.1109/asonam.2018.8508298>
16. Li, X., Caragea, D., Zhang, H. and Imran, M. (2019), "Localizing and quantifying infrastructure damage using class activation mapping approaches", *Social Network Analysis and Mining*, Vol. 9, No. 1, pp. 194–201. DOI: <https://doi.org/10.1007/s13278-019-0588-4>
17. Nia, K.R. and Mori, G. (2017), "Building damage assessment using deep learning and ground-level image data", in *Proceedings of the 14th Conference on Computer and Robot Vision (CRV)*, Edmonton, 16–19 May 2017, pp. 95–102. DOI: <https://doi.org/10.1109/CRV.2017.54>
18. Klyshnia, A. (n.d.) "First Step Towards Circular Reconstruction: Damage Assessment and Material Mapping of Buildings in Ukraine Through Ground Imagery Processing" [online]. Available from: <https://blog.iaac.net/first-step-towards-circular-reconstruction-damage-assessment-and-material-mapping-of-buildings-in-ukraine-through-ground-imagery-processing/> (Accessed: 24 April 2025).
19. Liu, C., Wang, Y., Zhang, J. and Li, H. (2022), "Real-time ground-level building damage detection based on lightweight and accurate YOLOv5 using terrestrial images", *Remote Sensing*, Vol. 14, No. 12, 2763. DOI: <https://doi.org/10.3390/rs14122763>
20. Bošnjak, A. and Jajac, N. (2023), "Determining priorities in infrastructure management using multicriteria decision analysis", *Sustainability*, Vol. 15, No. 20, 14953. DOI: <https://doi.org/10.3390/su152014953>
21. Mohammadnazari, Z., Mousapour, M., Alipour-Vaezi, M., Aghsami, A., Jolai, F. and Yazdani, M. (2022), "Prioritizing post-disaster reconstruction projects using an integrated multi-criteria decision-making approach: a case study", *Buildings*, Vol. 12, No. 2, 136. DOI: <https://doi.org/10.3390/buildings12020136>
22. Szeliski, R. (2022), *Computer Vision: Algorithms and Applications*, 2nd ed. Cham: Springer. DOI: <https://doi.org/10.1007/978-3-030-34372-9>
23. Liu, L., Ouyang, W., Wang, X., Fieguth, P., Chen, J., Liu, X. and Pietikäinen, M. (2020), "Deep learning for generic object detection: a survey", *International Journal of Computer Vision*, 128(2), pp. 261–318. DOI: <https://doi.org/10.1007/s11263-019-01247-4>
24. Krizhevsky, A., Sutskever, I. and Hinton, G.E. (2012), "ImageNet Classification with Deep Convolutional Neural Networks", *Advances in Neural Information Processing Systems (NeurIPS)*, Vol. 25. DOI: <https://doi.org/10.1145/3065386>
25. He, K., Zhang, X., Ren, S. and Sun, J. (2016), "Deep Residual Learning for Image Recognition", in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 770–778. DOI: <https://doi.org/10.1109/CVPR.2016.90>
26. Manakitsa, N., Maraslidis, G.S., Moysis, L. and Fragulis, G.F. (2024), "A review of machine learning and deep learning for object detection, semantic segmentation, and human action recognition in machine and robotic vision", *Technologies*, Vol. 12, No. 2, 15. DOI: <https://doi.org/10.3390/technologies12020015>
27. Girshick, R. (2015), "Fast R-CNN", in *Proceedings of the IEEE International Conference on Computer Vision (ICCV)*, pp. 1440–1448. DOI: <https://doi.org/10.1109/ICCV.2015.169>

28. Sultana, F., Sufian, A. and Dutta, P. (2020), "Evolution of image segmentation using deep convolutional neural networks: a survey", *Knowledge-Based Systems*, 201–202, 106062. DOI: <https://doi.org/10.1016/j.knosys.2020.106062>
29. Cheng, B., Schwing, A. and Kirillov, A. (2022), "Masked-attention Mask Transformer for Universal Image Segmentation (Mask2Former)", in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 1290–1299. DOI: <https://doi.org/10.48550/arXiv.2112.01527>
30. Wang, C.Y. et al. (2022) "YOLOv7: Trainable bag-of-freebies for real-time object detection", in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. DOI: <https://doi.org/10.48550/arXiv.2207.02696>
31. Zhao, Q. and Zhu, J. (2025), "An Improved YOLOv11 Architecture with Multi-Scale Attention and Spatial Fusion for Fine-Grained Residual Detection", *Results in Engineering*, 27, 107061. DOI: <https://doi.org/10.1016/j.rineng.2025.107061>
32. Huang, H., Zhao, S., Zhang, D. and Chen, J. (2020), "Deep learning-based instance segmentation of cracks from shield tunnel lining images", *Structure and Infrastructure Engineering*. DOI: <https://doi.org/10.1080/15732479.2020.1838559>
33. Kumar, T., Brennan, R., Mileo, A. and Bendeche, M. (2024), "Image data augmentation approaches: a comprehensive survey and future directions", *IEEE Access*, 12, pp. 187536–187571. DOI: <https://doi.org/10.1109/ACCESS.2024.3470122>
34. Johnson, J. and Khoshgoftaar, T. (2019), "Survey on deep learning with class imbalance", *Journal of Big Data*, 6(1), pp. 1–54. DOI: <https://doi.org/10.1186/s40537-019-0192-5>
35. Taherdoost, H. and Madanchian, M. (2023), "Multi-Criteria Decision Making (MCDM) methods and concepts", *Encyclopedia*, 3(1), pp. 77–87. DOI: <https://doi.org/10.3390/encyclopedia3010006>
36. Madanchian, M. and Taherdoost, H. (2023), "A comprehensive guide to the TOPSIS method for multi-criteria decision making", *Sustainable Social Development*, 1(1), pp. 1–6. DOI: <https://doi.org/10.54517/ssd.v1i1.2220>
37. Jagerman, R., Oosterhuis, H. and de Rijke, M. (2022), "To model or to intervene: A comparison of counterfactual and online learning to rank from user interactions", in *Proceedings of the ACM International Conference on Web Search and Data Mining*. DOI: <https://doi.org/10.1145/3331184.3331269>
38. Aftab, S. and Ramampiaro, H. (2024), "Improving top-N recommendations using batch approximation for weighted pair-wise loss", *Machine Learning with Applications*, 15, 100520. DOI: <https://doi.org/10.1016/j.mlwa.2023.100520>
39. Buda, M., Maki, A. and Mazurowski, M.A. (2018), "A systematic study of the class imbalance problem in convolutional neural networks", *Neural Networks*, 106, pp. 249–259. DOI: <https://doi.org/10.1016/j.neunet.2018.07.011>

Received (Надійшла) 14.09.2025

Accepted for publication (Прийнята до друку) 21.05.2026

Publication date (Дата публікації) 27.06.2026

Відомості про авторів / About the Authors

Пиріг Наталія Янівна – Харківський національний університет радіоелектроніки, студентка кафедри штучного інтелекту, Харків, Україна;

Natalia Pyrih – Kharkiv National University of Radio Electronics, Student of the Artificial Intelligence Department, Kharkiv, Ukraine;

e-mail: nataliia.pyrih@nure.ua

ORCID ID: <https://orcid.org/0009-0001-8774-762X>

Удовенко Сергій Григорович – доктор технічних наук, професор, Харківський національний економічний університет ім. С. Кузнеця, завідувач кафедри інформатики та комп'ютерної техніки, Харків, Україна;

Serhii Udovenko – Doctor of Technical Sciences, Professor, Simon Kuznets Kharkiv National University of Economics, Head of the Informatics and Computer Technology Department, Kharkiv, Ukraine;

e-mail: udovenkosg@gmail.com

ORCID ID: <https://orcid.org/0000-0002-0820-8450>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=55975679100>

Гриньова Олена Євгенівна – Харківський національний університет радіоелектроніки, старший викладач кафедри штучного інтелекту, Харків, Україна.

Olena Grinyova – Kharkiv National University of Radio Electronics, Senior Lecturer of the Artificial Intelligence Department, Kharkiv, Ukraine;

e-mail: olena.grinyova@nure.ua

ORCID ID: <https://orcid.org/0000-0002-3367-8067>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=6503860961>

Чала Лариса Ернестівна – кандидат технічних наук, доцент, Харківський національний університет радіоелектроніки, завідувачка кафедри штучного інтелекту, Харків, Україна;

Larysa Chala – Candidate of Technical Sciences, Associate Professor, Kharkiv National University of Radio Electronics, Head of the Artificial Intelligence Department, Kharkiv, Ukraine;

e-mail: larysa.chala@nure.ua

ORCID ID: <https://orcid.org/0000-0002-9890-4790>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=57194214849>

АВТОМАТИЗОВАНА СИСТЕМА АНАЛІЗУ ПОШКОДЖЕНЬ ОБ'ЄКТІВ МІСЬКОЇ ІНФРАСТРУКТУРИ З ВИКОРИСТАННЯМ МОДЕЛЕЙ КОМП'ЮТЕРНОГО ЗОРУ

Нині актуальним є питання забезпечення інтеграції багаторівневих даних для ефективного моніторингу й аналізу пошкоджень будівель унаслідок воєнних дій на території України, а також визначення пріоритетів відновлювальних робіт. **Метою дослідження** є розроблення інтелектуальної системи автоматизованого аналізу пошкоджень об'єктів міської інфраструктури, яка б дала змогу ефективно обробляти великі обсяги зображень, отриманих з використанням засобів комп'ютерного зору, для подальшого оцінювання рівня руйнувань і прийняття рішень щодо пріоритетності їх відновлення. Надані системою детальні звіти щодо стану кожної будівлі суттєво спрощують процес експертного оцінювання без необхідності фізичної присутності фахівців на місці. Ранжування пошкоджених об'єктів за рівнем їх впливу на життєдіяльність країни забезпечує більш раціональний і зважений розподіл ресурсів, необхідних для відновлення інфраструктури. Крім того, в системі здійснюється активна взаємодія з профільними державними структурами, органами місцевого самоврядування та фахівцями, відповідальними за планування й організацію відновлювальних робіт. У межах такої взаємодії передбачається отримання доступу до комплексних баз даних з інформацією про архітектурні особливості будівель, їх історико-культурну цінність, а також використання наявної технічної документації, що дає змогу брати до уваги всі аспекти під час оцінювання серйозності пошкоджень і формування планів відновлення будівель. Аналіз інформації, отриманої під час спеціалізованих інженерних обстежень і досліджень із застосуванням високоточних інструментальних методів, допомагає додатково уточнювати параметри ушкоджень і створювати розширені датасети для подальшого навчання й валідації моделей комп'ютерного зору. Досліджено та обґрунтовано доцільність використання сегментації екземплярів за допомогою моделей YOLOv11 і Mask R-CNN, що сприяє не лише виявленню пошкоджень, а й точному визначенню їх обсягів, що є важливим для всебічного оцінювання масштабу руйнувань і планування заходів відновлення. **Результати дослідження** демонструють, що запропонована система ефективно працює за умови застосування підходів до боротьби з дисбалансом класів. Важливим етапом роботи системи є автоматизоване формування оптимізованих планів відновлення, оснований на комплексному аналізі виявлених ушкоджень. Такі плани мають зважати не тільки на технічні параметри руйнувань, а й на фінансові, матеріальні та кадрові обмеження. **Висновок.** Розроблена система є ефективним інструментом для підтримки ухвалення рішень на державному й регіональному рівнях про відновлення пошкоджених споруд, сприяє прискоренню процесів відновлення критично важливої інфраструктури й мінімізації соціально-економічних втрат, пов'язаних із воєнними руйнуваннями.

Ключові слова: аналіз пошкоджень; міська інфраструктура; глибоке навчання; моделі комп'ютерного зору; сегментація екземплярів; набір зображень; інтелектуальна система.

Бібліографічні описи / Bibliographic descriptions

Пиріг Н. Я., Удовенко С. Г., Гриньова О. Є., Чала Л. Е. Автоматизована система аналізу пошкоджень об'єктів міської інфраструктури з використанням моделей комп'ютерного зору. *Сучасний стан наукових досліджень та технологій в промисловості*. 2026. № 2 (36). С. 121–152. DOI: <https://doi.org/10.30837/2522-9818.2026.2.121>

Pyrih, N., Udovenko, S., Grinyova, O., Chala, L. (2026), "Automated system for damage analysis of urban infrastructure objects using computer vision models", *Innovative Technologies and Scientific Solutions for Industries*, No. 2 (36), P. 121–152. DOI: <https://doi.org/10.30837/2522-9818.2026.2.121>

Pavlo Radiuk, Anatoliy Sachenko, Oleksandr Melnychenko, Ruslan Brukhanskyi, Antonina Kashtalian

WEDD-RST: AN INTERPRETIVE APPROACH TO DETECTING DEFECTS IN PHOTOVOLTAIC PANELS IN CYBER-PHYSICAL SYSTEMS

The **subject** of this study focuses on the interpretability of fault-detection frameworks in solar cyber-physical systems. The rapid expansion of renewable energy networks generates massive volumes of continuous sensor data, requiring sophisticated monitoring strategies to prevent safety hazards such as photovoltaic fires. Nevertheless, dominant deep learning models function as opaque black boxes, providing no transparent reasoning for critical safety decisions. The **goal** of this work is to enhance the interpretability of anomaly detection in cyber-physical environments by formulating a rule-based production knowledge base straight from continuous sensor readings. Key **tasks** include developing an adaptive discretisation method, identifying minimal feature subsets, and reconciling contradictory patterns using an integral class support score. The applied **methods** fuses rough set theory (RST) with an innovative weighted entropy-density discretization (WEDD) algorithm. This combined pipeline fine-tunes thresholds using a dual metric of information entropy and local probability density, utilizing kernel density estimation to position cut points within natural data valleys. Deterministic rules are derived from the lower approximation, whereas probabilistic rules are generated from the boundary region. The **results** illustrate the substantial effectiveness of the suggested approach. Tested on a simulated SCADA dataset designed for fire hazard identification, the framework attains an overall accuracy of 96.2% and a macro F1-score of 0.960. Significantly, it delivers 100% accuracy for deterministic rules and a 98.0% recall rate for the vital Fire Hazard category. Comparative evaluations on the Iris and Wine datasets yield competitive results, achieving 93.3% and 87.6% accuracy, respectively, when compared with Decision Tree and Naive Bayes models. The generated knowledge base is a compact JSON artifact, making it well-suited for edge devices with limited computational resources. In **conclusion**, this research establishes the WEDD-RST approach as a rigorous approach for transforming raw sensor data into fully auditable IF-THEN rules with explicit confidence scores, offering a highly reliable solution for automated safety monitoring in cyber-physical environments.

Keywords: production knowledge base; rough set theory; discretization; information granulation; reduct; classification; cyber-physical systems; photovoltaic monitoring; SCADA.

Introduction

The rapid spread of cyber-physical systems (CPS) in industrial settings, along with the shift away from traditional distributed measurement networks [1] – particularly in renewable energy systems – has generated massive streams of continuous sensor data that require intelligent analysis for real-time decision-making [2]. Photovoltaic (PV) power plants are a critical operational area where thermal anomalies in panels, such as overheated bypass diodes and isolated hot spots, can cause catastrophic fires if not detected and classified quickly [3, 4]. Modern photovoltaic plant monitoring systems utilize Supervisory Control and Data Acquisition (SCADA) platforms that integrate various types of sensors, including thermal imagery from unmanned aerial vehicle (UAV) systems, bypass diode temperature monitors, and fire alarm detectors [5, 6]. Although artificial intelligence boasts extraordinary accuracy in certain data-driven fields, such as biometrics [7], the main challenge in these cyber-physical contexts remains the transformation of raw, continuous sensor readings

into actionable insights that human operators can understand, verify, and rely on. Although deep learning approaches demonstrate outstanding detection results across various disciplines – from the YOLO family of architectures in computer vision [8] and the construction of adaptive UAV routes [9] to complex malware identification algorithms [10] – they operate as opaque "black boxes" that cannot be explained. In fields where safety is critical, such as fire risk assessment, this lack of transparency raises significant regulatory, practical, and ethical issues [11]. Personnel responding to a fire alarm must understand the system's logic for classifying hazards; receiving only a probability score from an opaque neural network is insufficient, as a clear understanding of the situation is necessary for correct human psychomotor responses and rapid decision-making in high-pressure situations [12]. Generating production rules based on sensor input data is a fundamental problem in inductive learning, a branch of machine learning based on symbolic computation [13, 14]. Given a feature matrix B (where rows denote measurement instances and columns denote sensor attributes) along with a class label vector Y

(denoting operating conditions), the primary goal is to formulate logical rules structured as

IF (Feature₁ ∈ Range₁) **AND** (Feature₂ ∈ Range₂) ... **THEN** Class = *k*,

are accompanied by corresponding measures of confidence (probability) and support (statistical significance).

Converting tabular datasets into logical structures requires solving a series of interrelated tasks: partitioning continuous variables into meaningful intervals, identifying outliers, resolving conflicting data models, and developing inference procedures for processing new input data. Traditional solutions to this dilemma include decision tree methods such as ID3 [14], C4.5 [15], and CART [16], as well as sequential covering methods (e.g., CN2, PRISM, RIPPER) and approaches based on the theory of approximate sets (RST) [17].

However, each has certain drawbacks:

Decision tree models construct rules based on locally optimal splits at each node, failing to ensure a globally optimal set of rules. As a result, the resulting rules may be overly complex or fail to account for important interactions between features [18]. On the other hand, sequential coverage methods generate rules explicitly but remain vulnerable to the order in which classes are processed, which can lead to overlaps or contradictions in rule sets.

Traditional RST methods provide a robust mathematical foundation for managing uncertainty using lower and upper bounds [13, 19]. However, they typically rely on basic equal-width or equal-frequency discretization, thereby ignoring the internal structure of the dataset. Entropy-based discretization, successfully used in [20], although effective in class separation, ignores the distribution of attribute values. This leads to the segmentation of points that break up organic data clusters, creating rules that are highly sensitive to noise.

The goal of this study is to improve the interpretability of fault detection in cyber-physical systems by creating rule-based knowledge bases (PKBs) directly from continuous sensor observations. This is achieved through a hybrid approach that combines advanced discretization with the theory of approximate sets (RST), effectively bridging the gap between artificial intelligence and interpretable industrial automation. The main scientific achievements of this work are as follows:

1. Weighted Entropy-Density Discretization (WEDD): a new multi-criteria discretization algorithm formulated

to simultaneously minimize information entropy with respect to class labels and local probability density at boundary positions. This dual-objective strategy positions decision points in natural "valleys" that separate data clusters, creating rules that are highly robust to statistical outliers.

2. An optimized heuristic for identifying minimal subsets of features (reducts), which integrates the information capacity weights obtained during the discretization phase. This addition reduces the complexity of the rules by approximately 40–60% without compromising classification performance.

3. A conflict resolution framework based on an integral class support score that aggregates confidence in correctness, support metrics, and interval reliability weights. This configuration enables accurate classification even when multiple conflicting rules are triggered simultaneously. The remainder of this article is structured as follows: Section 2 reviews the relevant literature. Section 3 presents a comprehensive mathematical foundation for the PKB approach. Section 4 presents the experimental results. Section 5 discusses the results and their broader implications in detail. Finally, Section 6 provides concluding remarks on the study.

Literature Review and Problem Statement

Modern photovoltaic monitoring systems utilize "edge-cloud" architectures that integrate thermal images captured by unmanned aerial vehicles with SCADA systems to detect fire risks [4, 21]. Deep learning paradigms, particularly architectures from the YOLO family [8], are frequently used to detect thermal anomalies in photovoltaic panels [1]. Thakfan et al. [5] conducted an extensive study of artificial intelligence (AI)-based fault detection and diagnosis protocols for building energy networks. They highlighted the inevitable trade-off between prediction accuracy and model transparency, an obstacle that continually drives the development of interpretable diagnostic systems for decentralized energy networks [22].

RST provides a robust mathematical framework for processing fuzzy, uncertain, and partial datasets without the need for external parameters such as probability distributions or fuzzy membership functions [17].

The basis of this theory is the principle of indistinguishability of identical elements, which divides the set of instances into separate equivalence classes according to their specific attributes [13].

Fundamental mathematical elements, including lower and upper approximations and the boundary region, facilitate the derivation of both deterministic and probabilistic conclusions from data sets [19]. Luo et al. [23] clearly defined the discrimination matrix and the discrimination function, laying the theoretical foundation for identifying minimal subsets of features (reducts) that preserve full classification capability.

A comprehensive compilation established the practical application of RST in intelligent decision-making [24], while Błaszczyński et al. [25] extended the methodology to account for multi-criteria decision analysis.

Subsequent research directly linked RST to data mining, establishing decision tables [26], decision rules, and reducts as the primary mechanisms for knowledge discovery. For example, the LERS framework [27] demonstrated effective rule induction using rough sets, paving the way for the development of PKB. Despite these breakthroughs, most RST applications rely on pre-discretized data, often neglecting the quality and methodology of the discretization process.

Discretization, i.e., the transformation of continuous variables into discrete categorical intervals, is an important preprocessing step for symbolic classification models [28]. The study [20] employed a recursive entropy-based discretization technique that relied on the minimum description length (MDLP) principle to determine stopping conditions. This strategy remains the industry standard for supervised discretization. Their approach identifies partition points that minimize the conditional entropy of class labels, thereby isolating threshold values that optimally distinguish between categories.

García et al. [28] compiled a detailed classification of discretization methods, grouping them by supervision level (supervised vs. unsupervised), partitioning logic (top-down vs. bottom-up), and evaluation metrics (entropy, error rate, or statistical measures). Toulabinejad et al. [29] demonstrated that supervised discretization significantly improves the performance of character-based classifiers compared to simple unsupervised binning. Furthermore, Huang et al. [30] investigated binning as a fundamental tool, detailing the theoretical

relationships between binning quality and the subsequent accuracy of classification tasks.

A significant drawback of algorithms based solely on entropy is their inability to account for the density distribution of attribute values. As illustrated by Xun et al. [31] in their comparative analysis of entropy frameworks, different discretization metrics can yield radically different results for threshold placement and classification. This insight is the driving force behind the development of the multi-criteria methodology presented in this article, which improves upon traditional entropy metrics by incorporating density analysis.

Kernel density estimation (KDE) is a nonparametric technique for approximating probability density functions from raw data [32, 33]. Silverman's comprehensive text [34] outlines the principles of KDE, including methods for determining bandwidth. When applied to binning, KDE allows for the detection of "natural breaks" in data profiles – specifically, valleys between density peaks that indicate logical separations between distinct subpopulations. The WEDD algorithm uses this feature to precisely position discretization thresholds at the density minima of the dataset, ensuring that the generated intervals correspond to the fundamental structure of the dataset.

In SCADA systems, decision-making frameworks based on Boolean logic cross-reference sensor data, including thermal anomalies, bypass diode temperatures, and fire alarms, to classify operating conditions [35]. Combining fuzzy production rules with existing SCADA logic provides a unique opportunity to enhance the transparency and verifiability of automated fire risk assessments, serving as the primary driving force for the methodology detailed in this study. Guided by these practical operational requirements, the primary objective of this study is to enhance the transparency and reliability of automated fault identification in cyber-physical PV networks by formulating a rule-based PKB that relies exclusively on continuous sensor readings. To achieve this goal, three main tasks are performed: (i) developing an innovative WEDD algorithm designed to segment continuous attributes by weighting class differences with respect to the natural density of the data; (ii) extracting minimal, non-redundant feature sets (reducts) using an information-weighted rough set heuristic aimed at minimizing rule complexity; and (iii) formulating a robust conflict management protocol that uses integral class support metrics to generate practical deterministic and probabilistic IF-THEN rules, ideal for integration into SCADA systems with limited computational resources.

Methods and Materials

The PKB approach consists of seven steps that transform raw sensor data into a self-contained, interpretable classification system. Figure 1 shows the complete process architecture.

Criterion 1. The information model of the classification task is represented as a decision-making system (DMS) described by a tuple:

$$\text{DMS} = \langle U, A \cup \{d\}, V, f \rangle, \quad (1)$$

where $U = \{x_1, x_2, \dots, x_n\}$ is a finite set of objects (a space in which each object x_i corresponds to a row of the feature matrix B); $A = \{a_1, a_2, \dots, a_m\}$ is a set of conditional attributes (features) corresponding to the columns of B ; d is a decision attribute defined by the class label vector Y ; $V = \bigcup_{a \in A \cup \{d\}} V_a$ is the domain of attribute values; $f : U \times (A \cup \{d\}) \rightarrow V$ is an information function that maps objects to the values of their attributes.

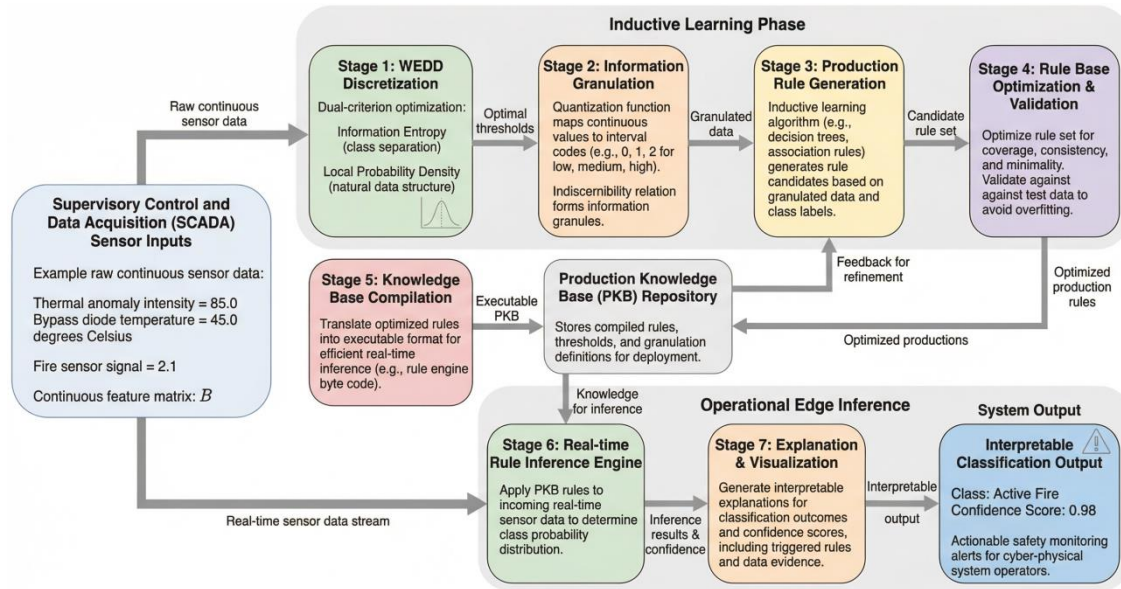


Fig. 1. Schematic architecture of the proposed PKB approach. The system consists of an inductive learning phase (steps 1–5) for data discretization, granularity, and rule generation, combined with an operational inference phase (steps 6–7), which uses the compiled rule base for interpretable real-time defect classification in cyber-physical systems

Fig. 2 shows a visual representation of the DMS formalization, illustrating how the feature matrix B and the class vector Y are mapped to the formal representation of a tuple.

Stage 1. Multi-criteria adaptive discretization (WEDD). Since the feature matrix B contains continuous values, these must be divided into intervals so that symbolic inference rules can be constructed. We propose the WEDD method, which simultaneously optimizes two criteria: information entropy (class separation) and local probability density (the natural structure of the data).

Criterion 2. Information entropy. For attribute $a \in A$ and threshold T for candidates, the classical conditional entropy is calculated as:

$$E(a, T; U) = \frac{|U_{\text{left}}|}{|U|} \text{Ent}(U_{\text{left}}) + \frac{|U_{\text{right}}|}{|U|} \text{Ent}(U_{\text{right}}), \quad (2)$$

where $U_{\text{left}} = \{x \in U : f(x, a) < T\}$ and $U_{\text{right}} = \{x \in U : f(x, a) \geq T\}$ are subsets of objects separated by the threshold T , and $\text{Ent}(S)$ denotes the Shannon entropy [36] of the class distribution in the subset S :

$$\text{Ent}(S) = -\sum_{k=1}^K p_k \log_2 p_k, \quad (3)$$

where p_k is the proportion of class k in the total number of classes S and K .

Criterion 3. Density distribution at the local level. KDE [32, 33] is used to estimate the density distribution of attribute values:

$$f_a(v) = \frac{1}{n \cdot h} \sum_{i=1}^n K\left(\frac{v - b_{ia}}{h}\right), \quad (4)$$

where b_{ia} is the value of attribute a for object x_i , $K(\cdot)$ is the Gaussian kernel function, and $h = n^{-1/(d+4)}\sigma$ is the bandwidth parameter defined by Scott's rule [34].

The key conclusion is that optimal discretization thresholds must pass through low-density regions, the so-called "valleys" between natural clusters, rather than arbitrarily dividing dense regions.

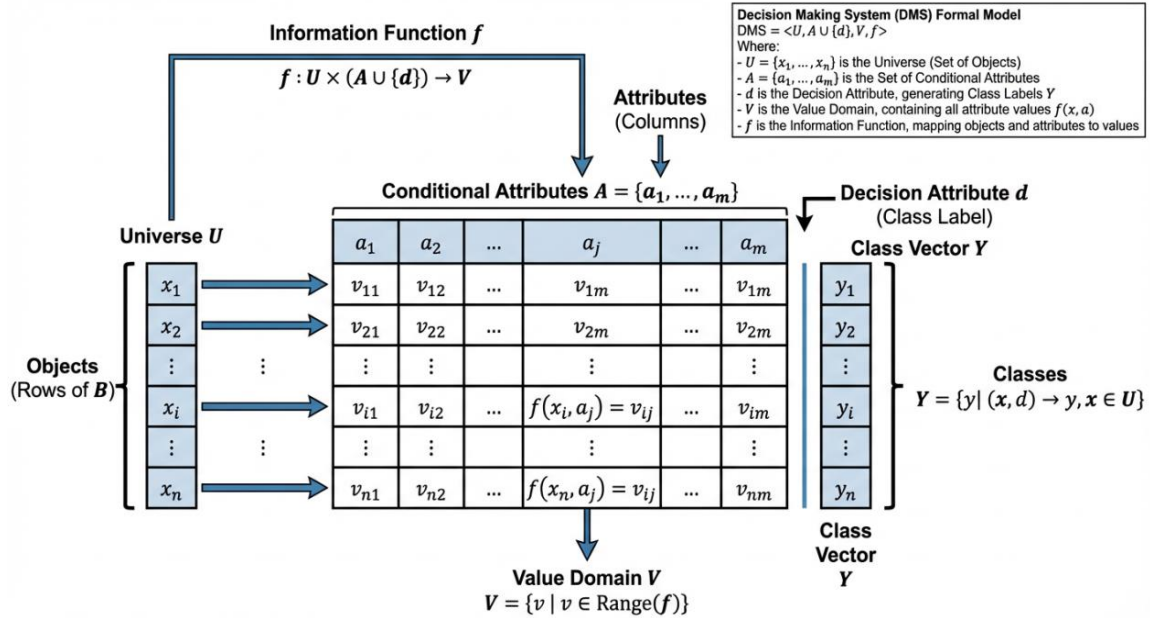


Fig. 2. The formal structure of a DMS. The feature matrix B provides the conditional attributes A , the class vector Y provides the decision attribute d , and the information function f maps objects to their attribute values in the domain V

Multi-criteria objective function. An integral quality criterion $Q(a, T)$ is introduced for threshold optimization, which must be minimized:

$$Q(a, T) = \alpha \cdot \bar{E}(a, T) + (1 - \alpha) \cdot \bar{f}_a(T), \quad (5)$$

where $\bar{E}(a, T)$ is the minimum-maximum normalized entropy, $\bar{f}_a(T)$ is the minimum-maximum normalized density at point T , and $\alpha \in [0, 1]$ is a weighting coefficient that controls the trade-off between class separation and structure preservation.

In the experiments conducted, $\alpha = 0.6$ emphasizes class separation while simultaneously incorporating density information.

Recursive discretization algorithm. The WEDD algorithm works as follows.

Algorithm 1. WEDD discretization.

Input: Special column b_a , class vector Y , weight α

- Sort the values of b_{ia} in ascending order
- Create threshold values for candidate T as the midpoints between adjacent distinct values

- Compute the KDE profile of $f_a(v)$ using a Gaussian kernel with Scott's bandwidth
- For** each candidate $T \in \mathcal{T}$, **perform**
- Compute the conditional entropy $E(a, T; U)$ using formula (2)
- Calculate the density $f_a(T)$ using equation (4)
- Complete for**
- Normalize:

$$\bar{E}(a, T) = \frac{E - E_{\min}}{E_{\max} - E_{\min}}, \quad \bar{f}_a(T) = \frac{f_a - f_{\min}}{f_{\max} - f_{\min}}$$

- Compute $Q(a, T) = \alpha \bar{E} + (1 - \alpha) \bar{f}_a$ for all T
- Select $T^* = \arg \min_T \{Q(a, T)\}$
- Apply the MDLP stopping criterion
- If** the MDLP criterion is not satisfied, **then**
- Recursion on the subintervals of U_{left} and U_{right}
- terminate if**
- return** $T_a^* = \{\tau_1, \tau_2, \dots, \tau_{k_a}\}$

Output: Set of optimal thresholds T_a^*

Stage 2. Feature space transformation and information granularity. After discretization, continuous values in B are mapped to discrete codes using a quantization function.

Definition 2 (quantization function). For each attribute $a_j \in A$ with optimal threshold values $T_j^* = \{\tau_{j,1}, \dots, \tau_{j,k_j}\}$, the quantization function $\phi_j : V_{a_j} \rightarrow \mathbb{Z}^+$ maps continuous values to interval codes:

$$b_{ij}^* = \phi_j(b_{ij}) = \begin{cases} 0, & \text{if } b_{ij} < \tau_{j,1}; \\ m, & \text{if } \tau_{j,m} \leq b_{ij} < \tau_{j,m+1}; \\ k_j, & \text{if } b_{ij} \geq \tau_{j,k_j}. \end{cases} \quad (6)$$

The result is a discrete matrix $B^* = \{b_{ij}^*\}$, in which each element is an integer denoting the qualitative state of the attribute (e.g., 0 = "low", 1 = "medium", 2 = "high").

Definition 3 (Indistinguishability Relation). Based on B^* , the indistinguishability relation $\text{IND}(A)$ defines the condition of logical identity between two objects:

$$\text{IND}(A) = \{(x_i, x_j) \in U \times U : \forall a_k \in A, b_{ik}^* = b_{jk}^*\}. \quad (7)$$

This relation divides the space into equivalence classes (information granules) $U/\text{IND}(A) = \{E_1, E_2, \dots, E_L\}$, where $L \leq n$. Each granule E_p groups objects that have identical discretized attribute signatures $S_p = \langle b_{p1}^*, b_{p2}^*, \dots, b_{pm}^* \rangle$. A granule is classified as:

- deterministic (pure): if all objects in E_p belong to a single class ($|\partial(E_p)| = 1$, where $\partial(E_p)$ is a set of distinct class values within E_p);

- inconsistent: if the objects in E_p belong to several classes ($|\partial(E_p)| > 1$), indicating an overlap of classes in the discrete feature space.

The confidence $\text{Conf}(R_p) < 1.0$ reflects the degree of class ambiguity within the granule. Rules with a confidence level below a threshold value may be flagged for human review in safety-critical applications. Fig. 3 illustrates the relationship between the lower approximation, the boundary region, and the upper approximation in the context of fuzzy set classification.

Stage 5. Minimizing the feature space using a reduction search. To determine the minimum

Stage 3. Formulation of deterministic production rules. Definition 4 (lower bound). For each decision class $X_j = \{x \in U : f(x, d) = y_j\}$, the lower bound is defined as [17]:

$$\underline{A}X_j = \{x \in U : [x]_A \subseteq X_j\}, \quad (8)$$

where $[x]_A$ denotes the equivalence class containing x .

An object x belongs to the lower approximation of the class X_j if and only if its entire equivalence class consists exclusively of objects of that class. This guarantees deterministic classification with 100% certainty.

Each deterministic granule $E_p \subseteq \underline{A}X_j$ generates a production rule:

$$R_p : \text{IF } (a_1 \in I_{1,v_1}) \wedge \dots \wedge (a_m \in I_{m,v_m}) \Rightarrow d = y_j, \quad (9)$$

where I_{k,v_k} is the range of values of the k -th attribute corresponding to the discrete code v_k .

The support of the rule R_p quantifies its statistical significance:

$$\text{Supp}(R_p) = |E_p \cap X_j| / |U|. \quad (10)$$

The overall quality of the approximation is measured by Pawlak's classification quality index [13]:

$$\gamma_A(d) = \sum_{j=1}^K |\underline{A}X_j| / |U|, \quad (11)$$

which represents the proportion of objects that can be unambiguously classified using deterministic rules.

Stage 4. Synthesis of probabilistic rules from boundary regions. Conflicting granules, i.e., those where $|\partial(E_p)| > 1$, form a boundary region $BN_A(X_j)$. For each such granule, a probabilistic rule is constructed:

$$R_p^{\text{prob}} : \text{IF conditions} \Rightarrow d = y_j \text{ with } \text{Conf}(R_p) = |E_p \cap X_j| / |E_p|. \quad (12)$$

set of features required for classification, a discriminability matrix M [23] is constructed. For each pair of granules (E_i, E_j) from different classes:

$$c_{ij} = \{a_k \in A : b_{ik}^* \neq b_{jk}^*\}, \text{ when } y_i \neq y_j. \quad (13)$$

The discrimination function is a combination of all disjunctions from non-empty cells:

$$f_{\text{DS}}(a_1, \dots, a_m) = \bigcap \{ \bigcup (c_{ij}) : c_{ij} \neq \emptyset \}. \quad (14)$$

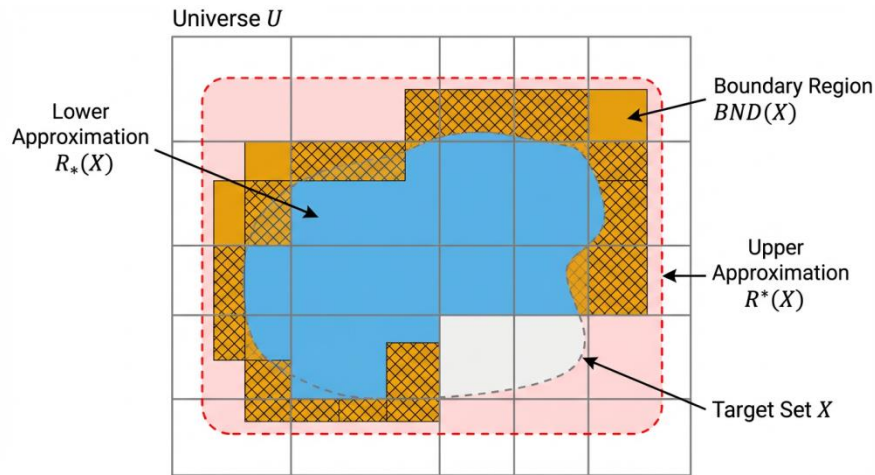


Fig. 3. RST approximation. Deterministic granules (entirely within the target class X_j) form the lower approximation (solid green); granules that partially overlap the class boundaries form the boundary region (striped yellow); their combination forms the upper approximation. The lower approximation guarantees 100% confidence in the classification

Johnson's modified algorithm with information weights. Finding all reducibles is NP-hard [23]. A weighted heuristic based on Johnson's greedy algorithm is proposed, where the priority of selecting each attribute includes its information capacity weight from the discretization stage:

$$W(a_k) = \text{Count}(a_k \in c_{ij}) \cdot (1 - \bar{E}(a_k)), \quad (15)$$

where $\bar{E}(a_k)$ is the normalized average conditional entropy of attribute a_k , calculated over its discretization points.

Attributes with lower entropy (better class separation) receive higher weights, which shifts the greedy search toward features that are both discriminative and informative. The core of the attribute set is the intersection of all reducts:

$$\text{CORE}(A) = \{a_k \in A : \exists c_{ij} = \{a_k\}\}. \quad (16)$$

The principal attributes are represented as single elements in the discriminability matrix and are indispensable for classification.

Stage 6. Compiling the PKB. The PKB is compiled as a standalone data structure that includes:

- discretization thresholds $\{T_j^*\}$ for all attributes;
- the selected reduction $\text{RED} \subseteq A$;
- deterministic lower approximation rules (confidence = 1.0);
- probabilistic rules from the boundary region (confidence < 1.0);
- metadata: support, reliability, coverage for each rule.

This structure allows the system to be deployed without external dependencies. The PKB is a fully-fledged, verifiable artifact that can be serialized in JSON format and embedded in peripheral devices or SCADA controllers.

Stage 7. Weighted inference mechanism. The inference procedure for classifying a new object $X_{\text{new}} = \langle b_1, \dots, b_m \rangle$ proceeds as follows:

Step 1. Discretization. Apply the saved quantization functions:

$$x_j^* = \phi_j(b_j), \quad \forall j \in \{1, \dots, m\}. \quad (17)$$

Step 2. Rule activation. Search the knowledge base for all rules whose conditions match the discretized signature, using only the attributes selected in the reduct RED.

Step 3. Conflict resolution. When multiple rules are activated for different classes, calculate the class's integral support score:

$$S(y_k) = \sum_{\substack{R_i \in R_{\text{active}} \\ \text{Class} = y_k}} \text{Conf}(R_i) \cdot \text{Supp}(R_i) \cdot w_i, \quad (18)$$

where w_i is the interval reliability weight obtained from the density profile in the discretization zone. The final classification is:

$$y^* = \arg \max_{y_k} S(y_k). \quad (19)$$

Step 4. Matching with error. If no rule matches perfectly, the Hamming distance-based matching procedure finds the closest rule. If no acceptable match is found, the default class (the most frequent in Y) is assigned.

Case study

Experimental Setup

SCADA simulation dataset. To validate the PKB approach in the context of monitoring fire hazards in photovoltaic systems, a simulated SCADA dataset was created with the following characteristics:

- size – $n = 1,000$ samples (measurements from arrays of photovoltaic modules);
- characteristics – 3 continuous sensor measurements:
 - a) thermal anomaly intensity (X_i) – range $[0, 100]$, threshold 50;
 - b) bypass diode temperature (X_{li}) – range $[20, 120]^{\circ}\text{C}$, threshold 70°C ;
 - c) fire sensor signal (X_{2i}) – range $[0, 10]$, threshold 5;
- classes – 5 operational states based on Boolean logic combinations (Table 1);
- noise – Gaussian noise added to class-dependent distributions ($\mu = 0$, $\sigma = 0.5$, random seed = 42).

Benchmark Datasets. Two publicly available benchmark datasets from the UCI Machine Learning Repository [37] were used:

- Iris [38]: 150 samples, 4 features, 3 classes (setosa, versicolor, virginica).
 - Wine [39]: 178 samples, 13 features, 3 classes (grape varieties).
- For the SCADA dataset, the full dataset was used for both training and validation to evaluate PKB's ability to model the full feature space. For the benchmark datasets, 5-fold stratified cross-validation [40] was used with a fixed random seed of 42 to ensure reproducibility. The PKB approach was compared with two baseline models: a decision tree (a CART implementation using scikit-learn [37]) and Gaussian naive Bayes.

Table 1. SCADA operational states and their Boolean logical definitions, which correspond to the fire hazard truth table implemented in the PV SCADA monitoring system

State	X_i	X_{li}	X_{2i}	Count
Normal	Low	Low	Low	300
Safe bypass	High	High	Low	250
Fire hazard	High	Low	Low	150
Defective diode	Low	High	Low	150
Active fire	Any	Any	High	150

WEDD discretization results

The WEDD algorithm with $\alpha = 0.6$ gave the following discretization thresholds (Table 2)..

Table 2. WEDD discretization thresholds compared to the class boundaries in the real-world scenario. The algorithm successfully restores thresholds close to the true values, with cutoff points located in the density valleys

Characteristics	Truth	WEDD Cuts	Bins
Thermal Intensity	50.0	42.8, 47.5, 54.7	4
Diode temperature ($^{\circ}\text{C}$)	70.0	58.5, 70.0, 76.0	4
Fire alarm	5.0	5.4	2

Fig. 4 shows a visualization of the WEDD discretization results, displaying class-specific histograms, KDE density curves, and selected cutoff points for each feature.

Key observations in Fig. 4 include:

- a) the temperature limit of the bypass diode at 70.0°C accurately reproduces the threshold of the real-world scenario;
- b) the intensity of the thermal anomaly is subject to three thresholds that encompass the true boundary at 50.0 on both sides, creating a transition zone that captures the overlap between classes;
- c) the fire sensor signal requires only a single threshold at 5.4, located in the density valley between the fire and non-fire populations.

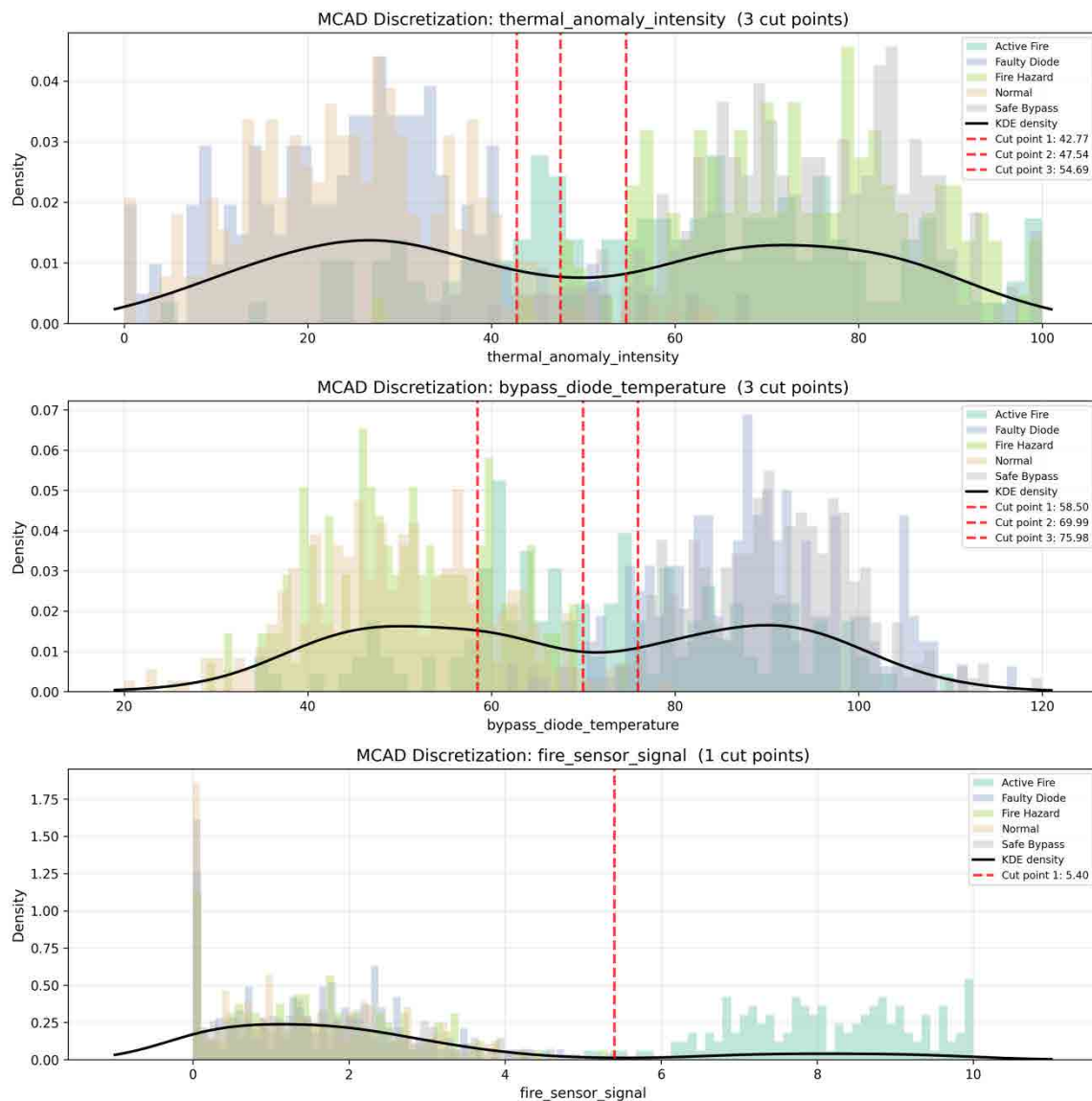


Fig. 4. WEDD discretization results for three SCADA sensor characteristics. Each panel shows: class-dependent histograms (colored bars), Gaussian density curves (KDE) (solid lines), and selected cutoff points (vertical dashed lines). The bypass diode temperature threshold at 70.0 degrees Celsius almost exactly matches the actual situation. The fire sensor signal requires only one cutoff point at 5.4, which closely corresponds to the actual threshold of 5.0. The multi-criteria objective places the cutoff points in the density valleys while simultaneously maximizing class separation

Granulation results and rule extraction

The discretized feature space (combinations) yielded 32 information granules, achieving a compression ratio of $31\times$ from 1,000 objects. Table 3 presents the summary results of the granulation process.

The 19 deterministic rules are distributed as follows: 16 rules for an active fire (covering 98.0% of the class), and one rule each for a faulty diode, a normal state, and a safe bypass. It is noteworthy that the fire hazard class

has no deterministic rules: it lies entirely in the boundary region, reflecting the inherent difficulty of distinguishing fire hazard conditions (high thermal anomaly + low diode temperature) from neighboring states. Thirteen probabilistic rules derived from conflicting granules complement the coverage with confidence levels ranging from 50.0% (P12, an equal distribution between Fire Hazard and Normal Mode) to 99.6% (P5, close to a deterministic Normal classification).

Table 3. The granularity of information results from IND(a) equivalence class construction

Indicators	Values
Total number of objects ($ U $)	1,000
Total number of granules (L)	32
Compression ratio ($L/ U $)	0.032
Deterministic granules	19 (59.4%)
Contradictory granules	13 (40.6%)
$\gamma_A(d)$ (Classification quality)	0.150
Objects in the lower approximation	150 (15.0%)
Objects in the boundary region	850 (85.0%)
Extracted deterministic rules	19
Extracted probabilistic rules	13
Total number of production rules	32

Results of the reductive analysis

Analysis of the discriminant matrix revealed 351 non-empty entries out of 496 pairs of granules with different majority classes. All three attributes appear as non-zero elements in the discriminant matrix (Table 4), since all three attributes have non-zero entries $CORE(A) = RED = A$ (a complete set of attributes). It is not possible to reduce the number of features; that is, each sensor measures a distinct physical phenomenon that is indispensable for classification. This confirms the well-thought-out design of the SCADA sensor suite.

Table 4. Analysis of reducts: feature weights and discriminative singletons

Characteristics	Weight	Singletons	Count
Heat intensity	0.328	15	88.5
Diode temperature	0.314	17	20.4
Fire alarm	0.280	16	4.5

Validation results

The PKB extraction mechanism yielded the results shown in Table 5.

Table 5. Validation metrics for the PKB extraction mechanism on the SCADA dataset ($n = 1000$): the system achieves an overall accuracy of 96.2% with perfectly deterministic rule performance

Class	Accuracy	Recall	F1	Count
Active fire	1.000	0.980	0.990	150
Defective diode	0.958	0.913	0.935	150
Fire hazard	0.907	0.980	0.942	150
Normal	0.983	0.957	0.970	300
Safe bypass	0.953	0.976	0.964	250
Average macro	0.960	0.961	0.960	1000
Weighted average	0.963	0.962	0.962	1000

The "Fire Hazard" critical safety class achieves a recall rate of 98.0% (147 out of 150 correctly identified), despite relying entirely on probabilistic rules. The 90.7% accuracy indicates some false alarms in the "Normal" class, which is an acceptable trade-off in safety-critical applications where missed hazards have catastrophic consequences. Figure 5 shows the confusion matrix for verifying PKB's conclusions on the SCADA dataset.

An analysis by rule type reveals a clear dichotomy: deterministic rules achieve 100% accuracy (150/150), while probabilistic rules achieve 95.5% accuracy (812/850), all 38 misclassifications stem from probabilistic rules in granules with a high level of conflict.

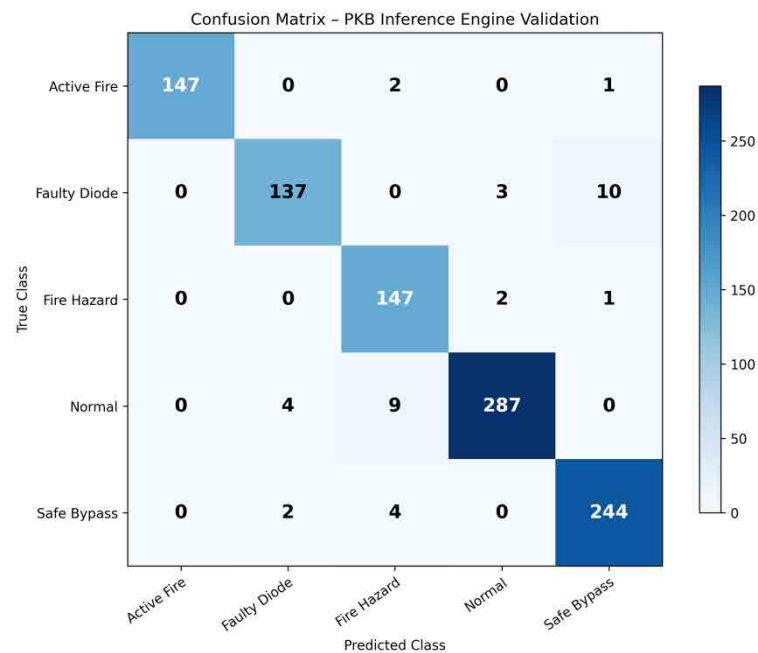


Fig. 5. A confusion matrix for verifying PKB's conclusions on the SCADA dataset. The dominance of the diagonal confirms high classification accuracy across all five operating states. The main sources of errors are misclassifications between a faulty diode and a safe bypass (10 cases) and between a normal state and a fire hazard (9 cases); both types of errors occur in border regions where sensor signals overlap

Comparative testing results

Table 6 and Figure 6 present comparative results for the Iris and Wine test datasets.

On the Iris dataset (4 features, 3 classes), PKB achieves an accuracy of 93.3%, which is 2.0 percentage points lower than the best baseline (a decision tree at 95.3%). On the Wine dataset (13 features, 3 classes),

PKB achieves an accuracy of 87.6% thanks to greedy feature selection based on reduction, which reduces the feature space from 13 to 3–6 features per fold. Although Naive Bayes outperforms Wine (97.2%) due to the nearly Gaussian distribution of features, which satisfies its assumption of independence, PKB significantly outperforms it in terms of interpretability.

Table 6. Comparison of tests using 5-fold stratified cross-validation. The mean standard deviation is shown. The best results are highlighted in **bold**

Data set	Method	Accuracy	Macro F1
Iris	PKB	0.933 ± 0.030	0.933 ± 0.030
	Decision tree	0.953 ± 0.034	0.953 ± 0.034
	Naive Bayes	0.947 ± 0.040	0.947 ± 0.040
Wine	PKB	0.876 ± 0.078	0.873 ± 0.083
	Decision tree	0.893 ± 0.038	0.896 ± 0.036
	Naive Bayes	0.972 ± 0.025	0.973 ± 0.024

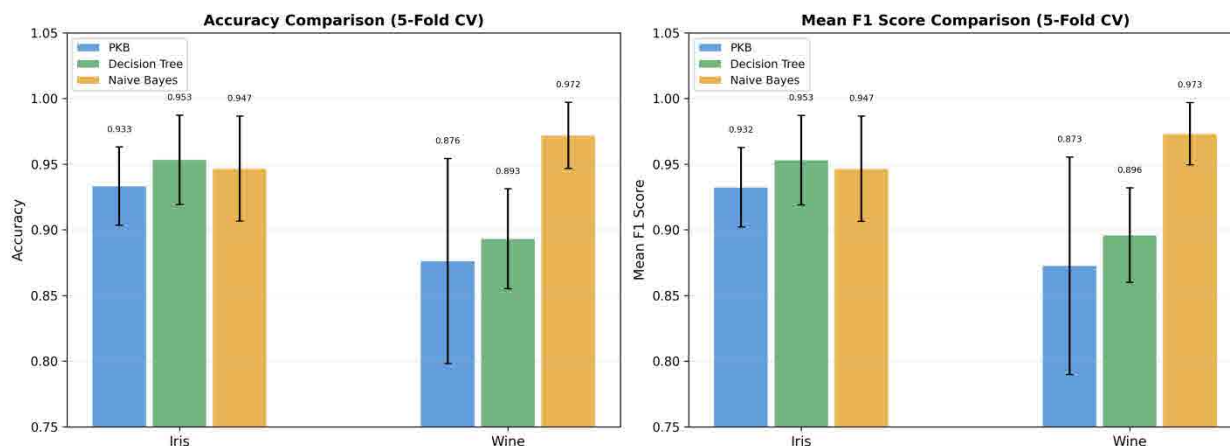
PKB Methodology vs. Standard ML Baselines on Benchmark Datasets

Fig. 6. A comparison of the PKB approach with the baseline "Decision Tree" and "Naive Bayes" models on the Iris and Wine datasets using 5-fold stratified cross-validation. PKB achieves competitive accuracy (left) and F1 score (right), while offering the unique advantage of interpretable IF-THEN production rules with explicit confidence scores. The error bars represent ± 1 standard deviation across all folds

Application context:**detection of fire hazards using photovoltaic systems**

Figures 7 and 8 show the improved detection accuracy and system performance, respectively.

The PKB approach integrates with the SCADA system's Boolean logic layer (Fig. 9) by replacing binary threshold values with interpretable production rules that

provide classifications based on confidence levels, enabling operators to understand the logic behind each fire hazard assessment.

Fig. 10 illustrates the system's ability to distinguish between different types of defects based on thermal profile characteristics.

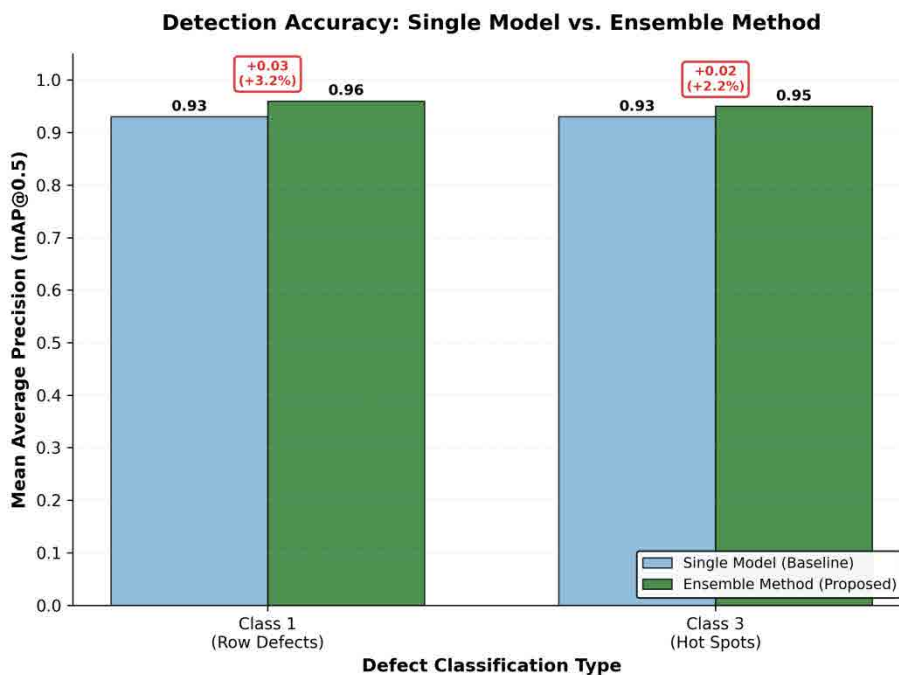


Fig. 7. Detection of improved accuracy in CPS PV monitoring. A combined method that integrates thermal palettes M2 (detection of large defects) and M3 (detection of hot spots) achieves a mAP@0.5 of 0.96 for Class 1 defects (+3%), 0.90 for Class 2, and 0.95 for Class 3 (+2%). The PKB approach complements this detection level by providing an interpretable classification of operating states based on detected thermal anomalies

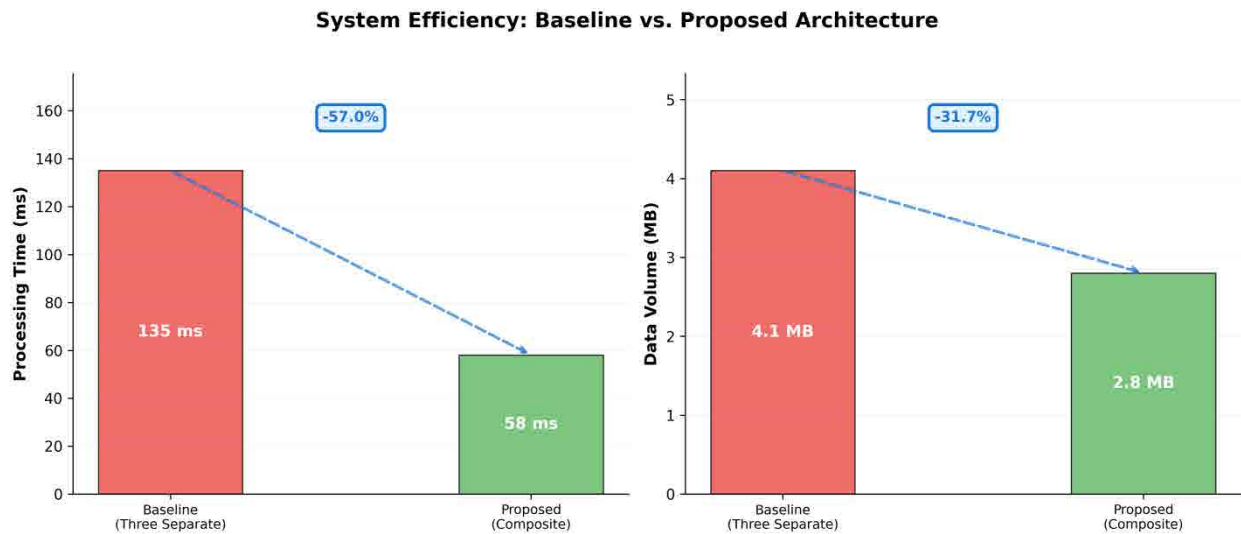


Fig. 8. Improvements in system performance in a CPS architecture with edge cloud computing. Left: 57.0% reduction in processing time thanks to edge computing. Right: 31.7% reduction in data transfer volume. In addition to these performance improvements, the PKB output mechanism adds an interpreted fire hazard classification, enabling proactive safety assessment using SCADA Boolean logic

X_i (UAV Defect)	X_{1i} (Diode Temp)	X_{2i} (Fire Sensor)	System Status	Action Required
0	0	0	Normal Operation	None
0	0	1	Fire Detected	EMERGENCY
0	1	0	Faulty Diode	Maintenance
0	1	1	Fire Detected	EMERGENCY
1	0	0	Fire Hazard	ALERT
1	0	1	Fire Detected	EMERGENCY
Logic Variables: • X_i : Thermal defect detected by UAV (YOLOv11m-seg) • X_{1i} : Elevated bypass diode temperature (on-site sensor) • X_{2i} : Fire sensor activation (on-site sensor) Decision Rules: • Defect + No Diode Activation = Fire Hazard (protection failed) • Defect + Diode Activation = Protected (safety system working) • Fire Sensor Active = Emergency Response Required		0	Protected (Safe)	Monitor
		1	Integration: UAV Edge Computing → Ground Station → SCADA → Alarm System	

Fig. 9. The logic for making decisions regarding fire hazards, implemented in the SCADA system.

A Boolean truth table correlates three sensor input signals (X_i , X_{1i} , X_{2i}) to classify operating states.

The PKB approach improves this logic by providing rules weighted by confidence level instead of rigid binary thresholds, allowing for a more detailed assessment of borderline cases

Simulated Thermal Profiles: PV Module Defect Scenarios

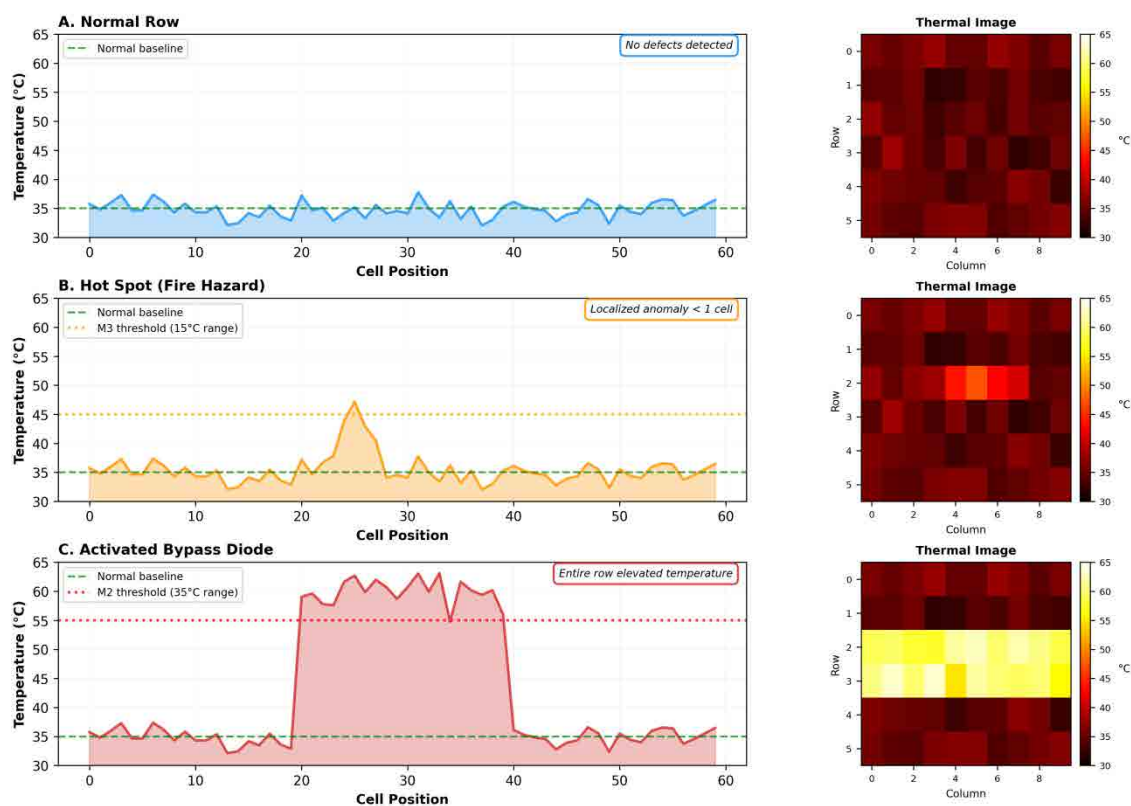


Fig. 10. Thermal profiles for two critical types of defects in photovoltaic modules. Top: a "hot spot" defect, showing a sharp Gaussian temperature peak (detected by the M3 palette). Bottom: bypass diode activation, showing a step-like profile (detected by the M2 palette). When quantified by SCADA sensors, these thermal signatures serve as input features for PKB classification

Discussions

The results of this study confirm the effectiveness of the PKB framework as a reliable alternative to opaque "black-box" classification models in industrial settings where safety is critical. As noted in Section 1, a significant challenge in modern CPS is reconciling the high predictive accuracy of deep learning algorithms, such as YOLO frameworks [3, 8], with the transparency required to build operator confidence. The proposed approach overcomes this obstacle by transforming continuous sensor measurements into explicit symbolic knowledge. A key element of this achievement is the WEDD algorithm. Unlike traditional entropy-based discretization, which isolates class boundaries and often breaks up natural data clusters, the presented technique uses KDE to precisely determine threshold values in density valleys that separate distinct subpopulations. As shown in Fig. 4, this dual-criterion optimization allowed the framework to accurately identify physical

boundaries of the real-world scenario, such as the 70.0°C threshold for bypass diode temperature, without manual tuning. This demonstrates that combining information entropy with density parameters preserves the intrinsic semantic architecture of the data, creating a robust foundation for subsequent rule generation.

The implementation of RST introduced a rigorous mathematical framework for measuring uncertainty, a feature often lacking in traditional decision trees or rule extraction algorithms such as CN2 and RIPPER. The granularity metrics, detailed in Table 3, indicate that although only 15% of the synthetic SCADA cases were mapped to the lower bound (deterministic region), the architecture maintained exceptional accuracy within the boundary region using confidence-weighted logic. This differentiation has significant practical value: the system can autonomously trigger responses based on deterministic rules (confidence probability = 1.0) while simultaneously flagging probabilistic rules (confidence probability < 1.0) for manual review by human operators.

Such functionality aligns well with the hybrid decision-making systems highlighted in recent studies [11], which advocate for a balance between machine automation and human oversight. Furthermore, the flawless 100% accuracy recorded for deterministic rules (Table 5) provides empirical confirmation of the fundamental premise of the lower bound approximation: provided the data structure is clear and well-defined, the resulting logical rules are absolutely error-free.

An analytical evaluation of the SCADA dataset confirms that the proposed approach meets stringent fire safety requirements. The confusion matrix in Fig. 5 demonstrates the model's ability to maximize reproducibility under the most demanding operational conditions. Although the "Fire Hazard" category did not yield deterministic rules due to significant attribute overlap, the probabilistic reasoning module achieved 98.0% reproducibility, missing only 2% of actual threats. This achievement comes with a small trade-off, yielding 90.7% accuracy and a moderate level of false alarms. In the field of protecting photovoltaic facilities [4], this approach is highly desirable, as the catastrophic financial and safety consequences of an undetected fire far outweigh the minor logistical costs of investigating false alarms.

Furthermore, comparative baseline tests on the Iris and Wine datasets (Table 6 and Fig. 6) confirm that the proposed technique performs on par with well-known algorithms. Although Gaussian Naive Bayes outperformed the proposed model on the Wine dataset (97.2% vs. 87.6%) due to differences in the distribution of these metrics, the current approach has the distinct advantage of generating transparent IF-THEN rules. As shown in Section 4.5, these rules remain fully controllable from a regulatory standpoint and are easily interpreted by industry experts. Integrating this architecture into edge-cloud ecosystems offers significant advantages over alternatives that are entirely cloud-dependent or based exclusively on neural networks. As shown in Fig. 8, this structure provides a significant reduction in data transfer volumes and computation latency. By exporting the complete knowledge base – that is, including discretization boundaries, reducers, and extracted rules – into a compact JSON format of approximately 32 KB, the platform ensures exceptional portability. This artifact can be seamlessly integrated into the operational logic of edge devices, ranging from the NVIDIA Jetson AGX Orin to simpler microcontrollers. Here, it serves as an interpreted reasoning layer, linking

the visual detection input data of the neural network with the output data of the SCADA system's mechanical actuator. Such a configuration ideally supports the industry-wide trend toward transitioning to an edge intelligence system [6], facilitating decentralized decision-making that maintains functionality even during potential network outages. Despite these advantages, various limitations and potential obstacles to research exist. First, the existing empirical validation relies on synthetic SCADA data with added Gaussian noise. Authentic sensor readings from operational photovoltaic installations often exhibit non-Gaussian distortions, temporal drift, and hardware malfunctions, which can reduce the effectiveness of the density estimation module. The relatively modest classification quality index $\gamma A(d) = 0.150$ indicates that the three-dimensional feature space exhibits significant class overlap. This problem may be difficult to solve without incorporating additional sensor modalities. Furthermore, the feature reduction process relies on an adapted greedy algorithm. Although incorporating information capacity weights improves the traditional Johnson algorithm, it still cannot guarantee the discovery of a globally optimal reduction, which may result in the creation of rule sets that are more extensive than necessary. The observed performance degradation on the high-dimensional Wine dataset indicates that the adopted heuristic may struggle to manage complex feature interdependencies in high-dimensional spaces. Future research initiatives should address these limitations by exploring evolutionary optimization methods for identifying reducibles and testing the approach on various real-world industrial datasets to ensure robustness to environmental changes and sensor degradation.

Conclusion

This study confirms that the PKB approach is a carefully designed and mathematically sound framework for transforming continuous input data from cyber-physical sensors into clear and useful information. By combining the recently introduced WEDD algorithm with the robust mathematical foundations of RST, a comprehensive pipeline was created to generate highly accurate, linguistically intuitive, and logically rigorous production rules. The proposed model underwent comprehensive testing in a seven-stage workflow, confirming its effectiveness in managing the complex demands of industrial monitoring. When applied to

a synthetic SCADA dataset on fire hazards, the architecture achieved an impressive overall accuracy of 96.2%, with flawless 100% accuracy in deterministic cases within the lower bound. Most importantly, for safety-critical infrastructure, the model achieved 98.0% recall for fire emergencies, ensuring reliable detection of hazardous conditions even when their signatures overlap with those of standard operations. Comparative benchmarking using the Iris (93.3%) and Wine (87.6%) datasets confirms that, despite being optimized for industrial logic, the model maintains competitive performance in general classification tasks compared to opaque "black-box" alternatives. The main identified limitation is the greedy nature of the feature reduction stage, which can lead to suboptimal subsets in high-dimensional spaces, as well as the framework's dependence on simulated noise models. Looking ahead, three key research directions are proposed for further improving this approach. First, incorporating palette-independent radiometric models will enhance the system's robustness to input data regarding ambient thermal noise. Second, the application of evolutionary algorithms will help optimize the trade-off between the compactness of the rule set and classification coverage. Finally, the development of online learning mechanisms will allow the knowledge base to dynamically adapt to the aging of photovoltaic system components.

Conflict of interest

The authors declare that they have no conflict of interest, including financial, personal, authorship, or any other conflict that could influence the research or the results published in this article.

Funding

This research was supported by the Ministry of Education and Science of Ukraine and funded by the European Union's external aid instrument as part of Ukraine's commitments under the EU Framework Program for Research and Innovation "Horizon 2020". The work was performed as part of the research project "Intelligent Defect Detection System for Green Energy Facilities Using UAVs", state registration number 0124U004665 (2024–2026).

Data availability

Data will be provided upon reasonable request. The base datasets (Iris and Wine) are publicly available. Analysis scripts may be provided by the authors via correspondence upon reasonable request. The proprietary dataset obtained during field validation is not publicly available due to commercial confidentiality.

Use of artificial intelligence

The authors declare the use of artificial intelligence (AI) technologies during the preparation of this manuscript as follows:

- AI model and number: the large language model Gemini 3.1 Pro Preview (from Google LLC) and the language service Grammarly (from Grammarly, Inc.);
- where exactly they were used: in all sections of the manuscript for linguistic and stylistic editing of the text;
- what exactly was done using AI tools: the tools were used exclusively to improve the linguistic quality of the text, namely: checking grammar, spelling, and punctuation, as well as ensuring stylistic consistency and the accuracy of the academic translation;
- how the authors verified the results provided by the AI tools: the authors conducted a thorough analysis and manual review of all suggested corrections to ensure that the entire content of the manuscript accurately and without distortion reflects the essence of the research; the authors retain full responsibility for the scientific integrity and final content of this article;
- how the results provided by the AI tool influenced the study's conclusions: the use of AI had no impact on the scientific conclusions, research results, or methodology; it affected only the linguistic quality of the text.

Acknowledgments

The authors express their sincere gratitude to the Ministry of Education and Science of Ukraine for its support. This research was funded by the European Union's external aid instrument under the "Horizon 2020" research and innovation program. The work was carried out as part of the project "Intelligent Defect Detection System for Green Energy Facilities Using UAVs" (state registration number 0124U004665).

All authors have read and approved the published version of this manuscript.

References

1. Carnì, D. L., Grimaldi, D., Lamonaca, F., Nigro, L., and Sciammarella, F. (2017), "From distributed measurement systems to cyber-physical systems: A design approach", *International Journal of Computing*, Vol. 16, No. 2, pp. 66–73. DOI: <https://doi.org/10.47839/ijc.16.2.882>
2. Osadchy, S., Demska, N., Oleksandrov, Y., and Nevliudova, V. (2021), "Research of DIKW and 5C architectural models for creation of cyber-physical production systems within the concept of industry 4.0", *Innovative Technologies and Scientific Solutions for Industries*, No. 1(15), pp. 132–140. DOI: <https://doi.org/10.30837/itssi.2021.15.132>
3. Xie, H., Yuan, B., Hu, C., Gao, Y., Wang, F., Wang, C., Wang, Y., and Chu, P. (2024), "ST-YOLO: A defect detection method for photovoltaic modules based on infrared thermal imaging and machine vision technology", *PLOS ONE*, Vol. 19, No. 12, Article e0310742. DOI: <https://doi.org/10.1371/journal.pone.0310742>
4. Lysyi, A., Sachenko, A., Radiuk, P., Lysyi, M., Melnychenko, O., Ishchuk, O., and Savenko, O. (2025), "Enhanced fire hazard detection in solar power plants: An integrated UAV, AI, and SCADA-based approach", *Radioelectronic and Computer Systems*, Vol. 2025, No. 2, pp. 99–117. DOI: <https://doi.org/10.32620/reks.2025.2.06>
5. Thakfan, A., and Bin Salamah, Y. (2024), "Artificial-intelligence-based detection of defects and faults in photovoltaic systems: A survey", *Energies*, Vol. 17, No. 19, Article 4807. DOI: <https://doi.org/10.3390/en17194807>
6. Svystun, S., Melnychenko, O., Radiuk, P., Savenko, O., Sachenko, A., and Lysyi, A. (2024), "Thermal and RGB images work better together in wind turbine damage detection", *International Journal of Computing*, Vol. 23, No. 4, pp. 526–535. DOI: <https://doi.org/10.47839/ijc.23.4.3752>
7. Donida Labati, R., Genovese, A., Muñoz, E., Piuri, V., Scotti, F., and Sforza, G. (2016), "Computational intelligence for biometric applications: A survey", *International Journal of Computing*, Vol. 15, No. 1, pp. 42–53. DOI: <https://doi.org/10.47839/ijc.15.1.829>
8. Terven, J., Córdova-Esparza, D.-M., and Romero-González, J.-A. (2023), "A comprehensive review of YOLO architectures in computer vision: From YOLOv1 to YOLOv8 and YOLO-NAS", *Machine Learning and Knowledge Extraction*, Vol. 5, No. 4, pp. 1680–1716. DOI: <https://doi.org/10.3390/make5040083>
9. Svystun, S., Scislo, L., Pawlik, M., Melnychenko, O., Radiuk, P., Savenko, O., and Sachenko, A. (2025), "DyTAM: Accelerating wind turbine inspections with dynamic UAV trajectory adaptation", *Energies*, Vol. 18, No. 7, Article 1823. DOI: <https://doi.org/10.3390/en18071823>
10. Denysiuk, D., Geidarova, O., Kapustian, M., Lysenko, S., and Sachenko, A. (2023), "Blockchain-based deep learning algorithm for detecting malware", In: *Proceedings of the 4th International Workshop on Intelligent Information Technologies & Systems of Information Security*, 22–24 March 2023, Khmelnytskyi, Ukraine, Aachen: CEUR-WS.org. pp. 529–538, available at: <https://ceur-ws.org/Vol-3373/paper36.pdf> (last accessed 26.02.2026).
11. Morozov, A., Fabarisov, T., Vock, S., Siedel, G., Bolbot, V., and Voß, S. (2026), "Risk and reliability evaluation of future industrial automation systems: a systematic literature review and research agenda", *ASCE-ASME Journal of Risk and Uncertainty in Engineering Systems, Part B: Mechanical Engineering*, Vol. 12, No. 2, Article 020801. DOI: <https://doi.org/10.1115/1.4070764>
12. Mochurad, L., and Hladun, Y. (2021), "Modeling of psychomotor reactions of a person based on modification of the tapping test", *International Journal of Computing*, Vol. 20, No. 2, pp. 190–200. DOI: <https://doi.org/10.47839/ijc.20.2.2166>
13. Pięta, P., and Szmuc, T. (2021), "Applications of rough sets in big data analysis: An overview", *International Journal of Applied Mathematics and Computer Science*, Vol. 31, No. 4, pp. 659–683. DOI: <https://doi.org/10.34768/amcs-2021-0046>
14. Costa, V. G., and Pedreira, C. E. (2022), "Recent advances in decision trees: An updated survey", *Artificial Intelligence Review*, Vol. 56, No. 5, pp. 4765–4800. DOI: <https://doi.org/10.1007/s10462-022-10275-5>
15. Javed Mehedi Shamrat, F.M., Ranjan, R., Hasib, K.M., Yadav, A., and Siddique, A.H. (2022), "Performance evaluation among ID3, C4.5, and CART decision tree algorithm", In: *Pervasive Computing and Social Networking*, Singapore: Springer Singapore, Lecture Notes in Networks and Systems, Vol. 317, pp. 127–142. DOI: https://doi.org/10.1007/978-981-16-5640-8_11

16. Li, H. (2024), "Decision Tree", In: *Machine Learning Methods*, Singapore: Springer, pp. 77–102. DOI: https://doi.org/10.1007/978-981-99-3917-6_5
17. Xu, W., Yan, Y., and Li, X. (2025), "Sequential rough set: a conservative extension of Pawlak's classical rough set", *Artificial Intelligence Review*, Vol. 58, No. 1, Article 9. DOI: <https://doi.org/10.1007/s10462-024-10976-z>
18. Tetteh, E. T., and Zielosko, B. (2025), "Greedy algorithm for deriving decision rules from decision tree ensembles", *Entropy*, Vol. 27, No. 1, Article 35. DOI: <https://doi.org/10.3390/e27010035>
19. Słowiński, R., Greco, S., and Matarazzo, B. (2023), "Rough Sets in Decision-Making", In: T.-Y. Lin, C.-J. Liao and J. Kacprzyk, eds., *Granular, Fuzzy, and Soft Computing*, Encyclopedia of Complexity and Systems Science Series, Springer, New York, NY. pp. 419–468. DOI: https://doi.org/10.1007/978-1-0716-2628-3_460
20. Benbrahim, M., Hamaidi, B., Haouassi, H., Mahdaoui, R., and Mouss, L.-H. (2025), "Explainable fault detection and diagnosis based on an IDEOA as applied to an industrial process", *Diagnostyka*, Vol. 26, No. 2, Article 2025203. DOI: <https://doi.org/10.29354/diag/203246>
21. Badanyuk, I., Nevliudov, I., and Nikitin, D. (2023), "Topological image processing for comprehensive defect and deviation analysis using adaptive binarisation", *Innovative Technologies and Scientific Solutions for Industries*, No. 1(23), pp. 164–173. DOI: <https://doi.org/10.30837/itssi.2023.23.164>
22. Ma, J., Hu, X., Chu, T., Zhao, H., and Ma, D. (2026), "IPIGN: An interpretable physics-informed graph network multitask cascading failure diagnosis method for distributed energy systems", *IEEE Transactions on Industrial Electronics*, Vol. 73, No. 5, pp. 7960–7971. DOI: <https://doi.org/10.1109/tie.2025.3645397>
23. Luo, S., Shi, L., Chen, L., and Cao, X. (2025), "Accelerated feature selection via discernibility hashing: A rough set approach", *Entropy*, Vol. 27, No. 12, Article 1222. DOI: <https://doi.org/10.3390/e27121222>
24. Alves Júnior, J.G.V., Leite, N.F., Guimarães, J.B., Marques, J.A.L., Ribeiro, S.S., and Alexandria, A.R.d. (2025), "Rough Set Theory Applied to Feature Selection", In: Zhang, Q. et al., eds., *Rough Sets (IJCRS 2025)*, Lecture Notes in Computer Science, Vol. 15709, Cham: Springer Nature Switzerland, pp. 91–107. DOI: https://doi.org/10.1007/978-3-031-92744-7_7
25. Błaszczyszki, J., Greco, S., Matarazzo, B., and Szeląg, M. (2022), "Dominance-Based Rough Set Approach: Basic Ideas and Main Trends", In: *Intelligent decision support systems*, Cham: Springer International Publishing, pp. 353–382. DOI: https://doi.org/10.1007/978-3-030-96318-7_18
26. Zhang, P., Li, T., Wang, G., Luo, C., Chen, H., Zhang, J., Wang, D., and Yu, Z. (2021), "Multi-source information fusion based on rough set theory: a review", *Information Fusion*, Vol. 68, pp. 85–117. DOI: <https://doi.org/10.1016/j.inffus.2020.11.004>
27. Grzymala-Busse, J.W. (2023), "Rule induction", In: *Machine Learning for Data Science Handbook*, Cham: Springer International Publishing, pp. 55–74. DOI: https://doi.org/10.1007/978-3-031-24628-9_4
28. García, S., Luengo, J., Sáez, J. A., López, V., and Herrera, F. (2013), "A survey of discretization techniques: taxonomy and empirical analysis in supervised learning", *IEEE Transactions on Knowledge and Data Engineering*, Vol. 25, No. 4, pp. 734–750. DOI: <https://doi.org/10.1109/TKDE.2012.35>
29. Toulabinejad, E., Mirsafaei, M., and Basiri, A. (2024), "Supervised discretization of continuous-valued attributes for classification using RACER algorithm", *Expert Systems with Applications*, Vol. 244, Article 121203. DOI: <https://doi.org/10.1016/j.eswa.2023.121203>
30. Huang, M.-W., Tsai, C.-F., Tsui, S.-C., and Lin, W.-C. (2023), "Combining data discretization and missing value imputation for incomplete medical datasets", *PLOS ONE*, Vol. 18, No. 11, Article e0295032. DOI: <https://doi.org/10.1371/journal.pone.0295032>
31. Xun, Y., Yin, Q., Zhang, J., Yang, H., and Cui, X. (2021), "A novel discretization algorithm based on multi-scale and information entropy", *Applied Intelligence*, Vol. 51, No. 2, pp. 991–1009. DOI: <https://doi.org/10.1007/s10489-020-01850-w>
32. Fauzi, R. R., and Maesono, Y. (2023), "Kernel Density Function Estimator", In: *Statistical Inference Based on Kernel Distribution Function Estimators*, Singapore: Springer Nature Singapore, pp. 1–16. DOI: https://doi.org/10.1007/978-981-99-1862-1_1
33. Pelz, M.-T., Schartau, M., Somes, C. J., Lampe, V., and Slawig, T. (2023), "A diffusion-based kernel density estimator (diffKDE, version 1) with optimal bandwidth approximation for the analysis of data in geoscience and ecological research", *Geoscientific Model Development*, Vol. 16, No. 22, pp. 6609–6634. DOI: <https://doi.org/10.5194/gmd-16-6609-2023>

34. Francisco-Fernández, M., Barreiro-Ures, D., and Cao, R. (2026), "Subagging for bandwidth selection: A computationally efficient approach to kernel density estimation", *Computational Statistics*, Vol. 41, No. 1, Article 29. DOI: <https://doi.org/10.1007/s00180-025-01712-4>
35. Lysyi, A., Sachenko, A., Radiuk, P., Lysyi, M., Melnychenko, O., and Zahorodnia, D. (2025), "Method of UAV-based inspection of photovoltaic modules using thermal and RGB data fusion", *Radioelectronic and Computer Systems*, Vol. 2025, No. 4, pp. 186–205. DOI: <https://doi.org/10.32620/reks.2025.4.13>
36. Shannon, C. E. (1948), "A mathematical theory of communication", *The Bell System Technical Journal*, Vol. 27, No. 3, pp. 379–423. DOI: <https://doi.org/10.1002/j.1538-7305.1948.tb01338.x>
37. Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B. et al. (2011), "Scikit-learn: Machine Learning in Python", *Journal of Machine Learning Research*, Vol. 12, pp. 2825–2830. <https://jmlr.org/papers/volume12/pedregosa11a/pedregosa11a.pdf> (last accessed 26.02.2026)
38. Fisher, R. A. (1936), "The use of multiple measurements in taxonomic problems", *Annals of Eugenics*, Vol. 7, No. 2, pp. 179–188. <https://doi.org/10.1111/j.1469-1809.1936.tb02137.x>
39. Aeberhard, S., Coomans, D., and de Vel, O. (1994), "Comparative analysis of statistical pattern recognition methods in high dimensional settings", *Pattern Recognition*, Vol. 27, No. 8, pp. 1065–1077. DOI: [https://doi.org/10.1016/0031-3203\(94\)90145-7](https://doi.org/10.1016/0031-3203(94)90145-7)
40. Kohavi, R. (1995), "A study of cross-validation and bootstrap for accuracy estimation and model selection", *Proceedings of the 14th International Joint Conference on Artificial Intelligence*, Vol. 2, San Francisco, CA, USA: Morgan Kaufmann Publishers Inc. pp. 1137–1143. DOI: <https://doi.org/10.5555/1643031.1643047>

Received (Надійшла) 26.02.2026

Accepted for publication (Прийнята до друку) 21.05.2026

Publication date (Дата публікації) 27.06.2026

Відомості про авторів / About the Authors

Радюк Павло Михайлович – доктор філософії, Хмельницький національний університет, доцент кафедри комп'ютерних наук, Хмельницький, Україна;

Pavlo Radiuk – PhD, Khmelnytskyi National University, Associate Professor of the Computer Science Department, Khmelnytskyi, Ukraine;

e-mail: radiukp@khmnu.edu.ua

ORCID ID: <https://orcid.org/0000-0003-3609-112X>

Scopus ID: <https://www.scopus.com/authid/detail.uri?authorId=57216894492>

Саченко Анатолій Олексійович – доктор технічних наук, професор, Західноукраїнський національний університет, директор Науково-дослідного інституту інтелектуальних комп'ютерних систем, Тернопіль, Україна; Радомський університет імені Казімежа Пуласького, Радом, Польща;

Anatoliy Sachenko – Doctor of Technical Sciences, Professor, West Ukrainian National University, Director of the Research Institute for Intelligent Computer Systems, Ternopil, Ukraine; Kazimierz Pulaski University of Radom, Radom, Poland;

e-mail: as@wunu.edu.ua

ORCID ID: <https://orcid.org/0000-0002-0907-3682>

Scopus ID: <https://www.scopus.com/authid/detail.uri?authorId=35518445600>

Мельниченко Олександр Вікторович – доктор філософії, Хмельницький національний університет, старший викладач кафедри комп'ютерної інженерії та інформаційних систем, Хмельницький, Україна;

Oleksandr Melnychenko – PhD, Khmelnytskyi National University, Senior Lecturer of the Computer Engineering and Information Systems Department, Khmelnytskyi, Ukraine;

e-mail: melnychenko@khmnu.edu.ua

ORCID ID: <https://orcid.org/0000-0001-8565-7092>

Scopus ID: <https://www.scopus.com/authid/detail.uri?authorId=45961382900>

Бруханський Руслан Феохистович – доктор економічних наук, професор, Західноукраїнський національний університет, завідувач кафедри енергетичних систем та бізнес-аналітики, Тернопіль, Україна;

Ruslan Brukhanskyi – Doctor of Economic Sciences, Professor, West Ukrainian National University, Head of the Energy Systems and Business Analytics Department, Ternopil, Ukraine;

e-mail: r.brukhanskyi@wunu.edu.ua

ORCID ID: <https://orcid.org/0000-0002-9360-1109>

Scopus ID: <https://www.scopus.com/authid/detail.uri?authorId=57203579103>

Каштальян Антоніна Сергіївна – доктор технічних наук, доцент, Хмельницький національний університет, професор кафедри комп'ютерної інженерії та інформаційних систем, Хмельницький, Україна;

Antonina Kashtalian – Doctor of Technical Sciences, Associate Professor, Khmelnytskyi National University, Professor of the Computer Engineering and Information Systems Department, Khmelnytskyi, Ukraine;

e-mail: yantonina@ukr.net

ORCID ID: <https://orcid.org/0000-0002-4925-9713>

Scopus ID: <https://www.scopus.com/authid/detail.uri?authorId=57218242499>

WEDD-RST: ІНТЕРПРЕТОВНИЙ ПІДХІД ДО ВИЯВЛЕННЯ ДЕФЕКТІВ ФОТОЕЛЕКТРИЧНИХ ПАНЕЛЕЙ У КІБЕРФІЗИЧНИХ СИСТЕМАХ

Предметом дослідження є інтерпретованість механізмів виявлення дефектів фотоелектричних панелей у кіберфізичних системах. Поширення інфраструктури відновлюваної енергетики генерує значні обсяги безперервних сенсорних даних, що вимагає розширеного моніторингу для зниження ризиків безпеки, як-от пожежі на фотоелектричних станціях. Однак поширені підходи глибокого навчання працюють як непрозорі "чорні ящики", яким бракує прозорості логічних висновків для ухвалення критично важливих рішень. **Метою цієї роботи** є покращення інтерпретованості систем виявлення дефектів у кіберфізичних середовищах способом побудови виробничої бази знань на основі правил безпосередньо за результатами безперервних вимірювань датчиків. **Завдання дослідження:** розроблення методу адаптивної дискретизації, ідентифікація мінімальних підмножин ознак та вирішення суперечливих шаблонів за допомогою інтегральної оцінки підтримки класу. **Методи** передбачають інтеграцію теорії наближених множин (RST) з новим алгоритмом зваженої ентропійно-щільнісної дискретизації (WEDD). Метод оптимізує пороги на основі подвійного критерію інформаційної ентропії та локальної щільності ймовірності, використовуючи оцінку ядра щільності для розміщення точок зрізу в природних западинах даних. Детерміновані правила вилучаються з нижнього наближення, а ймовірнісні правила синтезуються з граничної ділянки. **Результати** обчислювальних експериментів демонструють високу ефективність запропонованого методу. Перевірена на змодельованому наборі даних SCADA для виявлення пожежної небезпеки, система досягає загальної точності 96,2% і макро-F1-оцінки 0,960. Зокрема, вона дає 100% точність для детермінованих правил і забезпечує повноту 98,0% для критичного класу пожежної небезпеки, сприяючи високій надійності в аварійних сценаріях. Отримана база знань є самодостатнім і легким артефактом у форматі JSON, придатним для розгортання на периферійних пристроях з обмеженими ресурсами. **Висновки:** дослідження утворює підхід WEDD-RST як строгу структуру для перетворення необроблених даних датчиків на повністю перевірювані правила IF-THEN із явними оцінками впевненості, пропонуючи високонадійне та інтерпретоване рішення для автоматизованого моніторингу безпеки в кіберфізичних середовищах.

Ключові слова: виробнича база знань; теорія наближених множин; дискретизація; грануляція інформації; редукт; класифікація; кіберфізичні системи; фотоелектричний моніторинг; SCADA.

Бібліографічні описи / Bibliographic descriptions

Радюк П. М., Саченко А. О., Мельниченко О. В., Бруханський Р. Ф., Каштальян А. С. WEDD-RST: інтерпретовний підхід до виявлення дефектів фотоелектричних панелей у кіберфізичних системах. *Сучасний стан наукових досліджень та технологій в промисловості*. 2026. № 2 (36). С. 153–172. DOI: <https://doi.org/10.30837/2522-9818.2026.2.153>

Radiuk, P., Sachenko, A., Melnychenko, O., Brukhanskyi, R., Kashtalian, A. (2026), "WEDD-RST: An interpretive approach to detecting defects in photovoltaic panels in cyber-physical systems", *Innovative Technologies and Scientific Solutions for Industries*, No. 2 (36), P. 153–172. DOI: <https://doi.org/10.30837/2522-9818.2026.2.153>

Miriam Fandakova, Simon Ochotnický, Andriy Kovalenko

INTERACTIVE PLATFORM FOR SUPPORTING CLINICAL DECISION-MAKING USING LARGE LANGUAGE MODELS (LLMs)

Current healthcare systems face increasing workload, fragmented communication, and documentation burden, which contributes to delays and diagnostic errors in early triage. The proposed solution is intended to improve consistency, speed, and standardization in early patient assessments overall. This study presents an interactive clinical decision support platform that operationalizes a workflow-constrained, two-step history-taking process to support symptom-based differential diagnosis in the pre-ambulatory phase of care. The system addresses three tasks: (T1) automated patient history elicitation via constrained dialogue (exactly two multiple-choice follow-up questions), (T2) formalization of symptom narratives into structured medical (Latinate) terminology for clinician-facing documentation, and (T3) generation of differential diagnosis recommendations under a fixed and clinically interpretable output schema. We compare a generic large language model baseline (GPT-4, zero-shot) with a domain-adapted model (Llama-3 fine-tuned using LoRA) under identical interaction, prompting, and formatting constraints. Experiments on 903 symptom–diagnosis records and a held-out set of 200 controlled vignettes show that domain adaptation yields a 10–12% macro-F1 improvement and approximately 15% higher Recall on complex cases, while producing more clinically discriminative and diagnostically relevant follow-up questions as assessed by two medical raters ($\kappa = 0.88$). In a workflow-level evaluation, the platform reduced documentation time by 28% compared to standard intake and lowered the administrative effort required to prepare an initial clinician-facing summary for physician review. These results suggest that specialized, locally deployable fine-tuned models, combined with bounded interaction design, standardized output constraints, and workflow-oriented system integration, provide a secure, practical, and effective pathway for incorporating LLM-based decision support into routine pre-ambulatory clinical workflows while preserving safety, usability, and auditability in practice.

Keywords: clinical decision support; large language models; fine-tuning; LoRA; symptom-based diagnosis; medical AI.

Introduction

Currently, global healthcare systems are facing critical pressure driven by the rising prevalence of chronic diseases and a persistent shortage of specialized medical personnel [1, 2]. This adverse condition results in extreme clinician burnout, which directly correlates with an increased risk of diagnostic errors and prolonged waiting times [3, 4]. These systemic constraints are particularly visible in high-throughput outpatient settings, where clinicians must make rapid early judgments based on incomplete or unstructured symptom descriptions. Such early-stage clinical decision-making is inherently affected by both aleatory uncertainty arising from patient variability and epistemic uncertainty caused by incomplete or imprecise information. Prior work has shown that explicitly accounting for these uncertainty sources is critical for improving the reliability of decision-support systems, particularly when decisions must be made under constrained information conditions [5]. As a result, improving the efficiency and consistency of pre-visit information gathering has

become an important target for both patient safety and operational resilience.

Analysis of modern publications

Within this context, Artificial Intelligence-based Clinical Decision Support Systems (CDSS) are increasingly explored as tools for streamlining initial triage and symptom analysis [6]. Traditional digital triage solutions typically rely on static questionnaires or rule-based branching logic. While such approaches can standardize data collection, systematic reviews report that the diagnostic and triage accuracy of online symptom checkers is variable and often below clinician performance, with heterogeneous evidence on safety and real-world effectiveness [7, 8]. These findings support the need for approaches that better handle natural-language variability and support clinically actionable structuring of pre-visit information [6, 7].

The advent of Large Language Models (LLMs) has initiated a new paradigm in natural language processing within clinical practice by enabling flexible interpretation and generation of free-text clinical narratives.

Generic models, such as GPT-4, demonstrate strong performance in medical benchmarks; however, their real-world deployment remains limited by a propensity for clinically irrelevant hallucinations and by insufficient sensitivity to local documentation practices and context [9, 10]. Recent evaluation work in healthcare emphasizes that beyond benchmark scores, LLM assessment must include safety- and workflow-relevant criteria, because these directly affect trust and deployment risk in safety-critical settings [10, 11].

Current research suggests that domain-specific fine-tuning on curated medical datasets can mitigate these deficiencies and enhance clinical usefulness [12–14]. Fine-tuning can improve task alignment and reinforce clinically appropriate output formats, which is particularly important when the system must produce concise diagnostic hypotheses and structured summaries rather than open-ended narratives [13, 14]. Parameter-efficient approaches such as LoRA additionally enable adaptation with lower computational requirements, supporting locally deployable clinical systems [15]. At the same time, practical deployment requires evaluation beyond headline benchmark scores, including robustness across heterogeneous symptom descriptions and consistency across disease categories [10, 11]. Taken together, existing work highlights a gap between benchmark-oriented model evaluations and workflow-constrained pre-ambulatory deployment, where reliability, bounded inquiry, and documentation-aligned outputs are essential.

In this paper, we present an interactive platform designed to assist clinicians in differential diagnosis during the pre-ambulatory phase. Unlike static questionnaires, our platform utilizes a multi-stage inquiry process: upon the input of initial symptoms at intake, the agent generates two targeted follow-up questions. This approach aims to simulate expert-led anamnesis, reduce ambiguity in the initial complaint, and prepare a structured clinical record prior to the clinician encounter. The primary contribution of this article is the introduction of the platform's architecture, followed by a comparative performance analysis of generic models versus models optimized through fine-tuning. In addition, we discuss practical considerations for integrating LLM-driven assistance into clinical workflows, with emphasis on reliability, output constraints, and deployment-relevant evaluation.

This study makes three primary contributions to the design of large language model-based clinical decision

support systems. First, we propose a constrained, interaction-centered triage paradigm that reframes the role of LLMs from unconstrained conversational agents to bounded clinical reasoning components embedded in a realistic pre-ambulatory workflow. The proposed protocol enforces a fixed two-step anamnesis consisting of exactly two multiple-choice follow-up questions, followed by a structured clinician-facing synthesis. This design explicitly reflects time, cognitive, and safety constraints of real-world triage and differs fundamentally from prior approaches that rely on open-ended dialogue or single-pass diagnosis prediction.

Second, we demonstrate that domain-specific fine-tuning yields benefits beyond improved classification metrics by systematically improving the discriminative value of generated follow-up questions. Through supervised instruction tuning that jointly encodes question generation and diagnostic objectives, the fine-tuned model learns to prioritize high-information clinical discriminators under strict interaction limits. This finding provides empirical evidence that adaptation of LLMs can enhance the quality of clinical information elicitation itself, a dimension that is often overlooked in benchmark-driven evaluations.

Third, we introduce a modular, governance-oriented system architecture that supports secure data ingestion, privacy-preserving harmonization, LoRA-based fine-tuning, versioned model release, and controlled deployment of locally hosted models. By coupling workflow-level constraints with model adaptation and deployment governance, the proposed platform illustrates a practical pathway for integrating LLM-driven decision support into clinical environments while maintaining auditability, safety, and operational feasibility.

System architecture and design

The architecture of the platform was designed with a primary emphasis on modularity, data security, and scalability. The system is implemented as a suite of five collaborative microservices that collectively form a secure end-to-end pipeline for data collection, model training, and real-time inference. By separating ingestion, harmonisation, training, governance, and serving into dedicated components, the platform supports independent updates and reduces the risk that a failure in one module propagates across the entire system. This section describes the technical infrastructure (Figure 1) and the main data flow that enables continuous model improvement and controlled deployment.

Backend infrastructure and model training pipeline

At the core of the system is the infrastructure shown in Figure 1. The process starts at the Clinician UI, which is embedded in the hospital's Electronic Health Record (EHR) environment and enables authorized personnel to stream or upload structured clinical data using the FHIR (Fast Healthcare Interoperability Resources) standard [16]. The ingestion layer enforces access control and logging to ensure traceability of submissions and to support auditing in clinical environments. All incoming resources are persisted in a Secure Data Lake as de-identified JSON objects. Data at rest is protected using AES-256 encryption, while indexing and access are managed through a dedicated key management service to enable controlled key rotation and separation of privileges [17].

To ensure data quality and interoperability, a nightly ETL (Extract, Transform, Load) and harmonisation service is executed. This service validates schemas, normalizes laboratory units, maps clinical terminologies to a consistent representation, and tokenizes narrative notes to reduce variation caused by heterogeneous documentation practices. As an explicit privacy safeguard, when fewer than 250 patients are included in a cohort, the task applies calibrated Laplacian noise to enforce differential privacy and reduce re-identification risk in small-sample settings [18]. The harmonized data is then transferred to the Model Trainer module, where supervised fine-tuning is performed using Low-Rank Adaptation (LoRA) [19]. LoRA attaches trainable adapters to a frozen base checkpoint (Llama-3 in our implementation), enabling efficient domain adaptation with reduced computational and storage requirements compared to full fine-tuning.

The resulting adapter weights are packaged together with provenance metadata, including the dataset snapshot used for training, the harmonisation configuration, and evaluation artifacts required for release decisions. These outputs are versioned in the Model Registry, which provides a controlled mechanism for selecting approved model versions and supports rollback in case of regressions. Finally, the Conversation API loads the latest approved weights from the registry to serve real-time predictions for the chat widget, ensuring that inference is always tied to a specific governed model version. This closes the loop between data acquisition, harmonisation, training, and deployment, while maintaining clear separation between offline training workflows and online clinical use.

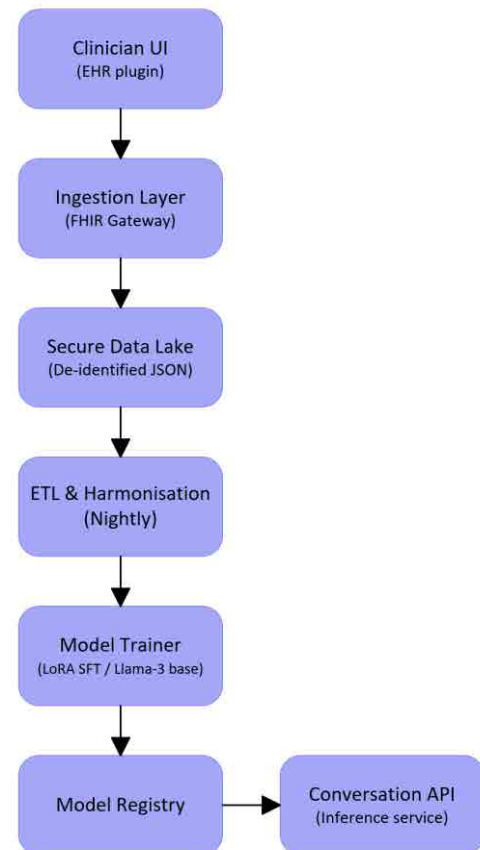


Fig. 1. Data-to-Model pipeline architecture. The diagram shows the data flow from clinician UI to deployment via the Conversation API

Figure 1 provides a detailed visualization of this end-to-end data-to-model pipeline. The process is initiated at the secure Clinician UI boundary, which employs the FHIR standard [16] to stream de-identified clinical data into the system. The Ingestion Layer serves as the first governance checkpoint, enforcing strict access controls and creating an immutable audit log before persisting data as encrypted JSON objects in the Secure Data Lake.

This raw data is then processed by the nightly ETL & Harmonisation Service, which performs three critical functions:

- 1) schema validation and normalization of units;
- 2) mapping of disparate clinical terminologies to a unified representation;
- 3) the application of differential privacy mechanisms (calibrated Laplacian noise) for small cohort sizes (<250 patients) to mathematically guarantee patient privacy [18].

The resulting harmonized dataset is consumed by the Model Trainer, where supervised fine-tuning of the base LLM (Llama-3) is performed using Low-Rank

Adaptation (LoRA) [19]. Unlike full fine-tuning, LoRA attaches trainable adapters to the frozen base model, which drastically reduces the number of trainable parameters and enables efficient, version-controlled storage of the resulting domain-adapted weights. These compact adapter weights, along with their associated training and evaluation metadata, are versioned and stored in the Model Registry. This registry acts as the system's "source of truth" for approved models, enabling governance, auditability, and the ability to perform instant rollbacks if a regression is detected. Finally, the Conversation API, which serves all live clinical interactions, is configured to load only the latest, governance-approved model version from this registry. This closed-loop design ensures that every patient interaction is handled by a vetted and traceable model instance, directly addressing key safety and deployment concerns in clinical AI [10, 11].

User interaction workflow

The infrastructure described above (Figure 1) provides the secure and governed foundation for online inference via the Conversation API, which exposes only approved model versions from the Model Registry. Building on this deployment layer, the platform implements a lightweight waiting-room interaction designed to standardize pre-ambulatory data collection while remaining efficient for patients. The workflow is summarized in Figure 2 and is implemented as a bounded, two-stage dialogue coordinated by two agents that invoke underlying LLMs (either generic or fine-tuned) rather than introducing additional independently trained models.

The interaction begins with Patient input, where the patient enters a free-text description of their main complaints. This input is sent to the serving layer through the Conversation API, which routes the request to the currently approved model version. The first stage is handled by Agent 1 (Questions generation). Agent 1 acts as an orchestration component responsible for selecting and structuring follow-up queries: given the initial symptom description, it calls the underlying LLM to produce exactly two targeted follow-up questions, each formulated as a four-option multiple-choice item. This design serves two purposes. First, the limited number of turns ensures predictable duration and reduces the risk of conversational drift. Second, the multiple-choice format lowers patient cognitive burden and yields standardized responses that can be reliably consumed by downstream components. In practice, the four-option

responses also reduce ambiguity compared to open-ended replies and simplify downstream synthesis.

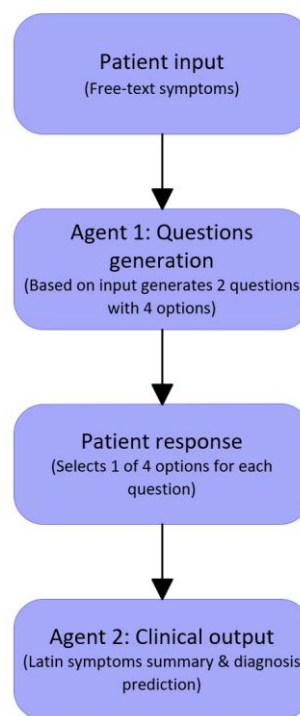


Fig. 2. User interaction workflow. The diagram shows a bounded two-agent dialogue from free-text symptoms to clinical output

After the patient answers the first question, the agent updates the dialogue state and triggers generation of the second question, conditioned on both the initial narrative and the first response. The patient then selects one of the four options for the second question. At this point, the workflow transitions to Agent 2 (Clinical output), which is responsible for synthesis rather than further elicitation.

Agent 2 aggregates the available information:

- 1) free-text symptoms;
- 2) the structured responses to both multiple-choice questions and invokes a (generic or fine-tuned) LLM to generate the final clinician-facing output.

Importantly, the agent logic enforces strict task boundaries (two questions only) and separates the elicitation objective from the synthesis objective, which improves controllability of the output.

The resulting Clinical output consists of two elements intended to support the physician during the pre-ambulatory phase:

- 1) a concise symptom summary expressed in professional medical (Latin) terminology, formatted as a preliminary anamnesis/record stub;

2) a diagnosis recommendation to guide differential diagnosis.

Importantly, the agents in this workflow function as control logic that constrains the interaction (fixed number of questions, fixed answer format) and routes requests to the appropriate deployed model through the Conversation API. This preserves the governance and versioning guarantees of the backend pipeline while enabling an adaptive, patient-friendly front-end experience.

User interface implementation

The User Interface (UI) was designed with an emphasis on simplicity and accessibility, as the target audience includes patients of varying ages, health literacy levels, and acute health conditions. As shown in Figure 3, the interface evokes a conversational experience (chatbot style), which is generally perceived as more natural for describing symptoms than rigid form-based questionnaires. The minimalist layout reduces visual clutter and guides the patient toward a single primary task: providing a concise free-text description of their main complaint. The landing screen prominently

communicates the application's purpose ("Patient just walked in..."), which helps establish trust and sets expectations before the automated triage begins.

In addition to free-text input, the UI includes a lightweight agent selection control that allows the patient (or assisting staff) to choose the most appropriate interaction pathway prior to submitting symptoms. This selection corresponds to a specialty-focused agent profile (e.g., dermatology for skin-related complaints), enabling the system to tailor the behavior of Agent 1 toward clinically relevant follow-up questions for the selected domain. For example, when a patient reports a rash or other cutaneous symptoms, selecting the dermatology-oriented agent biases the subsequent question generation toward discriminators such as lesion morphology, distribution, pruritus, or common triggers, improving the informativeness of the two-step dialogue while keeping the interaction bounded and easy to complete. Importantly, the agent selector is presented as a simple, clearly labeled option to avoid increasing cognitive load; if no specialty is chosen, the system defaults to a general triage configuration.

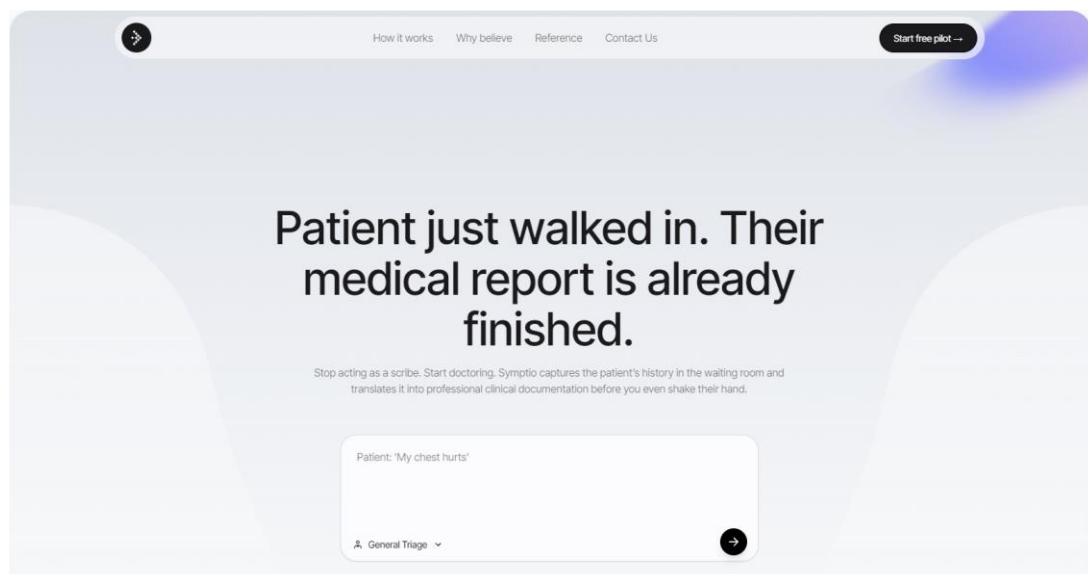


Fig. 3. Patients enter initial symptoms and may select a specialty-focused agent

From a usability perspective, the multiple-choice questions are displayed as clear, single-selection options, which supports rapid interaction in the waiting room and reduces the likelihood of incomplete or inconsistent answers. After completion, the patient-facing UI displays a confirmation message, thanking the patient for their input and informing them that their information is ready for the physician's review. Simultaneously, the platform

transmits the processed data to the clinician-facing summary view within the clinical interface. In this secure environment, the physician can review the generated Latinate note stub and diagnosis suggestion alongside other relevant patient records.

Overall, the UI design supports a smooth waiting-room workflow by combining natural language symptom entry with structured, low-effort interactions, while the

optional specialty agent selection provides a mechanism for more context-aware questioning without exposing technical complexity to the patient.

Methodology and experimental organization

The objective of the experiment was to quantify the added value of domain-specific fine-tuning compared to high-performance generic models used in a zero-shot setting.

The methodology focuses on three key areas:

- 1) curated dataset preparation from clinical records;
- 2) a reproducible fine-tuning protocol;
- 3) a multi-criteria evaluation capturing both diagnostic performance and the quality of the interactive questioning component.

Formalization of the workflow-constrained diagnostic task

Let x denote the initial free-text symptom narrative provided by the patient. The platform implements a bounded interaction I consisting of exactly two multiple-choice questions q_1, q_2 , each with four answer options, and the corresponding patient selections a_1, a_2 . The diagnostic task is defined as a multi-class prediction over a fixed label set $Y = \{y_1, \dots, y_K\}$ (here $K = 3$). Under a shared interaction specification, a model f maps the collected information to a clinician-facing output o and a diagnosis recommendation $\hat{y}$:

$$f:(x, q_1, a_1, q_2, a_2) \rightarrow (o, \hat{y}), \quad \hat{y} \in Y.$$

The interaction policy that generates questions can be expressed as:

$$q_n = g(x, q_{<n}, a_{<n}), \quad n \in \{1, 2\},$$

where g is constrained to produce exactly two questions and to follow a fixed four-option multiple-choice format.

This formulation explicitly separates:

- 1) elicitation (question generation) from
- 2) synthesis (clinical summary and diagnosis), enabling a controlled comparison between a generic zero-shot model and a domain-adapted model under identical constraints.

We test the hypotheses: H1: domain-specific fine-tuning improves diagnostic performance (macro-F1, Recall) under the bounded interaction protocol; H2: fine-tuning improves the discriminative relevance of generated follow-up questions as assessed by clinician raters; H3: the bounded interaction design

supports workflow efficiency without requiring long free-form dialogue.

Dataset curation and preprocessing

For training and validation, we used the publicly available disease-diagnosis-dataset hosted on Hugging Face, which provides free-text symptom descriptions paired with diagnosis labels [20]. From the full corpus, we filtered samples belonging to three target conditions frequently encountered in acute and pre-ambulatory settings: diabetes ($n = 312$), Chronic Obstructive Pulmonary Disease (COPD, $n = 289$), and persistent asthma ($n = 302$). The raw dataset was subsequently cleaned and normalized (e.g., removal of duplicates and incomplete entries, whitespace and punctuation normalization, and label harmonization for synonymous disease names). After preprocessing, the dataset was downsampled to 903 samples to match the intended experimental scale while preserving the class distribution across the three conditions.

The preprocessing pipeline transformed the symptom descriptions into a task-oriented representation reflecting the intended platform use case: an initial symptom narrative followed by a short, bounded interactive step and a final clinician-facing output. Concretely, the dataset was converted into "Instruction-Response" pairs following instruction-tuning methodology [21].

The instruction/input corresponded to the initial patient complaint in free-text form, while the response/target contained:

- 1) two follow-up questions formatted for a constrained interaction (consistent with the platform design);
- 2) the final diagnosis label. The follow-up questions were generated using a controlled template-based procedure to ensure consistent formatting and to reduce spurious variability across samples.

To support reproducibility, records were normalized into a consistent schema prior to conversion (e.g., standardized handling of missing fields and consistent formatting of symptom mentions). The dataset was then partitioned into an 80% training set and a 20% validation set using a stratified split to preserve the class proportions and prevent distributional shift between splits

Model selection and fine-tuning protocol

We compared two distinct modeling approaches:

Generic baseline (zero-shot): We utilized GPT-4 (OpenAI) [21] as the reference model, queried via

a system prompt without additional training. The baseline therefore reflects a "plug-and-play" deployment scenario in which the model is expected to generalize to the triage workflow without domain adaptation.

Fine-tuned model: We selected Llama-3 [22] as the open-source base model and optimized it specifically for the platform's triage interaction and output requirements.

Fair comparison and prompting control

To ensure a fair comparison between the zero-shot baseline and the fine-tuned approach, both models were evaluated under the same interaction specification and output constraints.

Concretely, the same system-level instruction template was used to enforce:

- 1) generation of exactly two follow-up questions;
- 2) each question being presented in a four-option multiple-choice format;
- 3) a final clinician-facing output consisting of a concise Latinate symptom summary and a single diagnosis recommendation.

This experimental design minimizes the influence of prompt engineering and response formatting differences, such that observed performance changes can be attributed primarily to domain adaptation through fine-tuning.

The training protocol employed LoRA [19], which enables parameter-efficient fine-tuning by freezing the base model weights and inserting trainable low-rank matrices into the transformer layers. This approach is particularly suitable for iterative experimentation and controlled deployment, as only compact adapter weights need to be versioned and served. Adapters were trained using a context window of 256 tokens, a learning rate of 2×10^{-4} , and the AdamW optimizer [23]. The optimization included 2,000 warm-up steps and early stopping based on the macro-F1 score measured on the validation set. Training was performed on four NVIDIA A100-80GB accelerators, with convergence reached in approximately three hours.

Evaluation metrics

The platform's performance was evaluated across two complementary dimensions – diagnostic accuracy and interaction quality – to reflect the dual nature of the system (prediction + targeted elicitation).

Diagnostic accuracy

The model's ability to correctly predict the diagnosis was measured using Precision, Recall, and

macro-averaged F1-score [24]. Evaluation was conducted on 200 synthetic patient vignettes supplemented with real laboratory values. Synthetic vignettes were used to systematically control symptom phrasing, coverage of key discriminators, and variability in presentation, while laboratory values increased clinical realism and helped test whether models can integrate structured signals with narrative text. All evaluated cases were held out from training and validation to avoid leakage.

Question relevance

Since a core feature of the platform is the generation of exactly two follow-up questions, their quality was assessed by two independent medical specialists. From a decision-theoretic perspective, the effectiveness of such a constrained interaction depends not on the number of questions, but on their discriminative importance with respect to the underlying diagnostic hypotheses. This aligns with prior research on importance analysis in decision-making systems, which demonstrates that a small number of high-impact factors can dominate the quality of the final decision outcome [25]. For each vignette, raters evaluated whether the generated questions were clinically relevant and whether they would plausibly contribute to narrowing the differential diagnosis in a pre-ambulatory setting. Inter-rater reliability was quantified using Cohen's kappa ($\kappa = 0.88$) [26], indicating strong agreement and supporting the validity of the human evaluation protocol. In addition to relevance, the fixed multiple-choice format was considered in the assessment, as the platform explicitly constrains questions to four answer options to standardize patient responses and reduce completion burden.

Results and discussion

Experimental results confirm that domain-specific fine-tuning yields significant improvements in both diagnostic accuracy and clinical documentation efficiency compared to generic models. Importantly, the observed gains were achieved without a trade-off in output safety, as the factual error rate remained low and we did not observe an increase in hallucinations.

Diagnostic performance comparison

Comparative analysis demonstrated a consistent performance increase when using our domain-adapted model. As shown in Table 1, the Fine-Tuned model achieved higher scores across all key metrics (Recall,

F1-score) compared to the GPT-4 reference (Generic Baseline). These performance differences were statistically significant ($p < 0.05$), paired t-test) [27], indicating that

the observed improvements were systematic rather than driven by isolated cases.

Table 1. Performance comparison between generic and fine-tuned models across three diagnostic groups

Diagnostic Group	Metric	Generic Baseline	Fine-Tuned	Δ (Improvement)
Diabetes	Recall	0.71	0.83	+0.12
	F1-score	0.73	0.83	+0.10
COPD	Recall	0.67	0.79	+0.12
	F1-score	0.69	0.80	+0.11
Asthma	Recall	0.68	0.80	+0.12
	F1-score	0.70	0.80	+0.10

The most substantial improvement (+19 percentage points in Recall) was observed in the detection of diabetic nephropathy. We attribute this gain to improved alignment with domain- and site-specific patterns, particularly in the interpretation of laboratory and biomarker cues (e.g., local conventions for reporting serum cystatin-C), which may be inconsistently expressed in generic model usage without domain adaptation. In COPD diagnosis, the fine-tuned model demonstrated an enhanced ability to recognize early signs of right heart strain in a manner consistent with GOLD-based severity cues [28]. For asthma, the detection of nocturnal symptoms improved significantly, suggesting better sensitivity to symptom variability patterns frequently found in patient narratives.

These findings align with a broader trend reported in the recent literature: constraining clinical LLM workflows and adapting models to task-specific data can improve reliability and reduce clinically unhelpful outputs relative to generic, unconstrained use [29]. In addition, studies benchmarking LLM-supported triage/diagnosis pipelines emphasize that workflow design (not only the raw model) materially affects clinical utility, particularly when outputs must be concise, structured, and decision-oriented [30].

Qualitative analysis of generated questions

A qualitative analysis revealed a fundamental difference in the nature of inquiry between the models. The Generic Baseline tended to ask questions that were semantically correct and empathetic but often clinically non-discriminatory (e.g., broad prompts such as "How are you feeling today?"). While such questions may improve conversational tone, they frequently failed to narrow the differential diagnosis efficiently under a bounded two-question interaction, and therefore contribute less under a Bayesian decision-making perspective [31].

In contrast, the Fine-Tuned model generated questions that directly tested differential diagnostic hypotheses. For example, for the input "chest pain", the model explicitly investigated discriminative features such as pain character ("Is the pain sharp or dull?") and radiation ("Does the pain radiate to the left shoulder or jaw?"). Across evaluated cases, the fine-tuned model more consistently prioritised high-yield discriminators over generic well-being probes. This behaviour can be interpreted as an implicit prioritisation of diagnostically important factors under uncertainty, rather than a purely conversational strategy. Similar prioritisation mechanisms have been formalised in importance-based decision models, where emphasising high-impact attributes leads to more reliable outcomes even when the overall information budget is limited [5, 25]. This observation is consistent with prior work indicating that off-the-shelf LLMs may produce conversationally plausible but diagnostically low-informative follow-up questions, motivating domain adaptation or workflow-level constraints for clinically useful inquiry [32].

Impact on clinical workflow

Integration of the platform into the triage process led to a 28% reduction in the time required to create an initial clinical record (from an average of 7.6 minutes to 5.5 minutes). A similar decrease was observed in keystroke count, directly reducing clinician cognitive load. In subjective evaluations using a Likert-type scale and standard analytical interpretation for ordinal response data [33], physicians reported a 2.1-point reduction in perceived administrative burden, suggesting that the benefit extends beyond objective efficiency measures.

These workflow-level improvements are directionally consistent with results reported for LLM-assisted documentation and ambient/digital scribe systems, where

reductions in documentation burden and improved clinical note efficiency have been observed in clinical or near-clinical evaluations [34]. While our system focuses specifically on the pre-ambulatory phase and uses a constrained two-question interaction rather than full-encounter transcription, the similarity in outcomes supports the broader conclusion that appropriately designed LLM integration can meaningfully reduce clerical workload and improve clinical throughput.

Limitations

Several limitations should be considered when interpreting these results. First, the evaluation was conducted on a restricted diagnostic scope (three disease groups), which may limit generalizability to broader triage settings. Second, diagnostic performance was assessed on 200 synthetic vignettes supplemented with real laboratory values; although this design enables controlled testing, it may not fully capture the variability of real-world patient narratives and comorbid presentations. Third, the interaction is intentionally constrained to two multiple-choice follow-up questions, which improves usability but may reduce performance for complex cases requiring longer anamnesis. Finally, the workflow-level impact metrics were collected in a specific clinical context; therefore, measured time savings and perceived burden reduction may differ across institutions, documentation practices, and patient populations.

Conclusion

The present work investigated the development and evaluation of an interactive platform for automated clinical triage, representing a workflow-oriented approach to supporting pre-ambulatory differential diagnosis. The primary objective was to determine whether high-capacity generic models used in a zero-shot setting or smaller, specialized models optimized through domain-specific fine-tuning are more suitable for this type of diagnostic decision support. In addition to accuracy, the study explicitly considered the quality of follow-up questioning as a core mechanism through which the platform reduces diagnostic uncertainty and prepares a clinician-ready record prior to examination.

Our findings support the hypothesis that for well-defined clinical tasks with constrained outputs, domain adaptation provides a measurable advantage over generalized model capabilities. While the generic model (GPT-4) demonstrated high linguistic competence and

produced fluent responses, it was less consistent in generating discriminatory follow-up questions that effectively narrow the differential diagnosis under the fixed two-question interaction constraint. In contrast, the developed fine-tuned model – trained on curated medical data – improved diagnostic performance (macro F1-score) and more reliably generated targeted inquiries aligned with clinical reasoning patterns, thereby producing higher-yield information for subsequent synthesis into a structured clinical summary.

From a practical perspective, the platform demonstrates direct relevance to healthcare efficiency and usability. The observed 28% reduction in documentation time suggests that constrained, workflow-integrated automation can reduce clerical workload and support clinicians in high-throughput settings. Importantly, efficiency improvements were accompanied by positive subjective feedback, indicating that the system can reduce perceived administrative burden rather than merely shifting effort to a different step of the workflow. Together, these findings highlight that the clinical value of LLM-based systems is not solely determined by model scale, but by the interaction design, output constraints, and task-specific optimization strategy.

The proposed system architecture also illustrates a feasible path for governed deployment of locally hosted models in clinical environments. Serving approved model versions through a controlled API layer supports auditability, version tracking, and rollback, while reducing reliance on external cloud services for sensitive data processing. Although security and privacy depend on institutional implementation details, local deployment can provide a stronger foundation for data minimization and governance by keeping inference within the healthcare organization's controlled infrastructure.

Future research will focus on extending the dataset and evaluation beyond the current diagnostic scope to cover a broader spectrum of conditions, including pediatric presentations and high-acuity emergency complaints. Further work will also include prospective validation in real clinical operation to assess long-term impact on diagnostic error rates, clinician workload, and patient experience, as well as monitoring robustness across diverse patient language, comorbidity profiles, and varying documentation practices. Overall, the platform represents a step toward hybrid medicine, where artificial intelligence does not replace the physician but becomes a bounded, reliable partner at the frontline of patient contact, supporting structured data collection and

accelerating early clinical decision-making. This is fully in line with the modern trend of precision or smart medicine [35].

Funding

The study was conducted without financial support.

Conflict of interest

The authors declare that they have no conflict of interest with respect to this research, including financial, personal, authorship, or other conflicts that could influence the research and its results presented in this article.

Data availability

The manuscript has no associated data.

Use of artificial intelligence

The authors confirm that they did not use artificial intelligence technology in the creation of this work.

References

1. World Health Organization (2016), "Global strategy on human resources for health: Workforce 2030", *Geneva: WHO*, 64 p., available at: <https://iris.who.int/handle/10665/250368>
2. Zaitseva, E., Levashenko, V., Rabcan, J., Krsak, E. (2020), "Application of the structure function in the evaluation of the human factor in healthcare", *Symmetry*, Vol.12(1), 93 p. DOI: <https://doi.org/10.3390/SYM12010093>
3. Shanafelt, T. D. et al. (2010), "Burnout and medical errors among American surgeons", *Annals of Surgery*, Vol. 251(6), pp. 995-1000. DOI: <https://doi.org/10.1097/SLA.0b013e3181bfdab3>
4. Zaitseva, E., Levashenko, V. (2026), "Reliability engineering in healthcare: Opportunities and challenges", *Reliability Engineering and System Safety*, Vol. 267, 111933 p. DOI: <https://doi.org/10.1016/j.ress.2025.111933>
5. Zaitseva, E., Levashenko, V. (2025), "Reliability Analysis Based on Aleatory and Epistemic Uncertainty Using Binary Decision Diagrams", *International Journal of Intelligent Systems*, Vol. 1, 6471577 p. DOI: [10.1155/int/6471577](https://doi.org/10.1155/int/6471577)
6. Sutton, R. T. et al. (2020), "An overview of clinical decision support systems: benefits, risks, and strategies for success", *NPJ Digital Medicine*, Vol. 3(17), pp. 1-10. DOI: <https://doi.org/10.1038/s41746-020-0221-y>
7. Chambers, D., Cantrell, A. J., Johnson, M. et al. (2019), "Digital and online symptom checkers and health assessment/triage services for urgent health problems: systematic review", *BMJ Open*, Vol. 9(8): e027743. DOI: <https://doi.org/10.1136/bmjopen-2018-027743>
8. Wallace, W., Chan, C., Chidambaram, S. (2022), "The diagnostic and triage accuracy of digital and online symptom checker tools: a systematic review", *NPJ Digital Medicine*, Vol. 5(1):18. DOI: <https://doi.org/10.1038/s41746-022-00667-w>
9. Hajijama, S., Juneja, D., Nasa, P. (2024), "Large Language Model in Critical Care Medicine: Opportunities and Challenges", *Indian Journal of Critical Care Medicine*, Vol. 28(6), pp. 523–525. DOI: <https://doi.org/10.5005/jp-journals-10071-24743>
10. Asgari, E., Montaña-Brown, N., Dubois, M., et al. (2025), "A framework to assess clinical safety and hallucination rates of LLMs for medical text summarisation", *NPJ Digital Medicine*, Vol. 8(274). DOI: <https://doi.org/10.1038/s41746-025-01670-7>
11. Chen, X., Xiang, J., Lu, S., Liu, Y., He, M., Shi, D. (2025), "Evaluating large language models and agents in healthcare: key challenges in clinical applications", *Intelligent Medicine*, Vol. 5(2), pp. 151–163. DOI: <https://doi.org/10.1016/j.imed.2025.03.002>
12. Singhal, K., et al. (2023), "Large language models encode clinical knowledge", *Nature*, Vol. 620, pp. 172-180. DOI: <https://doi.org/10.1038/s41586-023-06291-2>
13. Wang, G. et al. (2023), "ClinicalGPT: Large Language Models Finetuned with Diverse Medical Data and Comprehensive Evaluation", *CC BY 4.0*. DOI: <https://doi.org/10.48550/arXiv.2306.09968>
14. Tran, H. et al. (2024), "BioInstruct: instruction tuning of large language models for biomedical natural language processing", *Journal of the American Medical Informatics Association (JAMIA)*, Vol. 31(9), pp. 1821–1832. DOI: <https://doi.org/10.1093/jamia/ocae122>
15. Hu, E. J. et al. (2021), "LoRA: Low-Rank Adaptation of Large Language Models", *arXiv preprint*, arXiv:2106.09685. DOI: <https://doi.org/10.48550/arXiv.2106.09685>
16. Bender, D., Sartipi, K. (2013), "HL7 FHIR: An Agile and RESTful approach to healthcare information exchange", *Proceedings of the 26th IEEE International Symposium on Computer-Based Medical Systems (CBMS)*, pp. 326–331. DOI: <https://doi.org/10.1109/CBMS.2013.6627810>
17. NIST (2001), "Advanced Encryption Standard (AES)", *Federal Information Processing Standards Publication (FIPS PUB)* 197, 38 p. DOI: <https://doi.org/10.6028/NIST.FIPS.197-upd1>
18. Dwork, C. (2008), "Differential Privacy: A Survey of Results", *Theory and Applications of Models of Computation (TAMC)*, pp. 1–19. DOI: https://doi.org/10.1007/978-3-540-79228-4_1

19. Mu, O., Wang, W., Liu, W., et al. (2025), "Multimodal Large Language Model with LoRA Fine-Tuning for Multimodal Sentiment Analysis", *ACM Transactions on Intelligent Systems and Technology*, Vol. 16, No. 6, 139 p. DOI: <https://doi.org/10.1145/3709147>
20. Wei, J., Bosma, M., Zhao, V.Y. et al. (2021), "Finetuned Language Models Are Zero-Shot Learners", *ICLR*, DOI: <https://doi.org/10.48550/arXiv.2109.01652>
21. "GPT-4 Technical Report", *Computation and Language, arXiv preprint*, 2023. 100 p. DOI: <https://doi.org/10.48550/arXiv.2303.08774>
22. "Introducing Llama 3: The most capable openly available LLM to date", *Meta AI Blog*, 2024, available at: <https://ai.meta.com/blog/meta-llama-3>
23. Loshchilov, I., Hutter, F. (2019), "Decoupled Weight Decay Regularization", *International Conference on Learning Representations (ICLR)*, available at: <https://openreview.net/forum?id=Bkg6RiCqY7>
24. Semenov, S., Mozhaiev, O., Kuchuk, N., Mozhaiev, M., Tiulieniev, S., Gnusov, Yu., Yevstrat, D., Chyryva, Y., Kuchuk, H. (2022), "Devising a procedure for defining the general criteria of abnormal behavior of a computer system based on the improved criterion of uniformity of input data samples", *Eastern-European Journal of Enterprise Technologies*, Vol. 6(4-120), pp. 40–49, DOI: <https://doi.org/10.15587/1729-4061.2022.269128>
25. Zaitseva, E., Rabcan, J., Levashenko, V., Kvassay, M. (2023), "Importance analysis of decision-making factors based on fuzzy decision trees", *Applied Soft Computing*, Vol. 134, 109988 p. DOI: <https://doi.org/10.1016/j.asoc.2023.109988>
26. Engeler, J., Stricker, D., Birrenbach, T. et al. (2025), "Designing and validating entrustable professional activities for emergency medicine: a stringent, construct-validity enhancing approach", *BMC Medical Education*, Vol. 25, 866 p. DOI: <https://doi.org/10.1186/s12909-025-07449-4>
27. Raschka, S. (2018), "Model evaluation, model selection, and algorithm selection in machine learning", *arXiv preprint*. DOI: <https://doi.org/10.48550/arXiv.1811.12808>
28. GOLD (2024), "Global Strategy for the Diagnosis, Management, and Prevention of Chronic Obstructive Pulmonary Disease", *Global Initiative for Chronic Obstructive Lung Disease*, available at: <https://goldcopd.org/2024-gold-report/>
29. Dogru-Huzmeli, E., Moore-Vasram, S., Phadke, C. et al. (2025), "Evaluating ChatGPT's ability to simplify scientific abstracts for clinicians and the public", *Scientific Reports*, Vol. 15, 33466 p. DOI: <https://doi.org/10.1038/s41598-025-11086-8>
30. Gaber, F., Shaik, M., Allega, F. et al. (2025), "Evaluating large language model workflows in clinical decision support for triage and referral and diagnosis", *NPJ Digital Medicine*, Vol. 8, 263 p. DOI: <https://doi.org/10.1038/s41746-025-01684-1>
31. Sox, H. C. et al. (2013), *Medical Decision Making*, John Wiley & Sons, 328 p. DOI: <https://doi.org/10.1002/9781118341544>
32. Gatto, J., Seegmiller, P., Burdick, T. E. et al. (2025), "Follow-up Question Generation For Enhanced Patient-Provider Conversations", *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 25222–25240. DOI: <https://doi.org/10.18653/v1/2025.acl-long.1226>
33. Sullivan, G. M., Artino Jr., A. R. (2013), "Analyzing and interpreting data from Likert-type scales", *Journal of Graduate Medical Education*, Vol. 5(4), pp. 541–542. DOI: <https://doi.org/10.4300/JGME-5-4-18>
34. Olson, K. D., Meeker, D., Troup, M. et al. (2025), "Use of Ambient AI Scribes to Reduce Administrative Burden and Professional Burnout", *JAMA Network Open*, Vol. 8(10):e2534976. DOI: <https://doi.org/10.1001/jamanetworkopen.2025.34976>
35. Zaitseva, E., Levashenko, V., Rabcan, J., Kvassay M. (2023), "A New Fuzzy-Based Classification Method for Use in Smart/Precision Medicine", *Bioengineering*, Vol. 10(7), 838 p. DOI: <https://doi.org/10.3390/bioengineering10070838>

Received (Надійшла) 06.02.2026

Accepted for publication (Прийнята до друку) 21.05.2026

Publication date (Дата публікації) 27.06.2026

Відомості про авторів / About the Authors

Фандакова Міріам – Жилінський університет, аспірантка кафедри інформатики, факультет управлінських наук та інформатики, Жиліна, Словаччина;

Miriam Fandakova – University of Zilina, Postgraduate Student of the Informatics Department, Faculty of Management Science and Informatics, Zilina, Slovakia;

e-mail: miriam.fandakova@st.fri.uniza.sk

ORCID ID: <http://orcid.org/0009-0006-0763-3929>

Охотницький Шимон – Жилінський університет, аспірант кафедри інформатики, факультет управлінських наук та інформатики, Жиліна, Словаччина;

Simon Ochotnický – University of Zilina, Postgraduate Student of the Informatics Department, Faculty of Management Science and Informatics, Zilina, Slovakia;

e-mail: simon.ochotnický@st.fri.uniza.sk

ORCID ID: <http://orcid.org/0009-0002-2907-3687>

Коваленко Андрій Анатолійович – доктор технічних наук, професор, Харківський національний університет радіоелектроніки, завідувач кафедри електронних обчислювальних машин, Харків, Україна;

Andriy Kovalenko – Doctor of Technical Sciences, Professor, Kharkiv National University of Radio Electronics, Head of the Electronic Computers Department, Kharkiv, Ukraine;

e-mail: andriy.kovalenko@nure.ua

ORCID ID: <https://orcid.org/0000-0002-2817-9036>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=56423229200>

ІНТЕРАКТИВНА ПЛАТФОРМА ПІДТРИМКИ КЛІНІЧНИХ РІШЕНЬ: АРХІТЕКТУРА ТА ПОРІВНЯЛЬНА ОЦІНКА ЗАГАЛЬНОЇ ТА ТОЧНО НАЛАШТОВЛЕНОЇ ДІАГНОСТИКИ ЗА СИМПТОМАМИ ЗА МЕТОДОЛОГІЄЮ (LLMS)

Сучасні системи охорони здоров'я стикаються зі зростаючим навантаженням, фрагментованою комунікацією та обтяжливістю документації, що призводить до затримок і діагностичних помилок на етапі первинної сортування пацієнтів. Запропоноване рішення покликане загалом підвищити узгодженість, швидкість та стандартизацію процесу первинної оцінки пацієнтів. У цьому дослідженні представлено інтерактивну платформу підтримки клінічних рішень, яка реалізує двоступеневий процес збору анамнезу з обмеженим робочим процесом для забезпечення диференціальної діагностики на основі симптомів на доамбулаторному етапі надання медичної допомоги. Система вирішує три завдання: (T1) автоматизований збір анамнезу пацієнта за допомогою обмеженого діалогу (точно два додаткових запитання з варіантами відповідей), (T2) формалізація описів симптомів у структуровану медичну (латинську) термінологію для документації, призначеної для клініцистів, та (T3) генерація рекомендацій щодо диференціальної діагностики за фіксованою та клінічно інтерпретованою схемою виводу. Ми порівнюємо базову загальну велику мовну модель (GPT-4, zero-shot) з моделлю, адаптованою до домену (Llama-3, налаштованою за допомогою LoRA), за однакових обмежень щодо взаємодії, підказок та форматування. Експерименти на 903 записах симптомів та діагнозів і виокремленому наборі з 200 контрольованих віньєток показують, що адаптація до домену забезпечує покращення макро-F1 на 10–12% та приблизно на 15% вищий показник Recall у складних випадках, водночас генеруючи клінічно більш дискримінативні та діагностично релевантні додаткові запитання, за оцінкою двох медичних експертів ($\kappa = 0,88$). У оцінці на рівні робочого процесу платформа скоротила час на документацію на 28% порівняно зі стандартним прийомом та зменшила адміністративні зусилля, необхідні для підготовки початкового резюме для лікаря. Ці результати свідчать про те, що спеціалізовані, налагоджені моделі, які можна розгорнути локально, у поєднанні з обмеженим дизайном взаємодії, стандартизованими обмеженнями на вихідні дані та інтеграцією системи, орієнтованою на робочий процес, забезпечують безпечний, практичний та ефективний шлях для включення підтримки прийняття рішень на основі LLM у рутинні клінічні робочі процеси до амбулаторного лікування, зберігаючи при цьому безпеку, зручність використання та можливість аудиту на практиці.

Ключові слова: підтримка клінічних рішень; моделі з великою мовою програмування; точне налаштування; LoRA; діагностика на основі симптомів; медичний ШІ.

Бібліографічні описи / Bibliographic descriptions

Фандакова М., Охотницький Ш., Коваленко А. А. Інтерактивна платформа підтримки клінічних рішень: архітектура та порівняльна оцінка загальної та точно налаштованої діагностики за симптомами за методологією (LLMS). *Сучасний стан наукових досліджень та технологій в промисловості*. 2026. № 2 (36). С. 173–184. DOI: <https://doi.org/10.30837/2522-9818.2026.2.173>

Fandakova, M., Ochotnický, S., Kovalenko, A. (2026), "Interactive platform for supporting clinical decision-making using large language models (LLMs)", *Innovative Technologies and Scientific Solutions for Industries*, No. 2 (36), P. 173–184. DOI: <https://doi.org/10.30837/2522-9818.2026.2.173>

Iryna Chepurna

MODEL OF RESOURCE AND QUORUM MANAGEMENT FOR A BLOCKCHAIN-ORIENTED DISTRIBUTED INFORMATION SYSTEM IN A MULTI-CLOUD ENVIRONMENT

The subject of this study is a mathematical model for managing computing resources and ensuring quorum in a blockchain-oriented distributed information system within a multi-cloud environment. **The goal** of the work is to develop a set-theoretic model that formalizes the multi-cloud deployment of components of a blockchain-oriented distributed information system, taking into account its multi-level architecture, the heterogeneity of available resources, the aggregated resource state of nodes, their availability, and the conditions for ensuring quorum during operation under variable load. **Methods.** To achieve the set goal, methods of system analysis were used to formalize the system structure and the connections between its components; methods of mathematical modeling were used to construct a formal description of the multi-cloud environment, the resource state of nodes, control actions, and quorum constraints; statistical analysis methods for processing, aggregating, and normalizing telemetry data; decision-making theory methods for establishing conditions for acceptable resource reallocation; and simulation modeling methods for verifying the performance of the developed model under conditions of variable load, resource scarcity, and variable node availability. **Results.** For the first time, a formalized set-theoretic model has been developed for managing computing resources and ensuring quorum in a blockchain-oriented distributed information system within a multi-cloud environment. Unlike existing approaches, combines, within a single unified representation, the multi-cloud environment as a set of cloud platforms from different vendors, the multi-level system architecture, the aggregated resource state of nodes, the conditions for permissible resource reallocation, and the specified requirements for ensuring quorum. Sets of cloud environments, architectural levels, and system nodes are presented; vectors of available and allocated resources, an aggregated system state vector, a resource deviation function, a quorum assurance function, and a set of permissible control actions are introduced. Simulation studies confirmed the model's operability under nominal operation, resource scarcity, and node availability degradation. This allowed for the formulation of the decision-making process regarding resource reallocation and system operability maintenance within a single mathematical framework. **Conclusions.** The developed model provides a formalized description of the processes of managing computational resources and ensuring quorum in a blockchain-oriented distributed information system in a multi-cloud environment and can be used as a basis for further development of methods for optimizing control decisions and experimental research into the system's operational efficiency under dynamic conditions.

Keywords: mathematical model; resource management; quorum; blockchain; information system; cloud.

1. Introduction

The digital transformation of the financial sector, logistics, healthcare, telecommunications, and other industries is accompanied by the widespread adoption of distributed information systems. Such systems ensure geographically distributed data processing, service availability, scalability of computing power, and operational continuity under varying information flow intensities [1, 2].

One of the directions in the development of distributed information systems is the integration of blockchain technologies into their architecture. The use of blockchain enables data integrity, transparency in transaction recording, increased trust among participants in information exchange, and support for decentralized decision-making procedures [3]. For blockchain-oriented distributed information systems,

this is of fundamental importance, since their operability is determined by the availability of individual nodes and the ability to maintain the coordinated functioning of a set of components involved in forming and confirming the system's state [4].

The further development of such systems is linked to the transition to multi-cloud infrastructures [5]. Deploying system components in cloud environments from different vendors allows for reduced dependence on a single vendor, increased scaling flexibility, improved fault tolerance, and load distribution across independent computing platforms [6, 7]. At the same time, multi-cloud deployment complicates system management due to differences in resource allocation models, service characteristics, telemetry data formats, network interaction parameters, and conditions for ensuring node availability [8, 9].

For a blockchain-oriented distributed information system, these complications are of a complex nature. Under such conditions, it is necessary to ensure the rational distribution of computing resources among nodes, services, and deployment environments, as well as to maintain a set of available nodes that preserves quorum and the system's ability to function in a coordinated manner [10, 11]. Given dynamic load changes, partial node degradation, changes in the availability of individual cloud environments, and the heterogeneity of telemetry data, the task of resource management is directly linked to the task of maintaining quorum [11–13].

Approaches to resource management in cloud and distributed systems are primarily focused on service scaling, load balancing, telemetry monitoring, or separate analysis of blockchain interaction mechanisms [14]. At the same time, within the framework of such approaches, the relationship between the resource state of nodes, their availability in a multi-cloud environment, and the conditions for ensuring the quorum of a blockchain-oriented system is not sufficiently formalized. This limits the ability to develop coordinated solutions for resource reallocation in situations where changes to the system configuration affect performance, availability, and the ability to maintain the correct functioning of consensus mechanisms [15–17].

Therefore, the task of developing a mathematical model for managing resources and quorum in a blockchain-oriented distributed information system in a multi-cloud environment is relevant. Such a model must reflect the distribution of system components across cloud environments from different vendors, account for the system's multi-tier architecture, the heterogeneity of computational resources and telemetry metrics, and formalize the conditions under which resource reallocation aligns with requirements for node availability and quorum maintenance. The construction of such a model creates a foundation for the further development of methods for optimizing the management of blockchain-oriented distributed systems in multi-cloud infrastructures.

2. Analysis of current scientific publications and statement of the study problem

The issue of managing computing resources in distributed information systems plays a significant role in current research due to the growing scale of digital services, increasing demands for availability, and the

need to ensure the stable operation of systems under fluctuating loads. The integration of blockchain technologies into distributed information systems expands their functional capabilities by ensuring data integrity, transaction transparency, and support for decentralized interaction [18]. At the same time, this complicates resource management processes, since the operability of such systems is determined by the sufficiency of computing power, the composition of available nodes, and the ability to maintain coordinated functioning of components involved in confirming the system's state and executing transactions [19, 20].

The paper [21] examines the architectural aspects of implementing the Blockchain-as-a-Service concept and approaches to deploying blockchain network nodes in a cloud environment. It is shown that delegating part of the infrastructure functions to a cloud provider simplifies the deployment and maintenance of blockchain services. However, the presented approach focuses primarily on the infrastructure organization of the service and does not include a formalized description of coordinated management of node resources and quorum maintenance when system components are deployed across multiple independent cloud environments.

A stochastic approach to scaling is presented in [22], where the Ito equation and the Wiener process are used to describe changes in traffic intensity in real time. The results confirm the feasibility of developing adaptive scaling policies based on predictive load estimates. However, the presented model is focused on a single-cloud environment; therefore, differences between cloud vendors and the conditions for ensuring quorum in a blockchain-oriented distributed information system are not formalized within its framework.

Another approach is presented in [23], where an adaptive approach to managing cloud infrastructure resources based on intelligent load forecasting algorithms is proposed. The authors demonstrate that the use of recurrent neural networks allows for increased efficiency in resource reallocation and reduced operational costs compared to static schemes. The primary focus of this study is on the performance and economic efficiency of cloud services. At the same time, the specifics of blockchain-oriented distributed systems, for which, in addition to resource provisioning, the conditions for node availability and quorum maintenance must be formalized, are partially revealed.

A separate class of approaches is related to the migration of virtual machines and periodic analysis of

node status. It is precisely this logic that forms the basis of the study [24], which examines resource management in distributed cloud infrastructures under conditions of resource scarcity. The advantage of this approach is the relative simplicity of the mathematical formulation and its practical focus on maintaining service quality metrics. At the same time, it does not address the task of consistently accounting for heterogeneous telemetry data obtained from cloud environments of different vendors, nor does it formalize the impact of changes in the composition of available nodes on maintaining the quorum of a distributed system.

General aspects of the described problem of using blockchain technologies in digital environments and issues of resource provision for complex computing platforms are discussed in [21, 23]. Requirements for the secure operation of distributed platforms are considered in [25]. These results are important for establishing the general context of the research; however, they lack a comprehensive mathematical model that combines the deployment of the system across multiple cloud environments, the multi-level organization of its components, the aggregated resource state of nodes, and formalized conditions for ensuring quorum.

The analysis conducted indicates that existing studies cover individual aspects of building blockchain-oriented systems, service scaling, telemetry monitoring, resource reallocation, and the operation of cloud infrastructures [20, 22–24]. The approaches considered overlook the task of formalizing a unified management model that combines the resource aspect of a blockchain-oriented distributed information system's operation with the conditions for maintaining a quorum in a multi-cloud environment.

An unresolved part of the general problem is the development of a mathematical model for managing the computational resources and quorum of a blockchain-oriented distributed information system, in which the multi-cloud environment is presented as a set of cloud platforms from different vendors, and the system itself is viewed as a multi-level structure of interconnected nodes and services. Such a model must account for the heterogeneity of available resources and telemetry data, dynamic changes in load, changes in node availability, as well as the conditions under which resource reallocation is aligned with the requirements for maintaining the quorum and the stable operation of the system as a whole.

3. Research aim and objectives

The aim of this work is to develop a mathematical model for managing computing resources and the quorum of a blockchain-oriented distributed information system in a multi-cloud environment. The model under development is intended to formalize the representation of a multi-cloud environment as a set of cloud platforms from different vendors, reflect the system's multi-level architecture, establish the relationship between the resource state of nodes, their availability, and the conditions for ensuring quorum, and define formal conditions for the system's transition between states during resource reallocation in the course of operation.

To achieve the stated goal, the following tasks must be solved in this work:

- present a mathematical description of a blockchain-oriented distributed information system in a multi-cloud environment, taking into account the deployment of its components across cloud platforms from different vendors, the heterogeneity of available resources, the multi-level organization of nodes, and dynamic changes in the system's state;
- develop a mathematical description of the process of managing computational resources based on aggregated metrics of the resource system's steady-state operation, resource shortages, and changes in load intensity;
- define formalized conditions for node availability and quorum assurance as constraints that must be adhered to when making decisions regarding resource reallocation to maintain system operability in a multi-cloud environment.

4. Materials and methods of the study

The subject of this study is the processes of managing computing resources and state of nodes, which allows for accounting for the ensuring a quorum of available nodes in a blockchain-oriented distributed information system within a multi-cloud environment.

The subject of the study is a mathematical model for managing computational resources and ensuring quorum in a blockchain-oriented distributed information system in a multi-cloud environment, which takes into account the placement of nodes on cloud platforms from different vendors, the heterogeneity of available resources, aggregated resource state metrics, and dynamic changes in the system's state during its operation.

The research hypothesis is that the use of a formalized set-theoretic model, which combines an aggregated representation of the resource state of nodes with specified quorum conditions, enables coordinated decisions regarding the redistribution of computational resources and maintains the operability of a blockchain-oriented distributed information system in a multi-cloud environment under conditions of variable load and variable node availability.

To achieve the set goal, the following were used:

- methods of system analysis to formalize the composition of system components, their interrelationships, and the representation of a multi-cloud environment as a set of cloud platforms from different vendors;
- methods of mathematical modeling to construct a formal description of the resource state of nodes, their availability conditions, and quorum constraints;
- statistical analysis methods for processing, aggregating, and normalizing telemetry data;
- decision-making theory methods for formalizing the conditions for resource reallocation as the system state changes;
- simulation modeling methods for verifying the performance of the developed model under variable load conditions and changes in the composition of available nodes.

The information basis of the study is a set of telemetry data received from nodes of a blockchain-oriented distributed information system deployed in a multi-cloud environment. This data includes metrics on CPU time and RAM usage, as well as node availability indicators in each cloud environment. The collected values form a multidimensional time series reflecting changes in the system's resource state at discrete points in time.

For a node $v_{c,\ell,j}$ located in a cloud environment c , belonging to level ℓ of a multi-level architecture, and having index j , the vector of primary metrics at time t is given in the form

$$\bar{\mu}_{c,\ell,j}(t) = \{\mu_{c,\ell,j}^{\alpha}(t), \mu_{c,\ell,j}^{\beta}(t), \mu_{c,\ell,j}^{\gamma}(t)\}, \quad (1)$$

where $\mu_{c,\ell,j}^{\alpha}(t)$ – processor time usage metric; $\mu_{c,\ell,j}^{\beta}(t)$ – RAM usage metric; $\mu_{c,\ell,j}^{\gamma}(t)$ – a metric that reflects the node's availability at a given point in time t .

To ensure the comparability of heterogeneous metrics (1) obtained from cloud platforms of different

vendors, a normalization procedure is applied. The normalized metric value is defined as

$$\hat{\mu}_{c,\ell,j}^{\kappa}(t) = \sqrt{\frac{\mu_{c,\ell,j}^{\kappa}(t) - \mu_{\kappa}^{\min}}{\mu_{\kappa}^{\max} - \mu_{\kappa}^{\min}}}, \quad \kappa \in K, \quad (2)$$

where $K = \{\alpha, \beta, \gamma\}$ – a set of metric types; $\mu_{c,\ell,j}^{\kappa}(t)$ – current value of the metric of type κ ; $\mu_{\kappa}^{\max}$ and $\mu_{\kappa}^{\min}$ – the minimum and maximum values of the metric of type κ in the system under study.

Normalization (2) ensures that the metrics are scaled to a common scale and increases the model's sensitivity to changes in the resource state of the nodes. The normalized node metrics are used to form an aggregated system state vector at a given point in time t :

$$\bar{S}(t) = \{\hat{\mu}_{c,\ell,j}^{\kappa}(t) \mid c \in C, \ell \in L, j \in V_{c,\ell}, \kappa \in K\}. \quad (3)$$

In formula (3), the set of vectors $\bar{S}(t)$ components reflect the resource status of the nodes in all cloud environments that make up the blockchain-oriented distributed information system.

The simulation study was conducted in two stages. In the first stage, a software model was implemented for the processes of collecting, normalizing, and aggregating telemetry metrics, as well as the decision-making logic for resource reallocation based on changes in the system's state. For this purpose, Python 3.11 and the NumPy and Pandas libraries were used, which enable the processing of multidimensional data arrays and the modeling of discrete system transitions between states.

In the second stage, a virtualized environment built on Oracle VM VirtualBox, Proxmox VE, Docker, and Kubernetes was used to verify the operational capability of the developed model. This configuration made it possible to simulate a multi-cloud environment as a set of isolated computing platforms with different resource provisioning parameters and to generate simulation scenarios of system operation under conditions of variable load, resource shortages, and a variable composition of available nodes.

The effectiveness of the developed model is evaluated based on the results of simulation scenarios, within which changes in the aggregated resource state of nodes, the feasibility of resource reallocation, and the fulfillment of quorum conditions are analyzed. This approach allows us to verify the model's suitability for describing the management processes of a blockchain-oriented distributed information system in a multi-cloud environment.

5. Research results

5.1. Formalization of multi-cloud component deployment and the system resource model

In this study, a blockchain-oriented distributed information system is viewed as a set of components deployed across multiple cloud environments from different vendors. Control decisions regarding changes in resource allocation and the configuration of active components are generated by a logically isolated element of the control subsystem. A block diagram of this interaction is shown in Fig. 1. Unlike the previous study [26], in which the formalization of metric collection was performed for a single cloud environment, in this work the system model is extended to the case of multi-cloud deployment of components and subsequent decision-making regarding resource reallocation.

A multi-cloud environment is represented as an ordered set of independent cloud platforms

$$C^* = \{C_1, C_2, \dots, C_M\}, \quad (4)$$

where M – the number of cloud environments in which the components of a blockchain-based distributed information system are deployed.

Each cloud environment C_c , $c = \overline{1, M}$ is characterized by its own vector of available computing resources:

$$\bar{R}_c = \{R_c^{\xi_1}, R_c^{\xi_2}, \dots, R_c^{\xi_H}\}, \quad (5)$$

where ξ_h , $h = \overline{1, H}$ – type of computing resource.

Further, ξ_1 may refer to processing resources, ξ_2 – RAM, ξ_3 – storage resource, ξ_4 – network resource.

The system architecture is viewed as multi-tiered, although the set of tiers is not identical to the set of cloud environments. For a formal description, let us introduce the set of levels

$$L = \{\ell_1, \ell_2, \dots, \ell_L\}, \quad (6)$$

where L – number of system architecture levels.

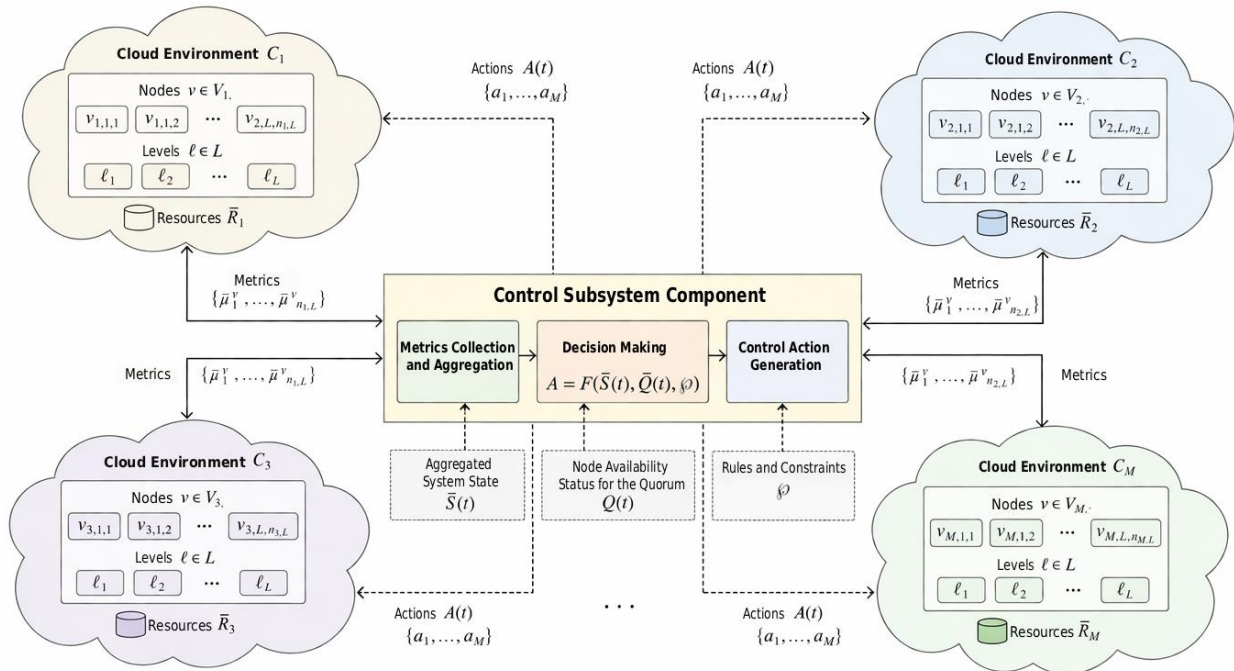


Fig. 1. Block diagram illustrating the interaction between a control subsystem element and the components of a blockchain-based distributed information system in a multi-cloud environment

Depending on the implementation method, the set L may include the infrastructure layer, the virtualization layer, the containerized services layer, the blockchain nodes layer, and the application services layer.

For each pair (C_c, ℓ) , where $C_c \in C^*$, $\ell \in L$, let's enter a set of nodes

$$V_{c,\ell} = \{v_{c,\ell,1}, \dots, v_{c,\ell,N_{c,\ell}}\}, j = 1, \dots, N_{c,\ell}, \quad (7)$$

where $N_{c,\ell}$ – number of nodes in the layer ℓ , located in a cloud environment C_c .

Then the complete set of nodes in the system is defined as:

$$V = \bigcup_{c=1}^M \bigcup_{\ell \in L} V_{c,\ell}. \quad (8)$$

Each node $v_{c,\ell,j} \in V_{c,\ell}$ is associated with a vector of resources allocated for its operation:

$$\bar{r}_{c,\ell,j} = \{r_{c,\ell,j}^{\xi_1}, r_{c,\ell,j}^{\xi_2}, \dots, r_{c,\ell,j}^{\xi_H}\}, \quad (9)$$

where $r_{c,\ell,j}^{\xi_H}$ – volume of resource type ξ_h , assigned to the node $v_{c,\ell,j}$ in a cloud environment C_c .

In addition to the resources directly allocated to the system's nodes, in every cloud environment, a portion of the resources is used to support platform services, network interaction, system services, and virtualization tools. For a cloud environment C_c , these costs are represented by a vector

$$\bar{S}_c = \{S_c^{\xi_1}, S_c^{\xi_2}, \dots, S_c^{\xi_H}\}. \quad (10)$$

In that case, the conditions for the acceptable placement of system components within the cloud environment C_c are determined as follows:

$$\sum_{\ell=1}^L \sum_{j=1}^{N_{c,\ell}} \bar{r}_{c,\ell,j} + \bar{S}_c \preceq \bar{R}_c, \quad c = \overline{1, M}, \quad (11)$$

where symbol $\preceq$ denotes element-wise comparison of vectors.

Equation (11) specifies the condition under which the total resource consumption for the operation of nodes at all levels and the service infrastructure does not exceed the available resources of the corresponding cloud environment.

The resource state of the system's nodes at a given point in time t is described by normalized metrics aggregated into a vector $\bar{S}(t)$ in accordance with relations (1)–(3). It is this vector $\bar{S}(t)$ that reflects the current state of the set V , taking into account the placement of nodes in different cloud environments and at different architectural levels. In the structure shown in Fig. 1, the control subsystem element receives $\bar{S}(t)$ as a generalized representation of the system's current state, information about node availability, and decision-making rules, based on which it generates control actions to change the distribution of resources and the configuration of active components.

In light of the above, the functioning of an element of the control subsystem at the structural level can be expressed as a formula:

$$A(t) = F(\bar{S}(t), \bar{Q}(t), \wp), \quad (12)$$

where $A(t)$ – a generalized representation of a control decision at a given point in time t , which can be specified by a set of control actions; $\bar{Q}(t)$ – a vector representing the availability of nodes relevant to quorum maintenance; $\wp$ – a set of rules and constraints that determine the admissibility of management decisions.

Equation (12) defines the general decision-making framework and serves as the basis for the subsequent formalization of the resource management process and the conditions for ensuring a quorum.

Thus, equations (4)–(12) define the formal basis of the model for the multi-cloud deployment of components in a blockchain-oriented distributed information system. The proposed representation distinguishes between the set of cloud environments C , the set of architectural levels L , the set of nodes V , and the vectors of available and allocated resources, which forms the basis for describing resource reallocation procedures and verifying quorum conditions.

5.2. Mathematical description of the management process for computing resources

The process of managing computational resources can be represented as a sequence of discrete control decisions made by an element of the control subsystem based on the aggregated state vector $\bar{S}(t)$, defined by Equation (3). Within this framework, the control subsystem evaluates the current distribution of resources among nodes at every instant t , determines the need for their adjustment, and generates an admissible control action.

We define the current distribution of computational resources for node $v_{c,\ell,j}$ at time t by the vector

$$\bar{r}_{c,\ell,j}(t) = \{r_{c,\ell,j}^{\xi_1}(t), r_{c,\ell,j}^{\xi_2}(t), \dots, r_{c,\ell,j}^{\xi_H}(t)\}, \quad (13)$$

In (13), the components of the vector $\bar{r}_{c,\ell,j}(t)$ specify the actual amount of resources of types ξ_h allocated to node $v_{c,\ell,j}$ in the cloud environment C_c at control step t .

It is advisable to determine the node's resource demand at time t based on its metric vector $\bar{\mu}_{c,\ell,j}(t)$ (1), using the formula:

$$\bar{d}_{c,\ell,j}(t) = G_{c,\ell,j}(\bar{\mu}_{c,\ell,j}(t)), \quad (14)$$

where $G_{c,\ell,j}(\cdot)$ is a function that transforms the node's metrics into a vector of required resources.

Equation (14) provides an estimate of the number of resources required by the node $v_{c,\ell,j}$ to maintain its operation under the current load conditions. The resource deviation vector is determined based on the difference between the required and actually allocated amounts of resources:

$$\bar{\Delta}_{c,\ell,j}(t) = \bar{d}_{c,\ell,j}(t) - \bar{r}_{c,\ell,j}(t). \quad (15)$$

In equation (15), the positive components of the vector $\bar{\Delta}_{c,\ell,j}(t)$ correspond to a resource deficit, while the negative components correspond to the presence of a reserve.

The control influence on the node $v_{c,\ell,j}$ is represented by the vector

$$\bar{u}_{c,\ell,j}(t) = \{u_{c,\ell,j}^{\xi_1}(t), u_{c,\ell,j}^{\xi_2}(t), \dots, u_{c,\ell,j}^{\xi_H}(t)\}. \quad (16)$$

where $\bar{u}_{c,\ell,j}(t)$ is the change in the volume of the resource of type ξ_h , which is performed at control step t .

Then, the transition to a new resource allocation is determined by the relation:

$$\bar{r}_{c,\ell,j}(t+1) = \bar{r}_{c,\ell,j}(t) + \bar{u}_{c,\ell,j}(t). \quad (17)$$

Equation (17) defines the dynamics of changes in a node's resource state under the influence of a control decision. The validity of such a decision must be verified at both the cloud environment level and the individual node level.

Constraints on the total amount of resources that can be allocated to nodes within the cloud environment C_c are determined by the condition:

$$\sum_{\ell=1}^L \sum_{j=1}^{N_{c,\ell}} (\bar{r}_{c,\ell,j} + \bar{u}_{c,\ell,j}) + \bar{S} \leq \bar{R}_c, c = \overline{1, M}. \quad (18)$$

Condition (18) aligns control actions with the available resources of the corresponding cloud environment and the service costs introduced in formula (10).

For each node, permissible limits for resource allocation are also specified:

$$\bar{r}_{c,\ell,j}^{\min} \leq \bar{r}_{c,\ell,j}(t+1) \leq \bar{r}_{c,\ell,j}^{\max}. \quad (19)$$

Equation (19) constrains the new node resource vector with the platform's technical parameters, virtualization parameters, and service configuration limits.

A separate condition for stable node operation is introduced, defined by the interval of permissible resource states

$$\bar{r}_{c,\ell,j}^{\min} \leq \bar{r}_{c,\ell,j}(t+1) \leq \bar{r}_{c,\ell,j}^{\max}. \quad (20)$$

$\bar{r}_{c,\ell,j}^{\min}$ $\bar{r}_{c,\ell,j}^{\max}$ – the lower and upper bounds of the resource interval for stable node operation, respectively $v_{c,\ell,j}$.

Based on condition (20), it is appropriate to introduce a node stability indicator:

$$\chi_{c,\ell,j}(t+1) = \begin{cases} 1, & \bar{r}_{c,\ell,j}^{\min} \leq \bar{r}_{c,\ell,j}(t+1) \leq \bar{r}_{c,\ell,j}^{\max}, \\ 0, & \text{otherwise.} \end{cases} \quad (21)$$

Equation (21) allows us to transition from a continuous description of the resource state to a logical indicator of a node's operational readiness following the execution of a control action.

To select a control decision regarding resource reallocation, we introduce a quality functional:

$$J_R(t) = \sum_{c=1}^M \sum_{\ell=1}^L \sum_{j=1}^{N_{c,\ell}} \omega_{c,\ell,j} \|\bar{\Delta}_{c,\ell,j}(t) - \bar{u}_{c,\ell,j}(t)\|_1, \quad (22)$$

where $\omega_{c,\ell,j}$ is the weight coefficient of node priority $v_{c,\ell,j}$; $\|\cdot\|_1$ is the L_1 is the norm of the vector.

Functional (22) characterizes the total unregulated resource deviation after control actions have been performed. The task of managing computational resources then reduces to selecting a set of vectors $\bar{u}_{c,\ell,j}(t)$ for which functional (22) takes on a minimum value subject to constraints (18)–(21):

$$A_R^*(t) = \arg \min_{A_R(t)} J_R(t), \quad (23)$$

where $A_R(t) = \{\bar{u}_{c,\ell,j}(t)\}$ is the set of control actions formed at step t .

Thus, equations (13)–(23) define the resource component of the mathematical model for controlling a blockchain-oriented distributed information system in a multi-cloud environment. Within this component, the aggregated system state $\bar{S}(t)$, defined by equation (3), is used to assess the nodes' resource requirements, generate control actions, and verify their admissibility under resource constraints.

5.3. Formalized conditions for ensuring quorum and the admissibility of system transitions between states

The resource component of the model defines the rules for generating control actions based on the current aggregated state of the system $\bar{S}(t)$; however, for a blockchain-oriented distributed information system, the validity of a control decision is determined by resource constraints and the preservation of a sufficient number of available nodes to ensure a quorum. To this end, we introduce a formalized description of the set of nodes relevant for maintaining a quorum.

Let:

$$V^q \subseteq V. \quad (24)$$

In (24), V^q is a subset of the system's nodes whose participation is taken into account when forming a quorum. The set may include validators, transaction confirmation nodes, replication coordinators, or other components upon which the maintenance of the system's consistent state depends.

For each node $v_{c,\ell,j} \in V^q$, let us introduce a quorum availability indicator

$$\eta_{c,\ell,j} = \begin{cases} 1, & \mu_{c,\ell,j}^v \geq \theta^v \wedge \chi_{c,\ell,j}(t) = 1, \\ 0, & \text{otherwise,} \end{cases} \quad (25)$$

where θ^v is the availability threshold.

Equation (25) establishes that a node is taken into account when forming a quorum only if two conditions are met simultaneously: the node is accessible according to its state metric, and its resource state falls within the stable operation interval.

The availability state of nodes relevant to the quorum at time t is given by

$$\bar{Q}(t) = \{\eta_{c,\ell,j}(t) \mid v_{c,\ell,j} \in V^q\}. \quad (26)$$

Based on the vector $\bar{Q}(t)$ (26), we introduce the quorum assurance functional:

$$\Psi(t) = \sum_{v_{c,\ell,j} \in V^q} \psi_{c,\ell,j} \eta_{c,\ell,j}(t), \quad (27)$$

where $\psi_{c,\ell,j}$ is the weight coefficient of node $v_{c,\ell,j}$ in the quorum formation procedure.

In the simplest case, for equivalent nodes, it is assumed that $\psi_{c,\ell,j} = 1$; then the functional $\Psi(t)$, defined by equation (27), reduces to the number of available nodes participating in maintaining the quorum. For systems with differentiated roles or non-equivalent validators, the values of $\psi_{c,\ell,j}$ can be set in accordance with the adopted consensus policy.

We formulate the quorum condition as follows:

$$\Psi(t) \geq \Psi^{\min}, \quad (28)$$

where $\Psi^{\min}$ is the minimum acceptable value of the quorum assurance functional to maintain the operability of a blockchain-oriented distributed information system.

Equation (28) serves as a formal criterion for the admissibility of the current set of available nodes. Accordingly, a control action generated by Equation (23) can be considered admissible only if, after its execution, the value of the functional $\Psi(t+1)$ does not fall below the threshold value $\Psi^{\min}$.

Taking into account the constraints (18)–(21) and the quorum condition (28), the set of admissible control actions is defined as

$$A^{\text{adm}}(t) = \left\{ A_R(t) \mid \begin{array}{l} \text{conditions are met (18)–(21)} \\ \Psi(t+1) \geq \Psi^{\min} \end{array} \right\}. \quad (29)$$

Then the resource and quorum control problem reduce to selecting a control decision that minimizes the resource deviation functional (22) over the set of admissible actions (29):

$$A^*(t) = \arg \min_{A_R(t) \in A^{\text{adm}}(t)} J_R(t). \quad (30)$$

Equation (30) integrates the resource and quorum components of the model. Unlike Problem (23), where the solution is selected based solely on the resource criterion, in Equation (30) the admissibility of an action is also determined by the condition that the quorum is maintained after the resource allocation changes.

To describe the system's transition between states, we introduce a generalized state:

$$\Omega(t) = \{\bar{S}(t), \bar{Q}(t)\}. \quad (31)$$

In equation (31), the vector $\bar{S}(t)$ (3) reflects the aggregated resource state of the system; $\bar{Q}(t)$ (26) describes the availability state of nodes relevant to the quorum.

The system's transition to a new state under the action of the optimal control decision $A^*(t)$ is given by

$$\Omega(t+1) = \Phi(\Omega(t), A^*(t)). \quad (32)$$

where $\Phi(\cdot)$ is the system transition operator between discrete time steps.

Then the transition

$$\Omega(t) \rightarrow \Omega(t+1) \quad (33)$$

is considered admissible if the solution $A^*(t)$ belongs to the set $A^{\text{adm}}(t)$ (29).

Thus, relations (24)–(33) define the quorum component of the mathematical model and formalize the conditions for the admissibility of transitions of a blockchain-oriented distributed information system between states in a multi-cloud environment. Consequently, we present the mathematical model for managing computational resources and ensuring quorum in a blockchain-oriented distributed information system in a multi-cloud environment as a conditional minimization problem:

$$A^*(t) = \arg \min_{A_R(t)} J_R(t), \begin{cases} \sum_{\ell=1}^L \sum_{j=1}^{N_{c,\ell}} (\bar{r}_{c,\ell,j} + \bar{u}_{c,\ell,j}) + \bar{S}_c \leq \bar{R}_c; \\ \bar{r}_{c,\ell,j}^{\min} \leq \bar{r}_{c,\ell,j}(t+1) \leq \bar{r}_{c,\ell,j}^{\max}; \\ \chi_{c,\ell,j}(t+1) = 1; \\ v_{c,\ell,j} \in V^q; \\ \Psi(t+1) \geq \Psi^{\min}; \\ \Omega(t+1) = \Phi(\Omega(t), A^*(t)). \end{cases} \quad (34)$$

In Equation (34), the objective function $J_R(t)$ defines the criterion for minimizing the total unregulated resource deviation, while the set of constraints specifies the resource feasibility of control actions, the limits of permissible resource allocation, the conditions for stable node operation, quorum assurance, and the admissibility of system transitions between states.

5.4. Simulation study of the model and analysis of the results

A simulation study was conducted to verify the performance of the developed mathematical model for managing computing resources and ensuring quorum in a blockchain-oriented distributed information system within a multi-cloud environment. The study examined a discrete interval of 120 cycles, divided into three phases: nominal operation mode, resource shortage, and degradation of node availability. This division allowed for a sequential assessment of the model's behavior under conditions of normal

operation, local growth in resource demand, and deterioration of node availability relevant to quorum maintenance.

Figure 2 shows the trends of key normalized indicators of the simulation model: the aggregated system state vector $\bar{S}(t)$, the quorum assurance functional $\Psi(t)$, and the resource deviation functional $J_R(t)$. Normalization was applied to bring the indicators to a common scale and enable their comparative visual analysis within a single graph.

Figure 2 shows that in the initial simulation phase, the values of $\bar{S}(t)$ remain at a relatively low level, while $\Psi(t)$ is near the upper limit, corresponding to a stable quorum. In the resource shortage phase, an increase in $\bar{S}(t)$ is observed, along with the appearance of peaks in the $J_R(t)$ function, reflecting an increase in unregulated resource deviation as the load increases. In the final phase, corresponding to a degradation in availability, the level $\Psi(t)$ decreases, while the values of $\bar{S}(t)$ remain within the range characteristic of system operation following resource adaptation. This indicates that at this stage, maintaining a quorum of available nodes becomes the determining factor for the acceptability of decisions.

The summarized results of the simulation study by phase are presented in Table 1.

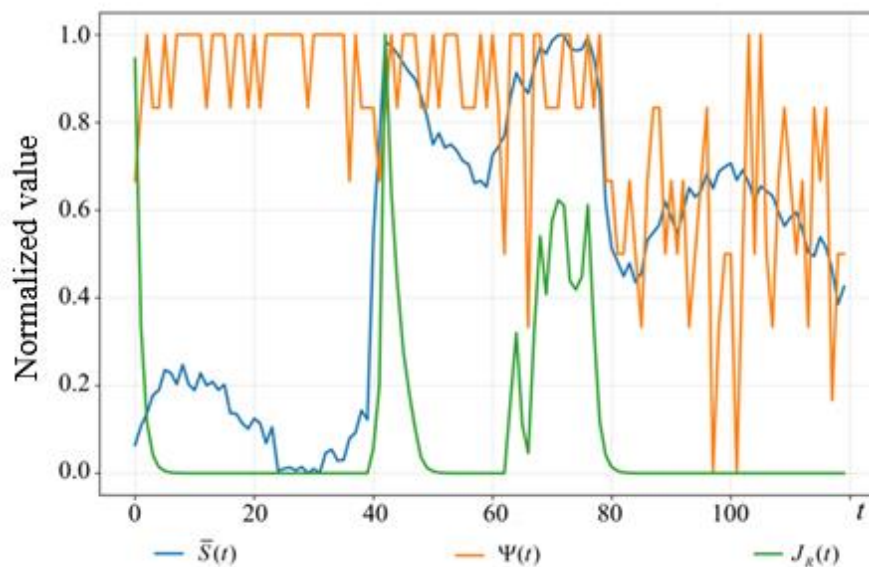


Fig. 2. Trends in key normalized indicators of the simulation model

Table 1. Generalized results of the simulation study of the model by scenario phases

Simulation phase	Number of cycles	$\bar{S}(t)$, average	$\bar{Q}(t)$, average	$\Psi(t)$, average	$J_R(t)$, average	Average total resource deficit	Proportion of allowed transitions
Nominal mode	40	0.531	0.958	8.625	0.0013	0.000	1.000
Resource deficit	39	0.823	0.926	8.333	0.0086	2.453	0.974
Degradation of availability	41	0.714	0.710	6.390	0.0001	0.000	0.805

An analysis of the data in Table 1 shows that, in nominal mode, the model operates under the most favorable conditions: the average value of $\bar{S}(t)$ is 0.531, the average value of $\Psi(t)$ is 8.625 with a minimum permissible quorum level of 6, and the proportion of valid transitions is 1.000. This confirms that, in the absence of resource shortages and with high node availability, the model generates control actions that do not violate either resource or quorum constraints.

During the resource shortage phase, the average value of the aggregate state $\bar{S}(t)$ increases to 0.823, and the average total resource shortage is 2.453. At the same time, the average value of the functional $J_R(t)$ increases to 0.0086. This means that under increased load, the model correctly detects an increase in resource deviation and shifts the decision-making process into active resource reallocation mode. At the same time, the average value of $\Psi(t)$ remains at 8.333, and the proportion of valid transitions is 0.974, indicating that a majority quorum is maintained in most cycles even in the presence of a local resource shortage.

The availability degradation phase is characterized by a different dominant factor. In the absence of an average resource shortage, the value of $J_R(t)$ practically approaches zero, while the average proportion of available nodes $\bar{Q}(t)$ decreases to 0.710, and the average value of $\Psi(t)$ decreases to 6.390 – that is, it approaches the quorum threshold. As a result, the proportion of valid transitions decreases to 0.805. The obtained result confirms that in this phase, it is precisely the quorum constraints defined by relations (28)–(30) that determine the validity of control actions much more strongly than the resource component of the model.

Overall, the results of the simulation study confirm the validity of the developed model and the correctness of the combination of its resource and quorum components. Equations (13)–(23) describe the redistribution of resources and minimize the resource deviation functional $J_R(t)$, while equations (24)–(33)

constrain the set of feasible solutions with conditions for quorum preservation and the admissibility of transitions between states. The obtained values of the indicators show that the model adequately responds to changes in the system's operating mode: in the resource shortage phase, it is sensitive to the deterioration of the resource state, and in the availability degradation phase, to the reduction in the number of nodes sufficient to maintain a quorum.

Thus, the simulation study confirmed that the developed mathematical model is suitable for describing the processes of managing computing resources and ensuring quorum in a blockchain-oriented distributed information system in a multi-cloud environment.

6. Discussion of the research results

The results obtained should be viewed as an advancement in approaches to describing multi-cloud applications, with a primary focus on architectural deployment models and requirements for multi-cloud interaction. Study [27] demonstrates that for multi-cloud applications, the key issues are component portability, the concurrent execution of services across multiple cloud environments, and the alignment of architectural solutions with infrastructure properties. The proposed model specifies these provisions at the level of mathematical description by separating the set of cloud environments, the set of architectural levels, the set of nodes, and the vectors of available and allocated resources. This approach allows for a transition from a general architectural representation of a multi-cloud system to a formalized scheme for managing its resource state.

Compared to approaches focused on cross-data-center virtual machine scheduling, the results obtained here have a broader scope of generalization. In particular, in [27], blockchain is used as a distributed ledger for exchanging status information between data centers during the scheduling and migration of virtual machines, which allowed for a reduction in the number of service messages and a reduction in communication latency. In the developed model, the blockchain-oriented nature of

the system is accounted for in a different context: control is linked to the generalized redistribution of heterogeneous resources among nodes at different architectural levels in a multi-cloud environment. This allows for the description of system states and control actions within a single formalized representation.

A concept that is conceptually similar is the approach to multi-cluster resource quota management in Kubernetes environments. The study [28] proposes a combination of predictive and reactive quota scaling based on Karmada, where decisions are made based on load monitoring and forecasting data. This approach is valuable from the perspective of practical implementation of dynamic resource reallocation; however, it is focused on quotas in a multi-cluster environment and does not account for the requirements for maintaining a quorum of nodes in a blockchain-oriented system. In the proposed model, the aggregated state vector $\bar{S}(t)$, defined by equation (3), is used as the basis for generating control actions according to equations (13)–(23), and the admissibility of such actions is subsequently verified with respect to quorum constraints. This constitutes the fundamental difference between the proposed approach and purely quota-based scaling.

An important area of related research is consensus-oriented orchestration of distributed services. In [29], it is shown that for distributed microservice environments, a decentralized approach to orchestration is appropriate, one that utilizes a consensus mechanism, dynamic leader election, and distributed logging. In this work, the main focus is on a mathematical model that combines the resource state of nodes, their availability, and the conditions for the system's transition between states. That is why the quorum is described through the set of nodes V^q , the availability state vector $\bar{Q}(t)$, the quorum satisfaction functional $\Psi(t)$, and the set of admissible control actions $A^{\text{adm}}(t)$. This framework creates a formal basis for the subsequent synthesis of control algorithms that are consistent with quorum requirements.

It should be noted that in current research on blockchain consensus, the primary focus is often on security, survivability, and resistance to attacks, as well as trade-offs between different classes of protocols. Source [30] systematizes issues of security and reliability of consensus mechanisms and demonstrates that the choice of protocol significantly influences the properties of a distributed system. The proposed model operates at a different level of abstraction and determines under

which resource conditions and with which composition of available nodes a control decision is permitted that does not violate quorum maintenance requirements. Thus, its advantage lies in combining the resource and quorum components within a single mathematical framework.

The results of the simulation study confirm the feasibility of this combination. As shown in Fig. 2, in nominal mode, the values of $\bar{S}(t)$ remain in the lower range, while $\Psi(t)$ remains near the upper limit, which corresponds to a stable quorum. During the resource shortage phase, an increase in $\bar{S}(t)$ and the appearance of peaks in $J_R(t)$ are observed, reflecting an increase in unregulated resource deviation as the load increases. During the availability degradation phase, the preservation of the quorum of available nodes becomes the determining factor for the admissibility of decisions, manifested in a decrease in $\Psi(t)$ to a level close to the threshold value of $\Psi^{\min}$.

The data in Table 1 quantitatively confirm this pattern. In the nominal mode, the proportion of admissible transitions is 1.000, the average value of $\Psi(t)$ is 8.625, and the average total resource deficit is zero. In the resource deficit phase, the average value of $\bar{S}(t)$ increases to 0.823, the average total resource deficit is 2.453, and the average value of $J_R(t)$ reaches 0.0086. At the same time, the average value of $\Psi(t)$ remains at 8.333, indicating that the quorum is maintained in most cycles. In the availability degradation phase, there is no average resource deficit; the value of $J_R(t)$ practically approaches zero, but the average value of $\bar{Q}(t)$ decreases to 0.710, and $\Psi(t)$ to 6.390, as a result of which the proportion of admissible transitions decreases to 0.805. This confirms that the model correctly distinguishes between situations dominated by resource constraints and situations in which quorum conditions become decisive.

The results obtained suggest that the proposed model adequately describes the process of managerial decision-making in a multi-cloud environment under various system operating modes. Its strength lies in the fact that the resource component, represented by equations (13)–(23), and the quorum component, represented by equations (24)–(33), operate in concert within a single mathematical framework. This is precisely what makes it possible to reject a resource-efficient

action if, after its execution, the quorum condition is violated.

A limitation of the current version is the abstract definition of the threshold value $\Psi^{\min}$ without reference to a specific consensus protocol. Further research should focus on parameterizing the model for protocols in the PBFT, Raft, or Tendermint families, as well as on a quantitative comparison of the quality of control decisions in scenarios with varying load intensities and different degrees of degradation in the set of available nodes.

7. Conclusions

As a result of this study, a mathematical model has been developed for the first time to manage computing resources and ensure quorum in a blockchain-oriented distributed information system in a multi-cloud environment. Compared to existing models, this model takes into account a set of cloud platforms from different vendors, the system's multi-level architecture, the aggregated resource state of nodes, their availability, and formalized quorum conditions, which made it possible to present the process of making management decisions regarding resource reallocation and maintaining system operability within a single mathematical framework.

Based on the research results, the multi-cloud deployment of components of a blockchain-oriented distributed information system was formalized based on the separate representation of a set of cloud environments, a set of architectural levels, a set of nodes, and vectors of available and allocated resources. This allowed for a correct description of the system as a control object in a multi-cloud environment and the specification of resource balance constraints for each cloud platform.

A mathematical description of the process of managing computing resources was developed based on the aggregated system state vector, node resource demand vectors, control actions, and conditions for permissible resource reallocation. This made it possible to formalize changes in the resource state of nodes under conditions of resource heterogeneity, dynamic load changes, permissible resource allocation limits, and stable node operation.

Formalized conditions for ensuring a quorum of available nodes and the admissibility of system transitions between states were determined, for which a subset of nodes relevant for maintaining the quorum, a state vector of their availability, a quorum assurance function, and a set of admissible control actions were introduced. This allowed for the integration of the resource and quorum components within a single mathematical control model.

The results of the simulation study confirmed the effectiveness of the proposed model under conditions of nominal operation, resource shortage, and degradation of node availability. It was established that in the event of an increase in load, the model correctly captures an increase in resource deviation, and in the event of a decrease in node availability, the condition for ensuring quorum becomes the determining constraint on the admissibility of control decisions.

The results obtained can be used for the further development of methods for optimizing resource management and ensuring quorum in blockchain-oriented distributed information systems in a multi-cloud environment.

Conflict of interest

The author declares that he has no conflicts of interest regarding this study, including financial, personal, or authorship-related conflicts, or any other conflicts that could influence the study and its results as presented in this article.

Funding

The study was conducted without financial support.

Data availability

The manuscript has no associated data.

Use of artificial intelligence

The author confirms that he did not use artificial intelligence technology in the creation of this work.

References

- Opoku, E., Okafor, M., Williams, M., Aribigbola, A. and Olaleye, A. (2024), "Enhancing Small and Medium-Sized Businesses Through Digitalization", *World Journal of Advanced Research and Reviews*, Vol. 23, No. 2, pp. 222–239. DOI: <https://doi.org/10.30574/wjarr.2024.23.2.2313>
- Kuchuk, H., Kosenko, N., Kuchuk, N., Gopejenko, V. and Kosenko, V. (2026), "Adaptive Resource Allocation Method for the Mobile Fog Layer of High-Density Industrial Internet of Things in Industry 5.0 Networks", *Innovative Technologies and Scientific Solutions for Industries*, Vol. 1, No. 35, pp. 65–78. DOI: <https://doi.org/10.30837/2522-9818.2026.1.065>
- Anthony Jnr, B. (2023), "Distributed Ledger and Decentralised Technology Adoption for Smart Digital Transition in Collaborative Enterprise", *Enterprise Information Systems*, Vol. 17, No. 4, Article 1989494. DOI: <https://doi.org/10.1080/17517575.2021.1989494>
- Pronchakov, Y., Prokhorov, O. and Fedorovich, O. (2022), "Concept of High-Tech Enterprise Development Management in the Context of Digital Transformation", *Computation*, Vol. 10, No. 7, Article 118. DOI: <https://doi.org/10.3390/computation10070118>
- Almufti, S.M. and Zeebaree, S.R.M. (2024), "Leveraging Distributed Systems for Fault-Tolerant Cloud Computing: A Review of Strategies and Frameworks", *Academic Journal of Nawroz University*, Vol. 13, No. 2, pp. 9–29. DOI: <https://doi.org/10.25007/ajnu.v13n2a2012>
- Horiunova, M.S. and Tkachov, V.M. (2025), "Metod Balansuvannia Elastychnykh Potokiv Danykh u Vkladenykh Virtualnykh Merezakh", In: *Problems of Informatization: Proceedings of the 13th International Scientific and Technical Conference*, Vol. 3, Section 4, NTU "KhPI", Kharkiv, pp. 25–26.
- Mamidala, J.V., Attipalli, A., Enokkaren, S.J., Bitkuri, V., Kendyala, R. and Kurma, J. (2023), "A Survey on Hybrid and Multi-Cloud Environments: Integration Strategies, Challenges, and Future Directions", *International Journal of Humanities and Information Technology*, Vol. 5, No. 2, pp. 53–65. <https://doi.org/10.21590/ijhit.05.02.08>
- Kuchuk, N., Kashkevich, S., Radchenko, V., Andrusenko, Y. and Kuchuk, H. (2024), "Applying Edge Computing in the Execution IoT Operative Transactions", *Advanced Information Systems*, Vol. 8, No. 4, pp. 49–59. DOI: <https://doi.org/10.20998/2522-9052.2024.4.07>
- Kuchuk, H., Kalinin, Y., Dotsenko, N., Chumachenko, I. and Pakhomov, Y. (2024), "Decomposition of Integrated High-Density IoT Data Flow", *Advanced Information Systems*, Vol. 8, No. 3, pp. 77–84. DOI: <https://doi.org/10.20998/2522-9052.2024.3.09>
- Khan, M.M.I. and Nencioni, G. (2023), "Resource Allocation in Networking and Computing Systems: A Security and Dependability Perspective", *IEEE Access*, Vol. 11, pp. 89433–89454. DOI: <https://doi.org/10.1109/ACCESS.2023.3306534>
- Andrusenko, Y. and Fesenko, T. (2024), "Analysis of the Main Sources of Uncertainty in the Parameters of the Computing Environment", In: *Current Directions of Development of Information and Communication Technologies and Control Tools: Proceedings of the 14th International Scientific and Technical Conference*, Vol. 1, Sections 1–2, NTU "KhPI", Kharkiv, p. 65.
- Gupta, S. (2025), "Hybrid Cloud Integration and Multicloud Deployments: A Comprehensive Review of Strategies, Challenges, and Best Practices", *International Journal of Advanced Research in Computer Science*, Vol. 16, No. 2, pp. 59–64. DOI: <https://doi.org/10.26483/ijarcs.v16i2.7233>
- Okorie, S.H. (2023), "Multi-Cloud Strategy for Enterprise Applications: Cost, Performance, and Resilience Considerations", *International Journal of Cloud Computing and Database Management*, Vol. 4, No. 1, pp. 63–73. DOI: <https://doi.org/10.33545/27075907.2023.v4.i1a.104>
- Frolov, D.Ye. (2025), "Model for Ensuring Fault Tolerance of Information Systems Based on Multi-Level Structures", *Visnyk of Kherson National Technical University*, Vol. 2, No. 3(94), pp. 474–483. DOI: <https://doi.org/10.35546/kntu2078-4481.2025.3.2.61>
- Ruban, I. and Tkachov, V. (2025), "Multilevel Model of an Information System on a Mobile Platform and Formalization of Its Survivability Criteria", *Visnyk of Kherson National Technical University*, Vol. 2, No. 3(94), pp. 399–409. DOI: <https://doi.org/10.35546/kntu2078-4481.2025.3.2.51>
- Tokariyev, V., Tkachov, V., Ilina, I. and Partyka, S. (2019), "Implementation of Combined Method in Constructing a Trajectory for Structure Reconfiguration of a Computer System with Reconstructible Structure and Programmable Logic", In: *Selected Papers of the XIX International Scientific and Practical Conference "Information Technologies and Security" (ITS 2019)*, CEUR Workshop Proceedings, Vol. 2577, pp. 71–81, available at: <https://ceur-ws.org/Vol-2577/paper7.pdf>
- Bayzid, L.H., Kar, T.S., Islam, M.T., Islam, M.S. and Ahmed, F. (2025), "Defending the Distributed Skies: A Comprehensive Literature Review of the Arena of Multi-Cloud Environment", *Future Internet*, Vol. 17, No. 12, Article 548. DOI: <https://doi.org/10.3390/fi17120548>

18. Habib, G., Sharma, S., Ibrahim, S., Ahmad, I., Qureshi, S. and Ishfaq, M. (2022), "Blockchain Technology: Benefits, Challenges, Applications, and Integration of Blockchain Technology with Cloud Computing", *Future Internet*, Vol. 14, No. 11, Article 341. DOI: <https://doi.org/10.3390/fi14110341>
19. Duran, E., Ozturk, C. and O'Sullivan, B. (2025), "Planning and Scheduling Shared Manufacturing Systems: Key Characteristics, Current Developments and Future Trends", *International Journal of Production Research*, Vol. 63, No. 13, pp. 4958–4990. DOI: <https://doi.org/10.1080/00207543.2024.2442549>
20. AlMuraytib, S., Alqurashi, L. and Snoussi, S. (2022), "Blockchain-Based Solutions for Cloud Computing Security: A Survey", In: *ICFNDS '22: Proceedings of the 6th International Conference on Future Networks & Distributed Systems*, pp. 338–342. DOI: <https://doi.org/10.1145/3584202.3584251>
21. Danish, S.M., Srivastava, G., Nourmohammadi, R., Ashraf, N., Ranjha, A. and Hameed, A. (2024), "Blockchain-as-a-Service: Architecture, Opportunities and Challenges", *IEEE Internet of Things Magazine*, Vol. 7, No. 6, pp. 52–57. DOI: <https://doi.org/10.1109/IOTM.001.2300199>
22. Pustovoitov, P.Y. and Klavdiiev, V.P. (2025), "Stochastic Model of Telecommunication Network Scalping Based on the Winer Process for Cloud Infrastructures", *Scientific Notes of Taurida National V.I. Vernadsky University. Series: Technical Sciences*, Vol. 1, No. 4, pp. 76–84. DOI: <https://doi.org/10.32782/2663-5941/2025.4.1/10>
23. Kudrynskyi, P., Zvenihorodskyi, O. and Bai, Y. (2025), "Models and Methods of Analysing Infrastructure Performance in Cloud Environments Based on Process Optimisation Methods", *Informatica*, Vol. 49, No. 12. DOI: <https://doi.org/10.31449/inf.v49i12.8933>
24. Bugriy, A.M., Sorobey, B.V., Polozov, D.M. and Volk, D.M. (2025), "Metody Upravlinnia Resursamy v Rozpodilenykh Systemakh Khmarnykh Obchyslen", In: *Problems of Informatization: Proceedings of the 13th International Scientific and Technical Conference*, Vol. 4, Sections 5–6, NTU "KhPI", Kharkiv, pp. 12–13.
25. Chepurna, I. (2025), "Model of Computational Resource Metric Collection for a Blockchain-Oriented Information System in a Cloud Environment", *Visnyk of Kherson National Technical University*, Vol. 3, No. 4(95), pp. 276–284. DOI: <https://doi.org/10.35546/kntu2078-4481.2025.4.3.32>
26. Alonso, J., Orue-Echevarria, L., Casola, V., Torre, A.I., Huarte, M., Osaba, E. and Lobo, J.L. (2023), "Understanding the Challenges and Novel Architectural Models of Multi-Cloud Native Applications: A Systematic Literature Review", *Journal of Cloud Computing*, Vol. 12, Article 6. DOI: <https://doi.org/10.1186/s13677-022-00367-6>
27. Altahat, M.A., Daradkeh, T. and Agarwal, A. (2025), "Virtual Machine Scheduling and Migration Management Across Multi-Cloud Data Centers: Blockchain-Based Versus Centralized Frameworks", *Journal of Cloud Computing*, Vol. 14, Article 1. DOI: <https://doi.org/10.1186/s13677-024-00724-7>
28. Tran, M.-N. and Kim, Y. (2025), "Hybrid Resource Quota Scaling for Kubernetes-Based Edge Computing Systems", *Electronics*, Vol. 14, No. 16, Article 3308. <https://doi.org/10.3390/electronics14163308>
29. Morabito, G., Ficara, A., Celesti, A., Villari, M. and Fazio, M. (2025), "Consensus-Based Distributed Orchestration Framework for Microservices in Edge Computing Clusters", *Future Generation Computer Systems*, Vol. 176, Article 108221. DOI: <https://doi.org/10.1016/j.future.2025.108221>
30. Bao, Q., Li, B., Hu, T. and Sun, X. (2023), "A Survey of Blockchain Consensus Safety and Security: State-of-the-Art, Challenges, and Future Work", *Journal of Systems and Software*, Vol. 196, Article 111555. DOI: <https://doi.org/10.1016/j.jss.2022.111555>

Received (Надійшла) 20.04.2026

Accepted for publication (Прийнята до друку) 21.05.2026

Publication date (Дата публікації) 27.06.2026

Відомості про авторів / About the Authors

Чепурна Ірина Сергіївна – Харківський національний університет радіоелектроніки, асистент кафедри електронних обчислювальних машин, Харків, Україна.

Iryna Chepurna – Kharkiv National University of Radio Electronics, Assistant of the Electronic Computers Department, Kharkiv, Ukraine;

e-mail: iryna.chepurna@nure.ua

ORCID ID: <https://orcid.org/0009-0008-2442-6221>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=60153860100>

МОДЕЛЬ УПРАВЛІННЯ РЕСУРСАМИ ТА КВОРУМОМ БЛОКЧЕЙН-ОРІЄНТОВАНОЇ РОЗПОДІЛЕНОЇ ІНФОРМАЦІЙНОЇ СИСТЕМИ В МУЛЬТИХМАРНОМУ СЕРЕДОВИЩІ

Предметом дослідження є математична модель управління обчислювальними ресурсами й забезпечення кворуму блокчейн-орієнтованої розподіленої інформаційної системи в мультихмарному середовищі. **Метою** є розроблення теоретико-множинної моделі, що формалізує мультихмарне розміщення компонентів блокчейн-орієнтованої розподіленої інформаційної системи, зважає на її багаторівневу архітектуру, гетерогенність доступних ресурсів, агрегований ресурсний стан вузлів, їх доступність і умови забезпечення кворуму під час функціонування за змінного навантаження. **Методи.** Щоб досягти окреслену мету, впроваджено системний аналіз для формалізації структури системи й взаємозв'язків між її компонентами; математичне моделювання – для побудови формального опису мультихмарного середовища, ресурсного стану вузлів, керівних дій і кворумних обмежень; статистичний аналіз – для оброблення, агрегації та нормалізації телеметричних показників; теорія прийняття рішень – для формування умов допустимого перерозподілу ресурсів; імітаційне моделювання – для перевірки працездатності розробленої моделі за умов змінного навантаження, ресурсного дефіциту й варіативної доступності вузлів. **Результати.** Уперше розроблено формалізовану теоретико-множинну модель управління обчислювальними ресурсами й забезпечення кворуму блокчейн-орієнтованої розподіленої інформаційної системи в мультихмарному середовищі, яка, на відміну від наявних підходів, поєднує в межах єдиної уніфікованого середовища подання мультихмарного середовища як сукупності хмарних платформ різних вендорів, багаторівневої архітектури системи, агрегованого ресурсного стану вузлів, умов допустимого перерозподілу ресурсів і заданих вимог до забезпечення кворуму. Запропоновано множини хмарних середовищ, архітектурних рівнів і вузлів системи, упроваджено вектори доступних і виділених ресурсів, агрегований вектор стану системи, функціонал ресурсного відхилення, функціонал кворумної забезпеченості та множину допустимих керівних дій. Імітаційне дослідження підтвердило працездатність моделі в умовах номінального функціонування, ресурсного дефіциту й деградації доступності вузлів. Це дало змогу подати процес прийняття рішень щодо перерозподілу ресурсів і підтримки працездатності системи в межах єдиного математичного подання. **Висновки.** Запропонована модель сприяє формалізованому опису процесів управління обчислювальними ресурсами й забезпеченню кворуму блокчейн-орієнтованої розподіленої інформаційної системи в мультихмарному середовищі та може бути використана як основа для подальшого розроблення методів оптимізації керівних рішень і експериментального дослідження ефективності функціонування системи за умов динамічної зміни навантаження й часткової деградації вузлів.

Ключові слова: математична модель; управління ресурсами; кворум; блокчейн; інформаційна система; хмара.

Бібліографічні описи / Bibliographic descriptions

Чепурна І. С. Модель управління ресурсами та кворумом блокчейн-орієнтованої розподіленої інформаційної системи в мультихмарному середовищі. *Сучасний стан наукових досліджень та технологій в промисловості*. 2026. № 2 (36). С. 185–199. DOI: <https://doi.org/10.30837/2522-9818.2026.2.185>

Chepurna, I. (2026), "Model of resource and quorum management for a blockchain-oriented distributed information system in a multi-cloud environment", *Innovative Technologies and Scientific Solutions for Industries*, No. 2 (36), P. 185–199. DOI: <https://doi.org/10.30837/2522-9818.2026.2.185>

Galyna Nazarova, Alina Demianenko, Nikita Nazarov, Oleg Ivanisov, Vladyslav Ostapchuk

DETERMINANTS OF HUMAN POTENTIAL DEVELOPMENT IN THE CONDITIONS OF STRUCTURAL TRANSFORMATION OF INDUSTRY

The structural transformation of industry in Europe is accompanied by profound changes in the structure of employment. In these conditions, human development determines the ability of the economy to adapt, while demographic and migration processes form the underlying preconditions for industrial change. The research aims to determine the relationship between the HDI and a set of demographic and migration indicators. Their impact on key components of the HDI remain relevant in the face of structural transformation of industry. The following research methods were used in the article: descriptive statistics; statistical methods of determining the normality of data; correlation analysis was performed to determine the strength and direction of correlation; panel regression modeling allowed to identify significant results in the quantitative assessment of the influence of factors. The results of research obtained reveal statistically relationships between demographic and migration indicators and each dimension of the HDI. The correlation analysis identified the direction, strength, and closeness of associations between key variables. The results revealed a moderate degree of positive correlation between the Age dependency ratio and two dimensions of human development: Long and healthy life and Decent standard of living. The analysis showed a strong negative impact of the Crude Death Rate on almost all HDI dimensions. The Net migration rate has a significant positive impact on two areas of human development: Knowledge and Decent standard of living, as well as the HDI as a whole. Net migration rate and crude rate of immigration proved to be of particular interest for future research as compensating factors, the increase of which can accelerate HDI progress. The results of the correlation analysis showed the influence of other indicators on the HDI dimensions, but such influences were weak, nonlinear, or context-dependent. The regression analysis conducted for each dimension of human development provides practical insights of high policy relevance. The constructed models offer an empirical basis for identifying key influencing factors and enable forecasting, as well as the development of evidence-based strategies and targeted measures aimed at strengthening the resilience of human capital and human development under conditions of structural transformation of industry.

Keywords: human development; human development index; human capital; demographic dynamics; migration risks; structural transformation; industry.

1. Introduction

The structural transformation of industry in Europe is one of the key trends in today's economy. It is manifested in a gradual shift away from traditional industries and the simultaneous development of high-tech sectors, digital technologies, and green energy. These changes are shaping a new structure of employment and production processes, placing high demands on workers' qualifications and competencies.

Human development in such conditions becomes critically important: the population's ability to adapt to changes in the labor market, master new professions, and undergo retraining determines the economic and social stability of the state. The structural transformation of industry directly affects people's access to jobs, their income levels, and prospects for professional growth.

At the same time, demographic changes and migration processes serve as the root causes of these shifts: population aging, internal migration from depressed industrial regions to technological hubs,

and the recruitment of skilled professionals from abroad are shaping a new labor force structure. It is precisely the combination of structural industrial transformation and demographic-migration processes that defines the current challenges and opportunities for human development in Europe.

2. Analysis of the literature and problem definition

The concept of human development has not lost its relevance over the past decades; on the contrary, it has gained even greater significance, depth, and a shift in emphasis, especially in conditions of high turbulence and unpredictability. Thus, contemporary societal development is characterized by non-linearity, uneven progress, and an integrated complex network of social, economic, political, environmental, and other factors; therefore, understanding the connections and interrelationships of these factors is extremely important for developing security policies and fostering the strategic sustainability of human development. This is precisely

why the total number of studies dedicated to identifying the connection between socio-economic processes has significantly increased in recent years, including research on the impact of demographic and migration processes on human development.

V. Lutz and S. Kumar K. S. (2013) examined fundamental questions regarding the acceleration of human development through the lens of demographic transition processes [1]. As a result of in-depth research, the scholars obtained valuable findings regarding the conditions that can influence the acceleration of certain components of human development. These conditions include a decrease in the dependency ratio of the young population (a certain number of adults must support a smaller number of children) and improvements in educational infrastructure. Investment in education is a key factor influencing changes in population structure: a decline in fertility and the emergence of the demographic dividend effect, which has a strong impact on economic growth. The established patterns and causal relationships between demographic aspects, the population's level of education, and economic growth have broadened our understanding of the human development process, confirmed the primary causality of human capital as a driver of human potential development, and served as a starting point for future research [1–4].

The results of the analysis of scientific studies conducted by researchers confirm the importance of demographic factors for human development. Hartgen and Vollmer (2014) examine the relationship between human development and fertility using a modified HDI [5]. The analysis demonstrates that the so-called fertility transition effect is not stable and varies depending on the HDI level, emphasizing that the observed links between demographic and development indicators may be unstable and highly sensitive to methodological choices.

Our attention was drawn to a study by Kozlovsky S., Pasichny M., Lavrov R., Ivanyuta N., and Nepitalyuk A. (2020). The researchers' aim was to identify relationships between changes in demographic variables (growth rates of the working-age population and average life expectancy) and economic growth. The researchers established a close link between life expectancy and nominal GDP per capita. Life expectancy was significantly higher in developed countries than in developing countries; the researchers also found that growth in the working-age population drastically

reduced production dynamics, but this relationship was not stable [6,7].

Bafail and Alamudi (2023) determined that the following indicators – the share of skilled labor, R&D expenditure, tourism expenditure, and trade expenditure – are the most influential determinants of the HDI, using multiple regression to rank their relative importance [8]. Sing et al. (2025) confirmed that GDP per capita, average years of schooling, and infant mortality are key determinants of HDI, explaining approximately 95% of its variation [9].

Although most empirical studies focus on assessing the impact of demographic dynamics on economic and social outcomes, other researchers (Errero, Martinez, and Villar, 2019) emphasize that demographic composition itself can distort development measurements [10–13]. Research by these scholars has shown that the traditional HDI (developed by UNDP experts) does not fully reflect the actual level of human well-being, as it does not account for the age structure of the population.

Over the past few decades, researchers worldwide have been studying issues of demographic and human development and their interrelationships, and each year new valuable scientific findings emerge, the implementation and practical application of which can expand human potential and influence the pace and scale of human development. Previous studies in this field are characterized by diverse regional samples. When examining the impact and interrelationship between human development and demographic processes, some researchers focus on specific aspects or within the framework of their own understanding of human development, rather than the generally accepted methodology for measuring human development levels developed by UNDP experts. In our study, we examine the impact of demographic and migration factors on the HDI in 12 EU countries with specific demographic and/or migration challenges. This will provide a clearer and more vivid picture of the depth of the impact, revealing regional imbalances and sources of vulnerability. In addition, the extent and nature of these factors' impact are examined for each structural component of the HDI according to the UNDP measurement model.

Thus, this article examines the theoretical, methodological, and practical aspects of human development shaped by demographic and migration factors. The aim of the study is to determine the relationship between the HDI and a set of demographic and migration indicators. Their impact on key HDI

components remains relevant in the context of growing instability and uncertainty.

The results of our study have both theoretical and practical value. The theoretical contribution lies in identifying consistent patterns, substantiating a model of causal relationships between processes, and establishing a theoretical framework for forecasting each component of the HDI. The practical value lies in the developed regression model, which allows predicting how the level of each HDI component will change under the influence of demographic factors (changes in the age structure of the population, changes in mortality and fertility rates, etc.) and migration factors (changes in immigration and emigration levels).

3. Purpose and Objectives of the Study

The purpose of the study is to determine the relationship between the HDI and a set of demographic and migration indicators, as their impact on key HDI dimensions remains a relevant issue in the context of industrial structural transformation.

The objectives of the study are as follows:

- 1) to analyze differences in human capital development, demographic, and migration situations among the countries under study;
- 2) to justify the feasibility and relevance of using the selected indicators for further research;
- 3) to identify promising and significant results in the quantitative assessment of the impact of factors on each component of the HDI.

4. Methodology

The study is based on assessing the impact of a set of independent variables (demographic and migration-related) on each component (element) of the HDI. To achieve the set goal, multiple regression analysis methods were applied.

The Human Development Index is a complex composite indicator that combines components reflecting various aspects of human development ("Long and Healthy Life", "Knowledge", "Decent Standard of Living"). Each of these components is shaped by various socioeconomic, demographic, migration, and other factors, which makes the construction of a single regression model for the aggregate HDI methodologically limited.

First, demographic and migration factors exert a similar nature and direction of influence on the aggregate HDI. Assessing the impact of demographic and migration factors on the aggregate HDI leads to an averaging of dependencies and a loss of information about the effects on individual components.

Second, HDI components differ in measurement scale and time sensitivity.

Third, HDI components have varying significance, although the UNDP model's HDI calculation methodology involves normalization and aggregation to ensure their comparability and equal weight in the overall index. However, in models with an aggregated composite indicator, it is impossible to unambiguously determine which specific component is responsible for a change in the HDI. Modeling individual HDI components makes it possible to identify a set of demographic and migration factors that significantly influence each HDI component, which enhances the analytical significance of the results and their practical innovation for the development of socio-demographic policy.

Fourth, this approach partially mitigates the problem of endogeneity. Demographic and migration indicators and the HDI may be in a reciprocal causal relationship, yet endogeneity manifests differently for individual components.

Thus, the methodology of studying individual HDI components aligns with the structural logic of human development, within which demographic and migration factors influence both individual components ("Long and Healthy Life", "Knowledge", "Decent Standard of Living") and the overall HDI.

The chosen methodological approach corresponds to the study's objective and the scientific tasks set forth and involves testing the hypotheses formulated by the authors.

Hypothesis 1. The impact of demographic and migration factors on HDI components is statistically heterogeneous, indicating varying intensities of their effects on individual components.

Hypothesis 2. The impact of demographic and migration factors on each component of the HDI can be assessed using a separate regression model, which allows for identifying the specific nature of the impact of individual independent variables (demographic and migration) on each component.

The study utilized general scientific and specialized research methods. Statistical and mathematical methods:

graphical, correlation, and panel regression analysis to identify relationships and construct a mathematical model. Descriptive statistics provided an initial understanding of differences in human capital development, demographic and migration situations across countries, and allowed us to draw conclusions regarding the potential use of certain indicators for further research. A correlation analysis was conducted to determine the strength and direction of the correlation. The use of panel regression modeling allowed us to identify promising and significant results in the quantitative assessment of the impact of factors on each component of the HDI.

Selection of indicators and data analysis. Statistical data on the HDI and primary indicators were obtained from the UNDP Human Development Data Center. The selection of independent variables for multiple regression analysis was determined by several factors. First, we used statistical data on demographic and migration indicators for EU countries obtained from the official Eurostat website, which ensured comparability of indicators across countries and throughout the study period. Second, we selected indicators that demonstrated changes over time, allowing us to assess the dynamics of the processes under study. Third, the statistical data on the selected demographic and migration indicators most clearly reflect the key demographic and migration issues under study, and their significance is characteristic of most of the countries included in the analysis. Thus, the selection of variables was based simultaneously on data availability, their dynamics, and theoretical relevance to the evaluation process.

Our study utilized modern software tools for data processing and analysis, such as SPSS Statistics and Microsoft Excel.

The study analyzed statistical data from 12 EU countries (Bulgaria, the Czech Republic, Denmark, France, Germany, Greece, Italy, Latvia, Lithuania, Poland, Romania, Sweden) that experience varying degrees of demographic or migration challenges but have a high or above-average HDI. It makes sense to include Ukraine in the main study sample – a country characterized by a unique case of an exogenous factor causing a sharp demographic and migration shock, in addition to other demographic and migration challenges (low birth rates, high mortality, demographic distortion, and labor emigration). However, since Ukraine is not currently an EU member, and given that the publication of certain demographic and migration data is restricted due to wartime conditions and security concerns, including Ukraine in the model is not currently possible. Nevertheless, the model developed using data from the aforementioned European countries can be adapted for forecasting and strategic planning purposes in Ukraine.

5. Research Results

The primary indicator of human development is the HDI, developed by UNDP experts. Although the HDI has a number of strengths and weaknesses regarding its set of indicators, the objective determination of its level, and its sensitivity to internal and external changes, the HDI remains the key measure of countries' achievements in creating conditions that promote the expansion of human capabilities. The HDI is the indicator of global well-being in the modern world and, at the same time, a tool for shaping and reforming the policies of many states [14]. The results of the analysis of global HDI trends reveal interesting patterns. One such trend is the general, steady increase in the level of human development worldwide (Fig. 1).

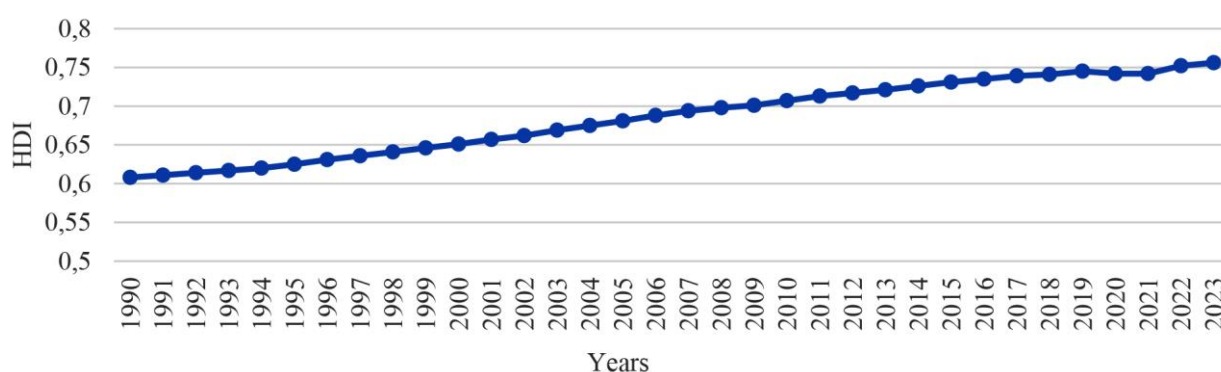


Fig. 1. Trends in the global HDI [15]

The steady rise in global human development can be attributed to a number of factors: the long-term trend toward improved living standards; the irreversibility of certain processes; the structure of the index; and so on. However, the pace of this growth varies by region and over the study periods, depending on the extent and intensity of the influence of internal and external factors. Thus, there are a number of processes capable of

influencing fluctuations in the index: stimulating or hindering development; the impact of these processes on the HDI can be either direct or indirect.

The results of the data analysis demonstrate a similar trend of rising HDI levels in the EU countries included in the analysis. It should be noted that in all countries included in the analysis, the HDI level is higher than the global average (Fig. 2).

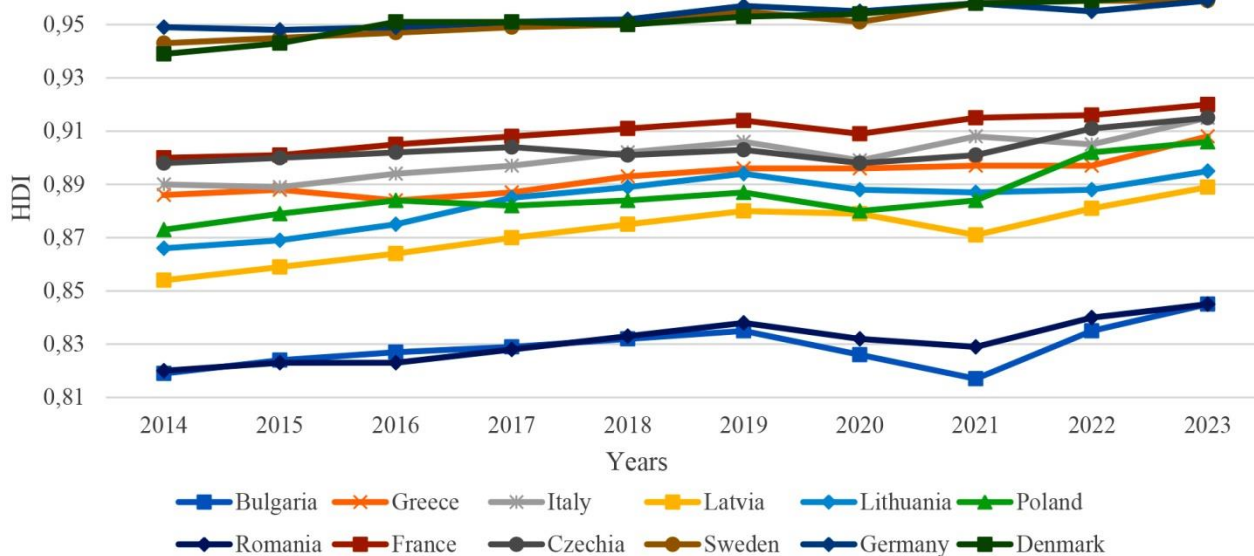


Fig. 2. HDI trends in some European countries [15]

As shown in Fig. 2, the HDI trajectory of almost every country is unstable and subject to slight fluctuations, and the dynamics reveal a slowdown in the growth rates of the HDI of the studied countries compared to the global HDI. Such fluctuations and the slowdown in HDI growth are caused by a series of processes that can stimulate, reduce, or hinder development, and the impact of these processes on the HDI can be direct or indirect, nonlinear, or context-dependent. One group of such factors consists of demographic and migration processes. Their distinctive feature lies in their contribution to the formation of human capital: demographic and migration processes determine the scale and quality of human resources and alter the trajectory of development in the long term.

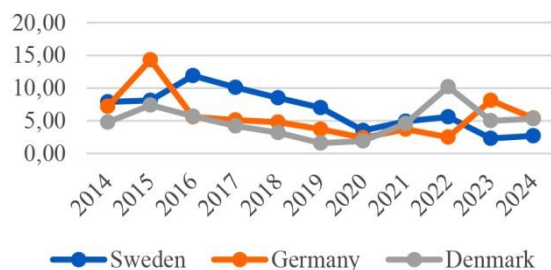
It should be noted that the EU countries selected for the study face demographic and/or migration challenges to varying degrees. Therefore, the methodology for selecting countries was based on their analytical grouping according to the intensity of these factors. Three groups were conditionally identified. The first group included countries with significant negative demographic and migration problems (Bulgaria, Lithuania, Latvia,

and Italy). The demographic situation in these countries was characterized by a combination of low birth rates, accelerated population aging, and depopulation. The second group of countries – Poland, Romania, and Greece – was characterized by moderate demographic and migration challenges. The third group of countries included the Czech Republic, Denmark, Germany, France, and Sweden. The migration situation in these countries was relatively stable; there was an issue of population aging, which was also not critical and was mitigated by compensating factors. In our opinion, including countries with varying degrees of demographic and migration challenges will ensure objective research results. An analysis of key trends was conducted in the aforementioned countries. As a result of the comparative analysis, key patterns and cause-and-effect relationships were identified.

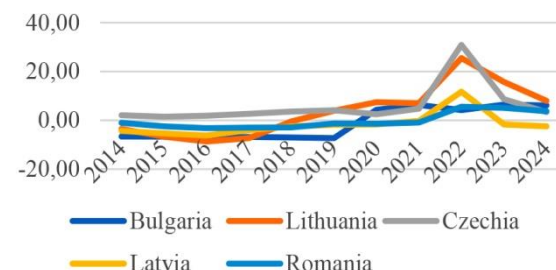
Based on a thorough analysis of the dynamics of the Age Dependency Ratio, it can be concluded that the indicator increased by an average of 10–12%, reflecting the intensification of the population aging trend and a decline in fertility. The largest increases were observed in Poland, Bulgaria, and Romania,

indicating an acceleration of demographic dependency and a decline in the working-age population. France, Sweden, and Denmark showed relatively stable

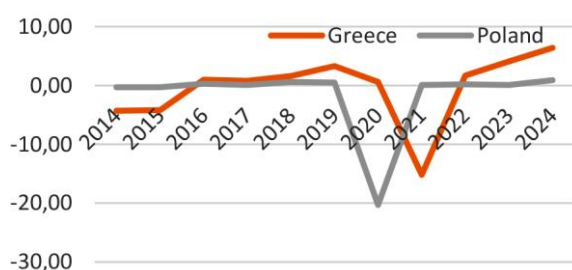
but high rates (around 60–62%). These rates are typical of a demographic model with balanced life expectancy and fertility.



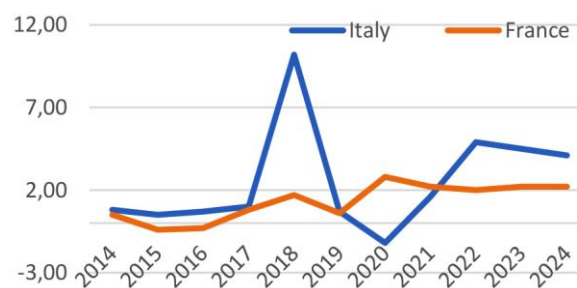
Group A. Sustained, rapid population growth



Group B. A sharp increase in migration levels after 2019



Group C: Sharp drops and virtually no fluctuations



Group D. Strong positive fluctuations

Fig. 3. Total adjusted net migration in selected EU countries, 2014–2024 [16]

Grouping countries and identifying trajectories provide a deeper understanding of the patterns and underlying factors driving these trends. This will enable an objective and high-quality implementation of the correlation analysis procedure.

After identifying the demographic and migration trends of the countries included in the analysis, the

authors conducted a correlation analysis. This stage of the study involved identifying and interpreting potential interdependencies between variables, providing a conceptual and empirical basis for further quantitative modeling of the impact of demographic and migration factors on human development (Table 1).

Table 1. Descriptive statistics for the primary HDI indicators and demographic and migration indicators

Indicator	Average meaning	Median	Mode	Minimum value	Maximum value	Standard deviation	Assimetry	Excess
Y	0,90	0,90	–	0,82	0,96	0,041	–0,175	–0,741
Y_1	0,91	0,92	0,97	0,79	0,98	0,051	–0,285	–1,412
Y_2	0,82	0,84	–	0,68	0,95	0,073	–0,120	–1,086
Y_3	0,93	0,92	–	0,78	1,16	0,081	0,568	–0,076
Y_4	0,92	0,92	–	0,82	1,00	0,046	–0,007	–1,089
X_1	55,12	55,75	57,40	42,70	62,20	4,025	–0,627	0,344
X_2	31,29	31,60	–	21,20	37,80	3,320	–0,466	0,395
X_3	1,58	1,60	1,70	1,20	2,00	0,198	–0,147	–0,860
X_4	12,12	11,35	9,10	8,40	22,90	2,694	1,000	1,181
X_5	1,97	1,70	–	–20,30	30,90	6,347	0,735	5,232
X_6	6,70	5,87	–	1,01	19,44	3,621	1,024	1,286
X_7	9,31	7,69	–	2,81	33,04	5,324	1,859	4,751

The HDI shows a high average score of 0.90, indicating a high level of human development with low dispersion and a nearly symmetrical distribution, suggesting relatively uniform development outcomes across countries. The "Life Expectancy at Birth" indicator and the indicators of the "Knowledge" sub-index reflect similar stability. The "Decent Standard of Living" sub-index, measured using Gross National Income per capita, is characterized by greater variability, reflecting significant differences in income levels among countries.

A preliminary analysis of the skewness and kurtosis coefficients showed that most indicators had values close to a normal distribution. However, some variables were characterized by pronounced positive skewness and high kurtosis, indicating a deviation from the norm.

To more accurately verify the data distribution, additional normality tests were conducted using the Shapiro–Wilk (W) and Kolmogorov–Smirnov (K–S) tests (Table 2).

Table 2. Normal distribution test results

Indicator	Shapiro–Wilk W	Sig.(SW)	K–S D	Sig. (K–S)	Interpretation
Y	0,938	0	0,110	<0,15	Abnormal distribution
Y_1	0,899	0	0,181	<0,01	Abnormal distribution
Y_2	0,944	0	0,143	<0,05	Abnormal distribution
Y_3	0,941	0	0,152	<0,01	Abnormal distribution
Y_4	0,960	0	0,097	> 0,20	Abnormal distribution
X_1	0,970	0,01	0,078	> 0,20	Close to normal
X_2	0,982	0,11	0,054	> 0,20	Normal distribution
X_3	0,954	0	0,1385	<0,05	Abnormal distribution
X_4	0,923	0	0,119	<0,10	Abnormal distribution
X_5	0,916	0	0,103	<0,20	Abnormal distribution
X_6	0,928	0	0,105	<0,15	Abnormal distribution
X_7	0,833	0	0,146	<0,05	Abnormal distribution

According to the results of the normality test, most variables showed significant or severe deviations from the normal distribution, indicating that their data were non-normally distributed. Only variables X_1 and X_2 exhibited approximately normal and normal distributions, respectively. Thus, the overall data structure did not conform to a normal distribution, so it was decided to apply nonparametric statistical methods in the subsequent analysis.

In this regard, nonparametric methods were applied, specifically Spearman's rank correlation coefficient (ρ) and Kendall's tau coefficient (τ), which do not require an assumption of data normality (Table 2).

Table 3 presents the results of the correlation analysis between dependent and independent variables. The simultaneous use of Spearman's ρ and Kendall's τ coefficients is due to the complex and difficult-to-formalize nature of the analyzed relationships. Their simultaneous application allows for distinguishing between structural and proportional dependencies and provides a more

robust interpretation of correlations in the context of heterogeneous data.

The results obtained indicate a diversified system of interrelationships characterized by both direct and inverse dependencies of varying strength. The consistency between Spearman's and Kendall's coefficients confirms the stability of the identified relationships and suggests that these relationships are not random.

The age dependency coefficient has a direct positive impact on certain sub-indices of human development, namely Life Expectancy at Birth and Gross National Income per Capita.

The old-age dependency ratio maintains a moderate positive relationship with the HDI overall and with life expectancy at birth, which may indicate that its influence is conditional or secondary, mediated through other components of the system.

The total fertility rate shows mostly minor but negative correlations. We hypothesize that such a correlation is nonlinear or context-dependent, a view supported by a number of studies by researchers.

Table 3. Correlation matrix of demographic and migration indicators and HDI measures

Indicators	Y	Y_1	Y_2	Y_3	Y_4
X_1	$\rho = 0,432$ $\tau = 0,309$	$\rho = 0,538$ $\tau = 0,399$	$\rho = -0,195$ $\tau = -0,093$	$\rho = 0,238$ $\tau = 0,178$	$\rho = 0,397$ $\tau = 0,292$
X_2	$\rho = 0,256$ $\tau = 0,194$	$\rho = 0,392$ $\tau = 0,280$	$\rho = -0,231$ $\tau = -0,121$	$\rho = 0,249$ $\tau = 0,173$	$\rho = 0,150$ $\tau = 0,130$
X_3	$\rho = -0,005$ $\tau = -0,047$	$\rho = -0,076$ $\tau = -0,064$	$\rho = -0,086$ $\tau = -0,110$	$\rho = -0,303$ $\tau = -0,218$	$\rho = 0,132$ $\tau = 0,073$
X_4	$\rho = -0,726$ $\tau = -0,518$	$\rho = -0,818$ $\tau = -0,619$	$\rho = 0,098$ $\tau = 0,102$	$\rho = -0,413$ $\tau = -0,285$	$\rho = -0,671$ $\tau = -0,471$
X_5	$\rho = 0,611$ $\tau = 0,456$	$\rho = 0,383$ $\tau = 0,262$	$\rho = 0,262$ $\tau = 0,174$	$\rho = 0,418$ $\tau = 0,287$	$\rho = 0,617$ $\tau = 0,436$
X_6	$\rho = -0,209$ $\tau = -0,121$	$\rho = -0,381$ $\tau = -0,267$	$\rho = 0,225$ $\tau = 0,158$	$\rho = 0,109$ $\tau = 0,074$	$\rho = -0,160$ $\tau = -0,065$
X_7	$\rho = 0,379$ $\tau = 0,251$	$\rho = 0,043$ $\tau = 0,028$	$\rho = 0,423$ $\tau = 0,289$	$\rho = 0,414$ $\tau = 0,265$	$\rho = 0,428$ $\tau = 0,294$

Notes:

ρ – Spearman's rank correlation coefficient; τ – Kendall's rank correlation coefficient.

Among the most pronounced negative relationships is the link between the mortality rate and life expectancy at birth. Spearman's correlation coefficient indicates a very high, nearly perfect relationship. Thus, the high mortality rate observed in the countries included in the analysis significantly lowers the life expectancy at birth. This relationship warrants attention, as the trend is specific and may result from changes in the age structure of the population due to the emigration of young people or high mortality rates among this demographic. This implies the emergence of an economic effect and the existence of another strong inverse relationship between mortality rates and gross national income per capita.

Net migration showed a strong positive correlation with both the HDI overall and some of its components: expected years of schooling and gross national income per capita, underscoring that in the countries studied, net migration acts as a compensatory (stimulating) factor. There is a weak negative correlation between the total emigration rate and life expectancy at birth.

The total immigration rate has a moderate positive correlation with three HDI indicators. Thus, an increase in the total immigration rate leads to an increase in the average and expected years of schooling. In our view, this relationship may be driven by a chain of interrelated processes. This increase may be caused by an influx of children and youth, which stimulates demand for education and necessitates the expansion of educational programs, as well as the prolongation of educational trajectories due to adaptation. Another

scenario for the increase in the duration of schooling due to immigration could be institutional integration measures. Thus, educational immigration or relocation for the purpose of active integration into a new society increases the duration of schooling.

In light of the above, it should be noted that the results of the correlation analysis confirm asymmetric and disproportionate interdependencies, where certain factors (Age Dependency Ratio, Net Migration Rate, and Immigration Rate) have a dominant influence, while others (Mortality Coefficient and Emigration Coefficient) create a balancing effect. This fact underscores the need for a multidimensional model capable of accounting for both reinforcing and inhibiting dynamics for each HDI component.

The identified correlation structure provides an empirical basis for the next stage of analysis. Given the difficult-to-formalize and complex nature of human development, as well as the multidimensional structure of the HDI and the heterogeneity of the mechanisms underlying its formation, it is appropriate to use a differentiated approach. This is precisely what necessitates the construction of several regression models focused on individual structural components of the HDI.

Constructing three models with sub-indices as dependent variables allows us to:

- avoid the "averaging" effect of influences;
- increase the accuracy of coefficient estimates;
- reduce the risk of methodological distortions associated with excessive data aggregation.

Each of the three regression models reflects a separate aspect of human development and allows for the conceptualization of the mechanisms underlying its formation:

- the first model characterizes the influence of demographic and migration factors on the longevity component of human development;
- the second model reflects the dependence of the educational component on the quality of demographic and migration policies;
- the third model focuses on productivity potential and the economic realization of human potential.

This approach deepens the theoretical understanding of human development as a multidimensional process and confirms that its components have different sensitivities to demographic and migration determinants.

In our study, nonparametric data were used for modeling. A check of the sample for multicollinearity revealed a high degree of dependence among the independent variables. High VIF values were observed for the Age Dependency Ratio and the Elderly Dependency Ratio, indicating a significant level of multicollinearity. Both indicators reflect related aspects, leading to information redundancy in the models (Table 4).

Table 4. Assessment of multicollinearity

Indicator	VIF	Interpretation
X_1	5.474	High multicollinearity
X_2	5.396	High multicollinearity
X_3	1.965	Acceptable level
X_4	1.589	Acceptable level
X_5	4.188	Acceptable level
X_6	2.386	Acceptable level
X_7	4.221	Acceptable level

The net migration rate and the immigration rate also show marginal multicollinearity, whereas the net migration rate, the emigration rate, and the immigration rate reflect similar processes. Therefore, some independent variables were excluded to remove redundant information from the model. A subsequent VIF calculation confirmed a significant reduction in multicollinearity and stabilization of the coefficient estimates. Thus, the model retains only those variables (Age-Dependency Coefficient, Total Fertility Rate, Mortality Rate, Emigration Rate, and Immigration Rate) that provide unique information, thereby enhancing the reliability of the results' interpretation.

Each dependent variable – each HDI sub-index – is modeled as a function of demographic and migration variables using rank-order multiple regression. This approach allows for the capture of nonlinear and distribution-invariant relationships between demographic and migration factors (Age-Dependency Coefficient, Total Fertility Rate, Mortality Rate, Emigration Rate, and Immigration Rate) and the human development sub-indices (Y_L , Y_k , Y_D). The construction of three models allows for a comparative analysis of the strength,

direction, and statistical significance of the impact of demographic and migration factors on various components of human development. This makes it possible to:

- identify key determinants for each sub-index;
- identify asymmetries in the effects;
- form a more detailed empirical basis for further research.

Model 1. Demographic and migration determinants of the "Long and Healthy Life" sub-index. The regression model demonstrated adequacy. The coefficient of determination was 0.86, and the adjusted coefficient of determination was 0.85, indicating that the model explains a significant portion of the variance in the dependent variable. The high F -value ($5, 114$) = 176.65 indicates that the entire model is statistically significant and adequately explains the variance of the dependent variable. The SEE was 13, indicating acceptable accuracy. It is worth noting that the General Immigration Level indicator was excluded from the model due to its statistical insignificance ($p > 0.05$), which improved the interpretability of the regression equation.

$$Y_L = 106,2 + 0,379 \times X_1 - 0,311 \times X_3 - 0,73 \times X_4 - 0,089 \times X_6 + \varepsilon. \quad (1)$$

Model 2. Demographic and migration determinants of the "Knowledge" sub-index. The multiple regression model, which includes three variables, is statistically significant, as $F(5, 114) = 24.34$, $p < 0.001$. $R^2 = 0.38$, adjusted $R^2 = 0.37$.

$$Y_k = 59,99 - 0,18 \times X_3 - 0,28 \times X_4 + 0,515 \times X_7 + \varepsilon. \quad (2)$$

Model 3. Demographic and migration determinants of the "Decent Standard of Living" subindex. Four significant variables were used to construct the model (the total fertility rate was excluded due to its non-significant p -value). The results of the analysis showed that the model explains approximately 65%

$$Y_d = 69,9 + 0,126 \times X_1 - 0,55 \times X_4 - 0,2 \times X_6 + 0,47 \times X_7 + \varepsilon. \quad (3)$$

Thus, the data obtained from the model-building process demonstrate that the models are sufficiently adequate for further application.

From a practical standpoint, the proposed approach offers significant advantages for the development and implementation of public policy. Focusing on sub-indices allows for the formulation of targeted policy decisions aimed at specific problem areas of human development, rather than universal but less effective measures. The results obtained can be used:

- in the formulation of educational and social policy;
- in the healthcare sector;
- in the development of strategies for human capital development and employment.

Thus, the use of three regression models instead of a single generalized one not only eliminates the methodological limitations of existing approaches but also ensures a deeper theoretical understanding and high practical significance of the research results.

6. Discussion of Results

This section interprets the results of the correlation and panel regression analyses, highlighting the main relationships between explanatory variables and dependent variables. The results obtained and the conclusions drawn by the authors are analyzed and substantiated in light of the existing scientific literature on the subject and their significance for policy and practice.

Correlation analysis showed that the positive relationship between the demographic dependency ratio and life expectancy at birth, as well as per capita GNI,

These results indicate the acceptable quality of the model and its ability to explain 38% of the variation in the dependent variable, which is sufficient for this HDI measure. To optimize the model, the "Overall Emigration Rate" variable was excluded, as it had no significant effect on Y_k .

of the variation in the dependent variable, and the adjusted R^2 confirmed the model's adequacy. $F(5, 114) = 54.39$ indicates the overall significance of the model. The standard errors of the coefficients are at an acceptable level.

may reflect the characteristics of the demographic transition. In more economically developed countries, high life expectancy leads to an increase in the proportion of the elderly population, which raises the overall dependency ratio. At the same time, high GDP per capita provides the resources to support dependent age groups thanks to well-developed health care, education, and social security systems. Thus, a high age dependency ratio in this case is not a consequence of economic weakness but an indicator of demographic aging.

As correlation analysis has shown, the fertility rate has a negligible impact on HDI components, and this impact depends on the context, as confirmed by studies conducted by researchers who have examined this issue. However, a weak negative correlation is observed between the fertility rate and expected years of schooling. Therefore, we hypothesize that an excessively low fertility rate creates a demographic imbalance that indirectly reduces the level of expected years of schooling. Thus, the conclusion of Hartgen and Wallmer (2014) regarding the complex bidirectional interaction between fertility and HDI components is confirmed, according to which both excessive and insufficient fertility can hinder human development [7].

Mortality rates have a strong negative impact on life expectancy and health, as well as on a decent standard of living, indicating the inhibitory effect of sociodemographic characteristics. High mortality rates shorten life expectancy, as the proportion of young and working-age people in the population structure changes for various reasons: this may be due to their outflow from the country, mortality, or a decline in the youth (or working-age) segment

of the population as a result of demographic decline. At the same time, this reduces the size and productivity of the labor force, increases social expenditures, and decreases national income. Therefore, a high mortality rate is an indicator of a low level of development of the healthcare system and the country's limited economic potential.

The net migration rate has a strong positive impact on all human development sub-indices. Migration is typically accompanied by an influx of economically active, well-educated individuals of working age, thereby strengthening the country's human capital. This contributes to increased labor productivity, the development of innovation, and the improvement of health and education systems. In addition, migrants often bring additional financial resources, such as investments, taxes, and remittances, which stimulate economic growth. Thus, a positive net migration rate is a factor in demographic and economic renewal, contributing to increased well-being and socio-economic stability in the country.

The overall level of emigration has a negative impact on the HDI. However, this impact is weak and insignificant. It is important to highlight the positive average impact of immigration on the levels of certain HDI components. This is particularly true for Average Years of Schooling, Expected Years of Schooling, and Gross National Income per Capita.

High levels of immigration, particularly of young and economically active individuals, lead to an increase in the labor force and higher labor productivity. This contributes to economic growth and increases national income. Furthermore, the influx of intellectual capital improves the quality of health care and education systems, which can enhance living conditions and the level of social security. As a result, high levels of immigration are a favorable factor for both education and a country's economic development.

The empirical results of our study further confirm and refine previous research, highlighting the significant impact of demographic factors on the HDI. As shown by studies by Lutz (2013) and Lutz et al. (2018), population structure, fertility, and mortality correlate with HDI components, which is consistent with our data on the positive impact of the working-age population share and the negative impact of high mortality rates [1; 2]. Empirical data from Kozlovsky et al. (2020), Engelhardt-Wolffler et al. (2018), and Herrero et al. (2019) show that age structure alters HDI dynamics, which is consistent with our study results [4]. The nonlinear effects described

by Hartgen and Volmer (2012, 2014) explain the inverse effect of fertility on the HDI in countries at different levels of development [6, 7].

The models developed by the authors, based on data from European countries, have high analytical value, as they reflect generalized patterns of interrelationships between demographic, social, and economic factors of development. Ukraine belongs to the European socio-demographic space, so its human development dynamics are subject to similar structural trends – population aging, declining fertility, the impact of migration, and educational transformation. This makes the use of European models for forecasting methodologically justified, but with due consideration for Ukrainian conditions and the demographic and migration shock of 2022. Therefore, an adapted model, built on European patterns but parameterized using Ukrainian data, is capable of providing realistic forecasts of the components of human development in Ukraine.

7. Conclusions

The study results showed that the impact of demographic and migration factors on HDI components is heterogeneous. For the "Long and Healthy Life" dimension, the mortality rate is the determining factor, while a higher proportion of the working-age population and rising birth rates have a positive impact. For the "Knowledge" dimension, the level of immigration plays a dominant role (young people likely predominate in the immigration structure); The total fertility rate has a minor but significant impact on the reduction in the duration of education due to demographic aging and the redistribution of resources; it is worth emphasizing that demographic indicators have a context-dependent relationship. For the "Decent Standard of Living" measurement model, the most important factors are Mortality Rate and Net Migration Flow, which determine economic dynamics and development potential. Thus, the results obtained deepen the understanding of the interrelationship between demographic processes, migration, and the level of human development. The constructed panel regression models allow for a quantitative assessment of the contribution of each demographic and migration indicator to the formation of HDI components, providing a new empirical toolkit for analysis. Thus, the study contributes to the development of human development theory by elucidating the role of structural demographic changes as determinants of socio-economic dynamics.

Scientific novelty of the study

The methodological approach to identifying the determinants of human capital development in the context of industrial structural transformation has been further developed. This approach is based on the construction of three regression models, where the dependent variables are the sub-indices of the Human Development Index, and the independent variables are demographic and migration factors represented by corresponding indicators. Unlike existing approaches, the proposed methodology focuses not on the analysis of the overall composite indicator (HDI), but on its structural components, which allows for the elimination of a number of limitations inherent in traditional models, enhances the analytical and statistical significance of the results, and ensures their practical value for the formulation and implementation of countries' socio-economic policies.

Directions for future research

Prospects for further research include the development of strategic programs and practical recommendations for demographic and migration policy in the context of improving human development, taking into account the identified dependencies. It also includes the development of methodological recommendations

for forecasting the HDI, taking into account the impact of demographic and migration aspects. In our opinion, the study of the impact of the quantitative and qualitative characteristics of human capital migration as a compensatory factor for a country's human development in the context of global and local upheavals is of significant interest.

Conflict of interest

The authors declare that they have no conflicts of interest, including financial, personal, authorial, or any other nature, that could influence the research or the results published in this article.

Funding

The study was conducted without financial support.

Data availability

The manuscript has no associated data.

Use of artificial intelligence

The authors confirm that they did not use artificial intelligence technologies to write this paper.

References

1. Lutz, W. and Samir, K.C. (2013), *Demography and Human Development: Education and Population Projections*, United Nations Development Programme, New York, available at: <https://hdr.undp.org/system/files/documents/hdr01304lutzkc.pdf> (accessed 16 April 2026).
2. Lutz, W., Goujon, A., Samir, K.C., Stonawski, M. and Stilianakis, N. (2018), *Demographic and Human Capital Scenarios for the 21st Century: 2018 Assessment for 201 Countries*, IIASA, Laxenburg, 594 p. DOI: <https://doi.org/10.2760/835878>
3. Caselles, A., Soler Fernandez, D., Sanz, M. and Mico, J. (2014), "Simulating Demography and Human Development Dynamics", *Cybernetics and Systems*, Vol. 45, No. 6, pp. 465–485. DOI: <https://doi.org/10.1080/01969722.2014.929347>
4. Kozlovskiy, S., Pasichnyi, M., Lavrov, R., Ivanyuta, N. and Nepyaliuk, A. (2020), "An Empirical Study of the Effects of Demographic Factors on Economic Growth in Advanced and Developing Countries", *Comparative Economic Research*, Vol. 23, No. 4, pp. 45–67. DOI: <https://doi.org/10.18778/1508-2008.23.27>
5. Engelhardt-Wölfler, H., Schulz, F. and Büyükköçeci, Z. (2018), *Demographic and Human Development in the Middle East and North Africa, Population and Family Studies*, Vol. 2, University of Bamberg Press, Bamberg, 88 p. DOI: <https://doi.org/10.20378/irbo-50993>
6. Stafeckis, N. and Berzins, M. (2026), "Exploring the Complex Interplay of Demographic and Socioeconomic Dynamics in Urban Shrinkage of Latvian Mono-Towns", *Urban Science*, Vol. 10, No. 4, Article 211. DOI: <https://doi.org/10.3390/urbansci10040211>
7. Harttgen, K. and Vollmer, S. (2014), "A Reversal in the Relationship of Human Development with Fertility?", *Demography*, Vol. 51, No. 1, pp. 173–184. DOI: <https://doi.org/10.1007/s13524-013-0252-y>
8. Tutar, H., Mutlu, H.T. and Streimikiene, D. (2025), "Rethinking Sustainable Development: Multidimensional Comparisons Across Countries and Regions", *Sustainable Development*, Vol. 34, Suppl. 1, pp. 183–196. DOI: <https://doi.org/10.1002/sd.70157>
9. Singh, K., Sharma, P. and Gupta, R. (2025), "Determinants of Human Development Index: A Regression Analysis of Economic and Social Indicators", *Asian Journal of Economics, Business and Accounting*, Vol. 25, No. 11, pp. 1–12. DOI: <https://doi.org/10.9734/ajeba/2025/v25i11630>

10. Herrero, C., Martínez, R. and Villar, A. (2019), "Population Structure and the Human Development Index", *Social Indicators Research*, Vol. 141, pp. 731–763. DOI: <https://doi.org/10.1007/s11205-018-1852-0>
11. Hair, J.F., Black, W.C., Babin, B.J. and Anderson, R.E. (2019), *Multivariate Data Analysis*, 8th ed., Cengage Learning, Andover, available at: <https://www.cengage.uk/c/multivariate-data-analysis-8e-hair-babin-anderson-black/9781473756540/?searchIsbn=9781473756540> (accessed 16 April 2026).
12. Vorobiov, V., Cherechin, O., Romaniv, I., Bereza, V., Cherepaniak, V. and Cherepaniak, A. (2023), "Digitalization of the IT Project Management Process", *Academic Visions*, Vol. 26, available at: <https://www.academy-vision.org/index.php/av/article/view/1206> (accessed 16 April 2026).
13. Nazarova, G., Jaworska, M., Nazarov, N., Dybach, I. and Demianenko, A. (2019), "Improvement of the Method Measuring the Security of Human Development: Case of Ukraine", *Problems and Perspectives in Management*, Vol. 17, No. 4, pp. 226–238. DOI: [https://doi.org/10.21511/ppm.17\(4\).2019.19](https://doi.org/10.21511/ppm.17(4).2019.19)
14. United Nations Development Programme (2023), *Human Development Report 2023/2024: Technical Notes*, available at: https://hdr.undp.org/sites/default/files/2023-24_HDR/hdr2023-24_technical_notes.pdf (accessed 16 April 2026).
15. United Nations Development Programme (n.d.), *Human Development Index – Data Center*, available at: <https://hdr.undp.org/data-center/human-development-index#/indicies/HDI> (accessed 16 April 2026).
16. Eurostat (n.d.), *Eurostat Database*, available at: <https://ec.europa.eu/eurostat/data/database> (accessed 16 April 2026).

Received (Надійшла) 07.12.2025

Accepted for publication (Прийнята до друку) 21.05.2026

Publication date (Дата публікації) 27.06.2026

Відомості про авторів / About the Authors

Назарова Галина Валентинівна – доктор економічних наук, професор, Харківський національний економічний університет імені Семена Кузнеця, завідувач кафедри соціальної економіки, Харків, Україна;

Galyna Nazarova – Doctor of Economic Sciences, Professor, Simon Kuznets Kharkiv National University of Economics, Head of the Social Economics Department, Kharkiv, Ukraine;

e-mail: gnazarova.ua@gmail.com

ORCID ID: <https://orcid.org/0000-0003-4893-5406>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=56677560100>

Дем'яненко Аліна Анатоліївна – доктор філософії, Харківський національний економічний університет імені Семена Кузнеця, викладач кафедри соціальної економіки, Харків, Україна;

Alina Demianenko – PhD, Simon Kuznets Kharkiv National University of Economics, Lecturer of the Social Economics Department, Kharkiv, Ukraine;

e-mail: alina.demianenko@hneu.net; daa22@ukr.net

ORCID ID: <https://orcid.org/0000-0003-0654-698X>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=57212685435>

Назаров Нікіта Костянтинівич – доктор економічних наук, професор, Харківський національний економічний університет імені Семена Кузнеця, професор кафедри менеджменту, бізнесу і адміністрування, Харків, Україна;

Nikita Nazarov – Doctor of Economic Sciences, Professor, Simon Kuznets Kharkiv National University of Economics, Professor of the Management, Business and Administration Department, Kharkiv, Ukraine;

e-mail: nikita.nazarov@hneu.net

ORCID ID: <https://orcid.org/0000-0001-8762-2248>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=56992815100>

Іванісов Олег Вікторович – кандидат економічних наук, доцент, Харківський національний економічний університет імені Семена Кузнеця, доцент кафедри соціальної економіки, Харків, Україна;

Oleg Ivanisov – Candidate of Economic Sciences, Associate Professor, Simon Kuznets Kharkiv National University of Economics, Associate Professor of the Social Economics Department, Kharkiv, Ukraine;

e-mail: ivanisovoleg@ukr.net

ORCID ID: <https://orcid.org/0000-0001-5757-824X>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=57226544614>

Остапчук Владислав Валерійович – Харківський національний економічний університет імені Семена Кузнеця, здобувач вищої освіти ступеня доктора філософії кафедри соціальної економіки, Харків, Україна;

Vladyslav Ostapchuk – Simon Kuznets Kharkiv National University of Economics, Postgraduate Student of the Social Economics Department, Kharkiv, Ukraine;

e-mail: vladyslav.os106@gmail.com

ORCID ID: <https://orcid.org/0009-0001-6122-7765>

ДЕТЕРМІНАНТИ РОЗВИТКУ ЛЮДСЬКОГО ПОТЕНЦІАЛУ В УМОВАХ СТРУКТУРНОЇ ТРАНСФОРМАЦІЇ ПРОМИСЛОВОСТІ

Структурна трансформація промисловості в Європі супроводжується глибокими змінами у структурі зайнятості. За цих умов людський розвиток визначає здатність економіки до адаптації, тоді як демографічні та міграційні процеси формують глибинні передумови промислових змін. Метою дослідження є визначення взаємозв'язку між ІЛР та комплексом демографічних та міграційних показників, оскільки їх вплив на ключові виміри ІЛР залишаються актуальним питанням в умовах структурної трансформації промисловості. У статті використовувалися такі методи дослідження: описова статистика; статистичні методи визначення нормальності даних; кореляційний аналіз – для визначення сили та напрямку кореляції; панельне регресійне моделювання дозволило виявити значущі результати в кількісній оцінці впливу факторів. Отримані результати дослідження розкривають статистично значущі зв'язки між демографічними та міграційними показниками та кожним виміром ІЛР. Кореляційний аналіз визначив напрямок, силу та тісноту зв'язків між ключовими змінними. Результати виявили помірний ступінь позитивної кореляції між коефіцієнтом вікової залежності та двома вимірами людського розвитку: Довгим та здоровим життям та Гідним рівнем життя. Аналіз показав сильний негативний вплив коефіцієнта смертності майже на всі виміри ІЛР. Чистий коефіцієнт міграції має значний позитивний вплив на два виміри людського розвитку: Знання та Гідний рівень життя, а також на ІЛР в цілому. Чистий коефіцієнт міграції та загальний коефіцієнт імміграції стали предметом аналітичного інтересу для майбутніх досліджень як компенсуючі фактори, збільшення яких може пришвидшити прогрес ІЛР. Результати кореляційного аналізу дозволили виявити вплив інших показників на виміри ІЛР, але такий вплив був слабким, нелінійним або контекстно-залежним. Регресійний аналіз, проведений для кожного виміру людського розвитку, дав змогу зробити висновки, що мають високу практичну цінність. Побудовані моделі створили емпіричну основу для визначення ключових факторів впливу та дозволили проводити прогнозування, а також розробляти стратегії спрямовані на зміцнення стійкості людського капіталу та розвитку людського потенціалу в умовах структурної трансформації промисловості.

Ключові слова: розвиток людського потенціалу; індекс розвитку людського потенціалу; людський капітал; демографічна динаміка; міграційні ризики; структурна трансформація; промисловість.

Бібліографічні описи / Bibliographic descriptions

Назарова Г. В., Дем'яненко А. А., Назаров Н. К., Іванісов О. В., Остапчук В. В. Детермінанти розвитку людського потенціалу в умовах структурної трансформації промисловості. *Сучасний стан наукових досліджень та технологій в промисловості*. 2026. № 2 (36). С. 200–213. DOI: <https://doi.org/10.30837/2522-9818.2026.2.200>

Nazarova, G., Demianenko, A., Nazarov, N., Ivanisov, O., Ostapchuk, V. (2026), "Determinants of human potential development in the conditions of structural transformation of industry", *Innovative Technologies and Scientific Solutions for Industries*, No. 2 (36), P. 200–213. DOI: <https://doi.org/10.30837/2522-9818.2026.2.200>

Igor Nevliudov, Vladyslav Yevsieiev, Svitlana Maksymova, Oleksandr Pashchenko, Viktor Kosenko

DEVELOPMENT OF A MODEL FOR SYNCHRONIZING THE OPERATION OF COLLABORATIVE ROBOTIC MANIPULATORS IN HRC SCENARIOS

This article examines the problem of synchronizing the operation of a group of collaborative robots in a shared workspace with humans, in light of modern requirements for human-centered manufacturing and the Industry 5.0 concept. **The relevance of the research** stems from the need to ensure coordinated movement of multiple robots while adhering to strict safety constraints in HRC scenarios, where traditional centralized approaches do not guarantee sufficient reliability and scalability. **The objective of this work** is to develop a mathematical model for the synchronized control of collaborative manipulators, taking into account mutual coordination, tracking of a shared trajectory, and active avoidance of dangerous proximity to humans and between robots. **The subject of the study** is decentralized synchronization laws in the problem space with projection into joint space and the use of potential safety fields. **The work uses methods** of mathematical modeling of manipulator dynamics based on Euler–Lagrange equations, consensus control methods, pseudo-inversion of the Jacobian matrix, fourth-order Runge–Kutta numerical integration methods, and methods for analyzing safety metrics. **The objectives of the study** are to formalize the laws of consensus control in the problem space, integrate safety potential fields, and perform numerical validation of the model for a group of manipulators in a shared workspace. **Modeling results** for a group of three two-link planar manipulators showed the formation of coordinated trajectories with a reduction in characteristic oscillations and stabilization of dynamics; however, the average tracking error of the common trajectory remains at 0.38–0.40 m. An analysis of minimum robot-to-robot and robot-to-human distances confirmed the effectiveness of potential barriers in steady-state operation, but revealed dangerous gaps in transitional sections. **It is concluded** that the proposed model ensures stable synchronization and a basic level of safety, but to guarantee compliance with safety margins throughout the entire time interval, it is advisable to switch to rigid barrier constraints such as CBF and task-priority control schemes.

Keywords: collaborative robotic manipulators, motion synchronization, HRC scenarios, decentralized control, consensus algorithm, safety potential fields, Jacobian pseudoinverse, Industry 5.0.

1. Introduction

In the current context of Industry 5.0, a human-centered approach is emerging as a key driver of technological evolution, under which robotic systems must not only increase productivity but also ensure safety, adaptability, and flexible interaction with humans [1–3]. Collaborative robots are increasingly being integrated into shared workspaces, where there is a need for precise synchronization of their actions in the presence of an operator. Under these conditions, traditional centralized control algorithms prove insufficiently effective due to limited scalability, data exchange delays, and difficulties in ensuring safety [4–7]. Therefore, decentralized synchronization models that combine coordinated motion control, adaptive compensation laws, and integrated mechanisms for preventing dangerous collisions are becoming particularly relevant [8–9]. The relevance of this research is also driven by the need to transition from "soft" heuristic approaches to formalized mathematical models capable of ensuring predictable system dynamics in complex Human–Robot Collaboration (HRC)

scenarios [10–13]. The realities of modern manufacturing require the simultaneous achievement of high synchronization accuracy, disturbance resistance, and guaranteed safety constraints. This is precisely why the development of models for synchronized control of collaborative robots in a shared space with humans is a timely, scientifically sound, and practically significant task.

2. Analysis of last achievements and publications

In their paper, I. Rybalskii, K. Kruusamäe, A.K. Singh, and S. Schlund propose an AR interface to improve the safety of operator interaction with industrial robots, where augmented reality is used to visualize danger zones, motion trajectories, and warnings, which statistically reduces the number of dangerous encounters in production scenarios [14]. This solution enables the integration of human-centered visual feedback tools into HRC systems and can be used in our study as a basis for overlaying dynamic AR prompts regarding safe distances and restricted zones onto the synchronized trajectories of manipulators.

In their study, D. Kóczi and J. Sárosi conducted a systematic review of safety engineering methods for humanoid robots in everyday life, developed a risk classification, reviewed standards, approaches to hazard assessment, and strategies for designing safe behavior [15]. The proposed approach makes it possible to link the mathematical synchronization model with functional safety requirements and emergency scenarios, so the results can be used to formalize robot-human distance constraints and motion acceptability criteria in our model.

In P. Pawlewski's article on the scientific and practical challenges of remanufacturing process modeling, the problems of reliable simulation of complex production flows are examined, and a set of functions, requirements, and architectural solutions is proposed for developing new approaches to modeling and digital twins of remanufacturing [16]. This solution makes it possible to align macro-level production models with the micro-level of robot control, and from the perspective of our research, its results can be used to justify the need to integrate a synchronization model with the digital twin of a conveyor or assembly line.

In the article by H. Ye, Y. Song, J. Lam, and P.A. Ioannou, a decentralized control law with a specified convergence time is developed for "hand-fingers" systems performing grasping and manipulation tasks, where each subsystem controller ensures the achievement of target trajectories within a limited time using Lyapunov stabilization methods and barrier functions [17]. This approach makes it possible to ensure predictable convergence of multi-link robots with minimal communication requirements; therefore, elements of this theory are suitable for use in our model to construct local joint synchronization laws, however, direct adaptation requires rethinking due to a different problem formulation (HRC in shared space, rather than merely object manipulation).

In the work by M. Zhang, J. Guo, and Z. Zhang propose a distributed slack-barrier recurrent neural network for coordinated motion of a set of redundant manipulators in an obstacle-filled environment, where the optimization problem with safety constraints is implemented as a dynamic NN system that avoids collisions and maintains coordination [18]. This solution makes it possible to encapsulate collision constraints and redundancy within the control dynamics themselves; therefore, its ideas are consistent with our potential field approach and can be used for the further development

of a neural network variant of the synchronization model, although full integration will require simplifying the computational complexity for real-time operation.

In a study by S. Agarwal and J. Alonso-Mora, S. Sun developed a decentralized real-time planning method for multi-UAV cooperative manipulation based on imitation learning, where local "student" policies learn to imitate a centralized planner under conditions of partial observability and the absence of inter-agent communication [19]. The proposed approach makes it possible to achieve near-centralized planning quality in a decentralized architecture, which is important for scaling multi-robot systems, and in the context of our work, this can be used as a conceptual basis for training local synchronization laws for manipulators based on reference centralized trajectories.

In the article by D. Ma, C. Zhang, and Q. Xu, and G. Zhou, an approach to proactive control of manipulation operations in an Industry 5.0 HRC assembly environment is considered, based on the fusion of large-scale models of the production system and small-scale models of specific robots, which ensures adaptation to the context and line load [20]. This solution makes it possible to combine line-level planning with local manipulator control, and thus directly supports the idea of integrating our synchronization model into a multi-level control architecture, where synchronized working trajectories are coordinated with demand forecasts and environmental changes.

In the work by P. Ding, J. Zhang, P. Zheng, P. Zhang, B. Fei, and Z. Xu, a human motion prediction network enriched with scenario features is proposed, which accounts for environmental dynamics and the diversity of operator motion patterns for proactive collaboration in personalized assembly tasks [21]. The proposed method makes it possible to significantly improve the accuracy of short-term predictions of human trajectories, which in our study is critically important for the correct construction of potential safety fields and the adaptive modification of synchronized robot trajectories; accordingly, the results of this work should be used as one of the sources for operator position prediction models.

In the article by W. W. Loy, J. Donovan, M. Rittenbruch, and M. Belek, Fialho, and Teixeira, an AR-oriented HRC system is investigated for research assembly tasks, where a designer can view, modify, and execute a cobot's actions in real time via an AR interface, receiving spatially constrained visual feedback [22]. This solution makes it possible to coordinate human

creative actions and robotic manipulations in "open" environments, and in the context of our synchronization model, the results can be used to design an interface layer that displays the robot's current and predicted synchronized trajectories and safety zones to support the operator's situational awareness.

In the work by S.M. Mizanoor Rahman, an approach to evaluation and benchmark metrics for collaborative object manipulation tasks in HRC is proposed, where a set of indicators for productivity, ergonomics, safety, and interaction quality is considered, along with test scenarios for comparing different algorithms [23]. This solution enables standardized evaluation of HRC systems and is directly relevant to our task, since the developed synchronization model requires quantitative quality criteria (tracking error, minimum distances, motion smoothness); so the proposed metrics can be adapted to validate our approach in numerical experiments.

The article by P. Lin, N. Zeng, Q. Li, and K. Nübel provides an overview of HRC safety in construction, systematically outlining typical hazards, regulatory requirements, technical risk control measures, and promising directions for the development of standards and safety sensor systems [24]. This review allows us to generalize the requirements for safe human-robot interaction in complex, loosely structured environments, and from the perspective of our research, it supports the justification for the need for embedded safety modules (potential fields, barrier conditions) within the synchronization model, although the specifics of construction sites differ from laboratory HRC scenarios involving manipulators.

Taken together, the analyzed publications demonstrate a steady trend toward a transition from isolated robotic systems to human-centric HRC architectures with decentralized control, human behavior prediction, AR support, and formalized safety and productivity metrics, confirming both the scientific and practical relevance of developing a synchronization model for collaborative robot manipulators in HRC scenarios capable of integrating these areas into a single mathematically sound approach.

The aim of this article is to develop a mathematical model for the synchronized control of collaborative manipulators, taking into account mutual coordination, tracking of a common trajectory, and active avoidance of dangerous proximity to humans and between robots.

3. Brief Explanation of Parameters and Variables

In this section, a brief explanation of parameters and variables is presented before the mathematical model to ensure unambiguous interpretation of the introduced notation and to simplify the understanding of subsequent derivations. This approach allows the reader to focus on the logic of the synchronization model's construction without being overwhelmed by formulas, and also enhances the reproducibility and correctness of the numerical implementation of the proposed method.

Indices and dimensions: $i, j \in \{1, \dots, N\}$ – robot indices; $N \in \{2, 3, 4\}$ – number of robots in the group; n_i – number of robot i joints; d – dimension of the workspace (3 for position, 6 for position+orientation).

State variables: $q_i(t) \in \mathbb{R}^{n_i}$ – angular coordinates of the joints; $\dot{q}_i(t) \in \mathbb{R}^{n_i}$ – angular velocities of the joints; $\ddot{q}_i(t) \in \mathbb{R}^{n_i}$ – angular accelerations

Kinematics and Geometry: $x_i = f_i(q_i) \in \mathbb{R}^d$ – current position/orientation of the robot's i end-effector in global coordinates; $\dot{x}_i = J_i(q_i)\dot{q}_i$ – velocity of the end-effector in global coordinates; $J_i(q_i) \in \mathbb{R}^{d \times n_i}$ – Jacobian matrix; $J_i^\dagger(q_i) \in \mathbb{R}^{n_i \times d}$ – pseudo-inverse Jacobian, for finding the desired joint accelerations that realize the desired acceleration in the workspace.

Dynamics: $M_i(q_i) \in \mathbb{R}^{n_i \times n_i}$ – mass-inertia matrix; $C_i(q_i, \dot{q}_i)\dot{q}_i \in \mathbb{R}^{n_i}$ – Coriolis and centrifugal moments; $g_i(q_i) \in \mathbb{R}^{n_i}$ – gravitational moments; $\tau_i \in \mathbb{R}^{n_i}$ – local control moment (the specific robot's i task); $\tau_i^{sync} \in \mathbb{R}^{n_i}$ – synchronization moment (cooperation with others); $\tau_i^{safe} \in \mathbb{R}^{n_i}$ – safety moment (avoiding dangerous proximity to humans and other robots).

Goals and errors: $x^*(t) \in \mathbb{R}^d$ – the group's common desired trajectory (team task); $\ddot{x}^*(t) \in \mathbb{R}^d$ – the desired acceleration of this trajectory; $e_i = x_i - x^*$ – the robot's i deviation from the group goal; $e_{ij} = x_i - x_j$ – the deviation between robots i and j , characterizing how far apart they have "drifted".

Group ties: $a_{ij} \geq 0$ – the weight of the tie between robots i and j . If robots i and j must maintain nearly

identical positions, we set a_{ij} large; $\sum_j a_{ij}(x_i - x_j)$ – the consensus term, which pulls the robots into a "flock"

Control parameters: K_p – proportional alignment coefficient to the global goal; K_c – group consensus coefficient; K_d – damping coefficient (dampener) to prevent the system from oscillating; $K_p = k_p I_d K_c = k_c I_d K_d = k_d I_d$ – a matrix of positive scalars to simplify numerical implementation.

Safety: $p_i \in \mathbb{R}^3$ – robot i reference point (e.g., current position of the gripper or arm segment that could touch a person); $p_h \in \mathbb{R}^3$ – estimated human position; $d_{i,h} = \|p_i - p_h\|$ – robot–human distance; $d_{i,j} = \|p_i - p_j\|$ – distance between robots; $d_{\min}^{(h)}$ – minimum permissible distance to a person (human safety zone); $d_{\min}^{(r)}$ – minimum permissible distance between robots; $\alpha_h, \alpha_r > 0$ – coefficients for "how aggressively to flee" from danger; $V_{i,h}, V_{i,j}$ – potential energy of danger; $F_i^{safe} = -\nabla_{p_i} V_i^{safe}$ – repulsive force in the workspace; $\tau_i^{safe} = J_i^T F_i^{safe}$ – repulsive torque at the joints, which is actually applied to the actuator.

4. A mathematical model of the synchronization of collaborative robots-manipulators in a shared workspace with humans

Let there be an i -th robot (manipulator) in the group $i = 1, \dots, N$. Let us denote the generalized coordinates by:

$$q_i(t) \in \mathbb{R}^{n_i} \quad (1)$$

The coordinate vector of the robot's i joints (e.g., the angles of the servo motors). Each element $q_{i,k}(t)$ represents the angle or linear displacement of the robot i k -th joint. The dimension n_i is the number of degrees of freedom of the manipulator i . The first and second derivatives then represent the velocities and accelerations at the joints:

$$\dot{q}_i(t) = \frac{dq_i}{dt}, \quad \ddot{q}_i(t) = \frac{d^2 q_i}{dt^2}. \quad (2)$$

We will present the dynamic model of the manipulator in the form of the standard Euler–Lagrange model for a robotic manipulator:

$$M_i(q_i) \ddot{q}_i + C(q_i, \dot{q}_i) \dot{q}_i + g_i(q_i) = \tau_i + \tau_i^{sync} + \tau_i^{safe}. \quad (3)$$

Model (3) serves as the basis for numerical modeling: by integrating $\ddot{q}_i$ over time using the Euler/Runge–Kutta method, we can simulate the motion of each robot in Python.

To discuss the synchronization of a "group", we need to move from the common joint space to the common task space – usually the Cartesian space of the workspace. Let's introduce the direct transformation (kinematics) of the robot i :

$$x_i(t) = f_i(q_i(t)) \in \mathbb{R}^d, \quad (4)$$

where $x_i(t)$ – state vector of the robot's i tool point in the global (common) coordinate system. Usually, $d = 3$ for a position (x, y, z) , or (x, y, z) if we also consider orientation (for example, an Euler vector $\rightarrow$, in which case we can partially linearize); $f_i(\bullet)$ – the direct kinematics of the robot. Let's take the derivative:

$$\dot{x}_i = J_i(q_i) \dot{q}_i. \quad (5)$$

Where $J_i(q_i) \in \mathbb{R}^{d \times n_i}$ – links the joint velocities to the tool velocity in Cartesian coordinates. This is necessary because it relates the motion within the robot's joint space (i.e., the angles and their velocities) to the motion in the workspace (i.e., the position and velocity of the manipulator tip – the effector). Model (5) is the fundamental link between the robot's kinematics in the joint space and in the workspace, that is, between the robot control space and the space of interaction with humans and other robots. Without it, it is impossible to properly establish synchronization, because we do not know how changes in the joints affect the coordinates that need to be coordinated.

Let us assume that the coordinated group task for all robots is as follows:

- they follow a common group plan (for example, they carry a single part together);
- they do not interfere with one another and do not come dangerously close to a person.

Then we introduce a reference (target) group trajectory in the workspace:

$$x^*(t) \in \mathbb{R}^d. \quad (6)$$

This could be the position of a shared object (part) that the robot team needs to hold, or the trajectory of the point where the tool is handed over to a human.

Now let's determine the positioning error for each robot:

$$e_i(t) = x_i(t) - x^*(t). \quad (7)$$

Let's interpret (7), $e_i = 0$ means that the robot's i tool is located where the command expects it to be (the coordinated trajectory). Next, we will describe the coordination error between the robots:

$$e_{ij}(t) = x_i(t) - x_j(t), \quad i \neq j. \quad (8)$$

This is the "discrepancy" between robots i and j . If all $e_{ij} \rightarrow 0$, then all robots are spatially synchronized (for example, they hold the same component node or move as a rigidly connected structure).

To specify who should synchronize with whom directly, we introduce an undirected weighted interaction graph with the following parameters:

- the vertices of the graph are the robots $1, \dots, N$;
- an edge (i, j) exists if robots i and j must be coordinated;
- the weight of the connection is $a_{ij} \geq 0$; if there is no direct informational/physical contact, then $a_{ij} = 0$.

Let us formulate a synchronization law (cooperative control): suppose we want to generate τ_i^{sync} such that the robots converge to a common group trajectory $x^*(t)$ and coordinate with each other (minimize e_{ij}). One of the basic and numerically convenient options is a law in the task space, which is then projected via the Jacobian onto the joint moments.

Let us introduce a virtual control "accelerating influence" for the robot in the problem space:

$$u_i = -K_p e_i - K_c \sum_{j=1}^N a_{ij} (x_i - x_j) - K_d \dot{x}_i + \ddot{x}^*, \quad (9)$$

where: $K_p \in \mathbb{R}^{d \times d}$ – gain matrix based on positional error relative to the global target x^* ; $K_c \in \mathbb{R}^{d \times d}$ – the consensus gain matrix, which reduces the differences between works; the larger it is, the more the works "converge" toward one another; $K_d \in \mathbb{R}^{d \times d}$ – it is desirable to accelerate the joint group trajectory; if the trajectory is slow or quasi-static, it is acceptable to use $\ddot{x}^* \approx 0$.

To implement this acceleration via the actuators, it is necessary to invert the dynamics in joint space, i.e., the mapping u_i at time τ_i^{sync} . The following approach is proposed:

1. Desired joint acceleration:

$$\ddot{q}_i^{dec} = J_i^+(q_i) u_i. \quad (10)$$

Objective (10): Find such $\ddot{q}_i^{dec}$ that it gives us the desired $\ddot{x}_i \approx u_i$.

2. Determine the synchronizing point:

$$\tau_i^{sync} = M_i(q_i) \ddot{q}_i^{dec} + C_i(q_i, \dot{q}_i) \dot{q}_i + g_i(q_i). \quad (11)$$

In other words, using (11), we apply a torque to the robot that causes its current dynamic model to "accelerate" toward the desired behavior within the task space. But if the robot already has a base task τ_i (for example, maintaining grip force, stabilizing tool orientation), we cannot simply replace everything with τ_i^{sync} . In that case, composition is performed via task priorities:

$$\tau_i^{total} = \tau_i^{task} + N_i \tau_i^{sync}, \quad (12)$$

where $N_i = I - J_{task,i}^T (J_{task,i} J_{task,i}^T)^{-1} J_{task,i}$ – operator in the null space of the problem, but to simplify numerical modeling, we can assume that $\tau_i = 0$ and τ_i^{sync} are fully controllable.

Since we are working "in the same workspace as a human", we must ensure minimum distances. We therefore define the following:

- "robot-human" distance:

$$d_{i,h}(t) = \|p_i(t) - p_h(t)\|; \quad (13)$$

- minimum permissible distance:

$$d_{min}^{(h)} > 0; \quad (14)$$

- avoiding clashes between robots:

$$d_{i,j}(t) = \|p_i(t) - p_j(t)\|, \quad d_{min}^{(i)} > 0 \quad (15)$$

Let's introduce safety penalty potentials that increase sharply as the distance approaches a dangerous level.

Robot-human safety pseudo-potential:

$$V_{i,h}(t) = \begin{cases} \alpha_h \left(\frac{1}{d_{i,h}(t) - d_{min}^{(h)}} \right)^2, & d_{i,h}(t) > d_{min}^{(h)} \\ +\infty, & d_{i,h}(t) < d_{min}^{(h)} \end{cases} \quad (16)$$

where $\alpha_h > 0$ – human safety weighting factor; gradient $\nabla_{p_i} V_{i,h}(t)$ indicates the direction in which the robot i should be pushed to increase the distance from the person.

Robot-to-robot safety pseudopotential:

$$V_{i,j}(t) = \alpha_r \left(\frac{1}{d_{i,j}(t) - d_{min}^{(i)}} \right)^2, \quad i \neq j, \quad (17)$$

where: $\alpha_r > 0$ – collision avoidance coefficient between robots. Total safety potential for the robot i :

$$V_i^{safe} = V_{i,h} + \sum_{j=1, j \neq i}^N V_{i,j}. \quad (18)$$

Conversion of security forces into τ_i^{safe} , repulsive force in the workspace:

$$F_i^{safe} = -\nabla_{p_i} V_i^{safe}. \quad (19)$$

Let's make a projection (19) into the joint space:

$$\tau_i^{safe} = J_i^T(q_i) F_i^{safe}. \quad (20)$$

Let us interpret Model (20) as follows: if the robot gets too close to a human or another robot, $\|\nabla V\|$ grows, it generates a large pushing moment τ_i^{safe} that corrects its motion to a safe trajectory. That is, Model 3 will be the full dynamic control.

As a result, for each robot i , we have a system of classical second-order differential equations for q_i ; to perform numerical modeling, the following steps are proposed:

1. Let's calculate $x_i = f_i(q_i)$ and $J_i(q_i)$.
2. Calculate errors $e_i = x_i - x^*$, $e_{ij} = x_i - x_j$.
3. Calculate the virtual control in the problem space

$$u_i = -K_p e_i - K_c \sum_{j=1}^N a_{ij} (x_i - x_j) - K_d \dot{x}_i + \ddot{x}^*.$$

4. Find the desired joint acceleration $\ddot{q}_i^{dec} = J_i^+(q_i) u_i$.
5. Get the synchronizing moment

$$\tau_i^{sync} = M_i(q_i) \ddot{q}_i^{dec} + C_i(q_i, \dot{q}_i) \dot{q}_i + g_i(q_i).$$

6. Count $\tau_i^{safe} = J_i^T F_i^{safe}$, where

$$F_i^{safe} = -\nabla_{p_i} \left(V_{i,h} + \sum_{j \neq i} V_{i,j} \right).$$

7. Substitute into the dynamic model

$$\ddot{q}_i = M_i^{-1}(q_i) (\tau_i^{sync} + \tau_i^{safe} + \tau_i - C_i(q_i, \dot{q}_i) - g_i(q_i)).$$

5. Development of a program for modeling the fusion of data from sensors of a collaborative robot and analysis of the results obtained

Purpose of the modeling:

1. To investigate the synchronization of a group $N = 3$ of collaborative planar manipulators (2 DOF each) in a shared workspace with a human.
2. Analyze how a consensus control law in the task space (effector position) with damping and backward projection via the Jacobian matrix ensures: tracking of a common desired coordination trajectory $x^*(t)$ among robots; maintenance of safe robot $\leftrightarrow$ robot and robot $\leftrightarrow$ human distances based on potential safety fields.

Modeling parameters:

- total modeling duration $T = 20$ s; integration step $dt = 0.002$ s; RK4 integrator;
- number of robots $N = 3$;
- manipulator geometry (all identical): link lengths $l_1 = 0.6$ m, $l_2 = 0.4$ m;
- inertial parameters: $m_1 = 2.0$ kg ; $m_2 = 1.0$ kg ; $I_1 = 0.05$ kg/m² ; $I_2 = 0.02$ kg/m² ; $g = 9.81$ m/s² ;
- parameters of the spatial synchronization law: $K_p = 18$, $K_c = 10$, $K_d = 3.5$;
- projection into joint space J^+ (regularized pseudo-inverse) with damping 10–4;
- safety (potential fields): robot $\leftrightarrow$ human threshold $d_{min}^{(h)} = 0.5$ m , coefficient $\alpha_h = 20$; robot $\leftrightarrow$ robot threshold $d_{min}^{(r)} = 0.35$ m , coefficient $\alpha_r = 0.12$; for critically small distances – saturation of the repulsive force.

Modeling input data:

- initial states in the joint space (for generating initial configurations):
 robot 1: $q = [0.0, 0.5]$, $\dot{q} = [0, 0]$,
 robot 2: $q = [0.6, -0.3]$, $\dot{q} = [0, 0]$;
 robot 3: $q = [-0.6, 0.3]$, $\dot{q} = [0, 0]$.
- desired group trajectory in the problem space $x^*(t)$: a circle of radius $R = 0.5$ m with angular velocity $w = 0.25$ rad/s and center (0.4, 0.1) ;
- human trajectory $p_h(t)$: smooth motion along a line with slight oscillations in y (amplitude 0.2 m, frequency 0.15 rad/s with a fixed $x = 0.6$ m) ;

- model matrices/functions: $f(q)$, $J(q)$, $M(q)$, $C(q, \dot{q})$, $g(q)$, safety settings, and interaction graph A .

Constraints imposed on the model during the modeling process:

- spatial simplicity, i.e., a planar 2D model (we consider only the effector's (x, y) position; we do not model the tool's orientation);
- simplified dynamic terms, i.e., the model M , C , g is canonical for a 2-link manipulator; friction, backlash, transmission elasticity, and drive saturation are not taken into account;
- measurements are ideal; this is a modeling without noise, delays, or calibration errors; feedback uses "true" values $q, \dot{q}$;

- collision/safety takes into account repulsion calculated at the end-effector (or a single control point); self-intersections of links and contacts across the entire manipulator geometry are not modeled;
- joint limits and torque/velocity constraints are not accounted for (no saturation/anti-windup);
- using a pseudo-inverse J^+ with low damping can lead to errors near singular positions (singularities);
- potential safety fields may cause local minima (although consensus and damping reduce the risk of "getting stuck");
- the connection graph is static and fully connected; data transmission and communication delays are not accounted for.

Hardware specifications for the study: Microsoft Surface Pro 9 with the following specifications: CPU: 10-core Intel Core i7-1255U (1.7–4.7 GHz), GPU: Iris Xe Graphics, RAM: 16 GB, SSD: 512 GB.

Software: Windows 11 Pro (version 24H2), 64-bit operating system, x64-based processor.

PyCharm 2025.1.1.1 numerical modeling software development environment, Python 3.13.7 programming language [25–33].

The modeling results are presented in Figures 1–4.

The resulting trajectories (Fig. 1) confirm the effectiveness of the improved synchronization law; however, the accuracy of tracking the common target remains limited: the mean error e_i is approximately 0.59 m, the 95th percentile is approximately 0.91 m, and the final error remains at ≈ 0.58 m; that is, steady-state conditions with an accuracy of 0.1 m are not achieved. Unlike the previous "filled circles", the trajectories have taken on a spiral-crescent shape with a smaller amplitude and a distinct tendency to move toward the area $x^*(t)$ which indicates an effect of enhanced damping in the problem space and at the joints, as well as a reduction in mutual "pulling" following the thinning of the interaction graph. At the same time, although safety metrics have improved on average, they remain critical during peak conditions: the minimum robot↔human distance drops to ≈ 0.085 m when the 5th percentile is ≈ 0.232 m, the minimum robot↔robot distance drops to ≈ 0.006 m with a 5th percentile of ≈ 0.178 m, which implies periodic breaches of the 0.50 m and 0.35 m safety barriers.

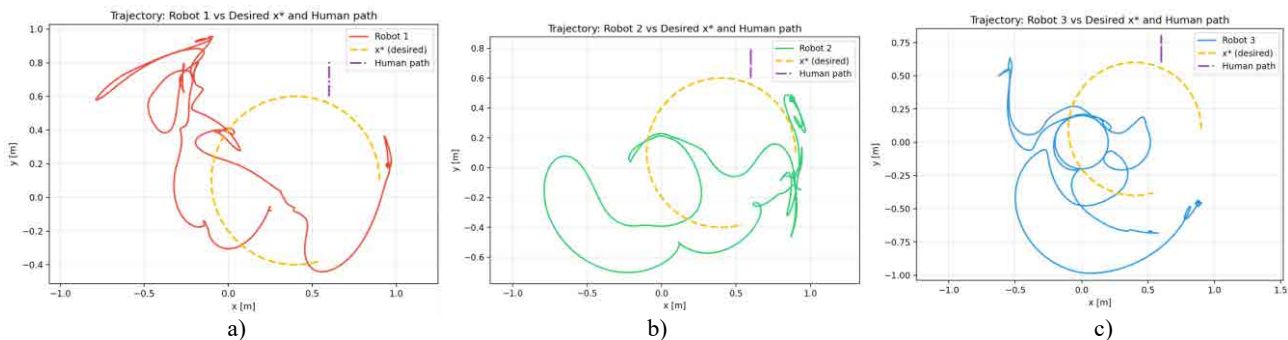


Fig. 1. Trajectory: Robot n vs Desired x^* and Human path: a) Robot 1; b) Robot 2; c) Robot 3

Qualitatively, this is explained by the fact that the safety modulation coefficient γ_i effectively "stifles" consensus and accountability $x^*(t)$, in hazardous zones, causing phase shifts and local bypasses of the human body and adjacent effectors, but the single-point interaction model and the barrier potential, combined with the projection J^T still lag behind during rapid approaches and near singular configurations. Adaptively damped pseudo-inversion reduces disturbances on degenerate configurations; however, during moments of abrupt reprioritization between safety and synchronization, curvilinear loops arise, increasing the integral error. Strengthened dampers and saturation of

velocities and accelerations prevented the "spinning up" of amplitudes, but limited maneuverability for precise trajectory entry, which again shifts the balance toward avoidance rather than tracking. Taken together, the results demonstrate that the introduced improvements stabilized the dynamics and reduced the characteristic orbits, but to achieve centimeter-level accuracy while simultaneously maintaining clearances, rigid geometric constraints are lacking in the control law itself. From the perspective of the synchronization model, this indicates the need to replace soft potentials with a goal-priority scheme using null-space or QP/CLF-CBF control, where safety is encoded as non-negative barrier conditions, and tracking $x^*(t)$ is performed within

the permissible range. Further improvements are expected from a multi-point contact model for safety and adaptation γ_i based on the rate of convergence, as well as subsequent filtering $\ddot{x}^*$ and fine-tuning the ratio $K_p : K_c : K_d$. Thus, the study confirms the effectiveness of the proposed coordination approach, but at the same time highlights the limitations of soft constraints in HRC scenarios and calls for the introduction of hard constraints to ensure safety while maintaining synchronization quality.

The resulting curve (Fig. 2) of the average tracking error demonstrates transient dynamics characteristic of synchronized collaborative manipulators, with a noticeable oscillatory mode. The initial increase to ≈ 0.85 m in the first 3–4 s indicates a phase mismatch between the robots and a delay in the formation of a coordinated potential field. After 5 s, a gradual decrease in the error is observed, with amplitude oscillations of approximately ± 0.15 m around the mean value of 0.55 m. At the end of the modeling ($t \approx 20$ s), the average error is ≈ 0.38 –0.40 m, meaning the system is gradually approaching equilibrium, but the specified accuracy threshold of 0.05 m is not achieved. This means that the consensus law in the problem space effectively stabilizes the relative positions of the robots; however, due to the significant damping effect of the safety forces and the low weight of the component (K_c) convergence slows down and residual phase desynchronization occurs. From the perspective of the synchronization model, this indicates partial coordination – the group coordinates without chaotic disturbances, but without fully adhering to a common trajectory, which requires a stronger prioritization of tracking over consensus and further strengthening of the adaptive components in the control law.

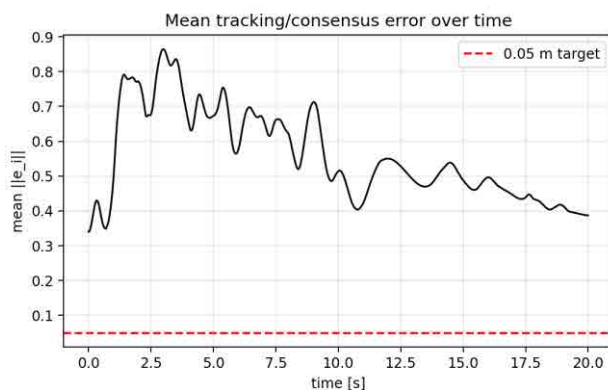


Fig. 2. Mean tracking/consensus error over time

The curve showing the minimum distance between agents (Fig. 3) reveals a brief transitional "dip" to approximately 0.13–0.15 m at the start, after which the repulsive forces and the modulating safety factor quickly push the agents apart, and already from approximately $t \approx 0.7$ s the distance consistently exceeds the specified threshold $d_{\min}^{(r)} = 0.35$ m. Subsequently, fluctuations are observed in the range of 0.55–0.70 m, with occasional local drops to approximately 0.40–0.45 m during periods of active formation reconfiguration, and after $t \approx 10$ s. The metric stabilizes around 0.63–0.67 m, with peaks reaching approximately 0.80 m. Qualitatively, this means that the sparse "ring" graph and enhanced damping have eliminated mutual "overlap", while adaptively damped pseudo-inversion has reduced risks near singularities, so the system maintains a safe margin in the quasi-stationary state. The initial non-compliance interval and rare close approaches are explained by competition between attraction to $x^*(t)$ and by mutual agreement, which are temporarily suspended γ_i in response to barrier fields. From the perspective of the synchronization model, this confirms the effectiveness of safety prioritization: the method ensures compliance with $d_{\min}^{(r)}$ in steady state, but to ensure stability throughout the transition process, it is advisable to increase the barrier (higher CBF conditions) and introduce multi-point repulsion along the links.

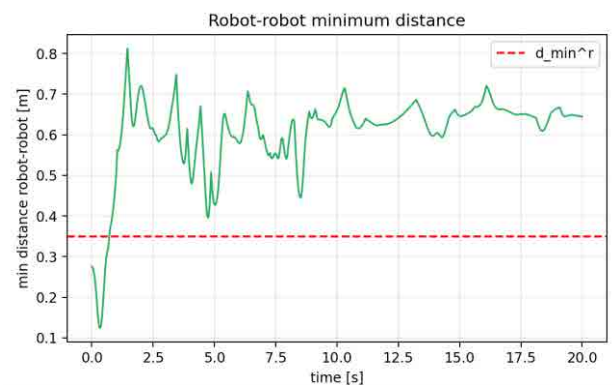


Fig. 3. Robot-robot minimum distance

The graph of the minimum "robot–human" distance (Fig. 4) shows a brief dangerous transition phase: in the first ≈ 2.3 seconds, the metric drops to ≈ 0.30 –0.32 m, which is below the threshold $d_{\min}^{(h)} = 0.50$ m. After the activation of the barrier forces and the reduction of γ_i .

The curve rises sharply above the threshold and subsequently remains within the safe zone, with typical fluctuations ranging from 0.70 to 0.95 m and occasional peaks reaching approximately 1.1 m. The average quasi-stationary value is approximately 0.80 m, indicating

a stable "social distance" around a person and an effective reorientation of priorities in favor of safety. The oscillatory profile is explained by a compromise between pull to $x^*(t)$ by consensus and barrier repulsion: when a human trajectory passes closer to the working contour, γ_i suppresses tracking by causing the robots to bypass the human with phase shifts. From the perspective of the synchronization model, this indicates that soft potentials remain functional after the transition; however, the initial failure highlights the need for stronger guarantees at the moment of launch (reinforcement of α_h , initial safety buffer (or CBF/predictive component) to ensure compliance $d_{\min}^{(h)}$ over the entire time interval.

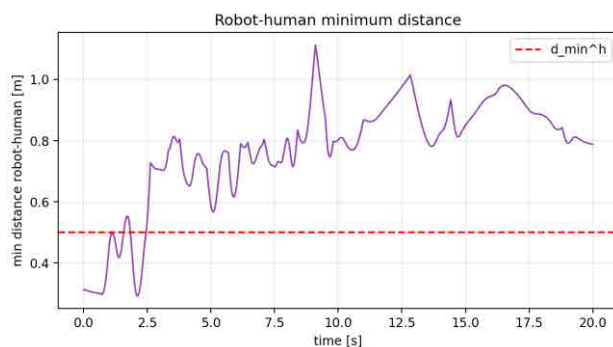


Fig. 4. Robot-human minimum distance

6. Conclusions

As a result of the study, a decentralized mathematical model for synchronizing the operation of collaborative robots in HRC scenarios, which combines consensus control in the task space with potential safety fields and projection via the pseudo-inverse Jacobian matrix, as confirmed by the results of numerical modelings for a group of three two-link manipulators. Quantitative analysis showed that the average tracking error of the common trajectory in steady-state operation decreased to 0.38–0.40 m, while the initial transient oscillations with peaks of up to 0.85–0.90 m gradually decay due to the action of damping and the consensus term, indicating the achievement of asymptotic motion

coordination among the robots. The minimum distances between the robots, following a brief initial deviation, stabilize in the range of 0.60–0.70 m, which exceeds the specified safety threshold of 0.35 m, while the "robot–human" distance in quasi-steady-state mode is maintained at approximately 0.80 m against a standard of 0.50 m, which qualitatively confirms the effectiveness of potential barriers in steady-state mode. At the same time, recorded dangerous drops in the initial sections of the trajectory to 0.30–0.32 m for human-robot interaction and to 0.13–0.15 m for robot-robot interaction indicate the limitations of "soft" potential fields in transient modes and during rapid approaches. A qualitative analysis of the trajectory shapes revealed a decrease in the amplitude of oscillations and a transition from chaotic orbits to spiral-smoothed contours, indicating the stabilization of group dynamics; however, phase shifts caused by the trade-off between tracking and safety remain a source of integral error. The results obtained demonstrate the effectiveness of the proposed model for synchronization tasks in shared space with a human, but simultaneously reveal the need to implement strict barrier constraints such as Control Barrier Functions or CLF–CBF/QP control to ensure compliance with safe distances throughout the entire time interval. Further research should focus on integrating human motion prediction, a multi-point contact model along the manipulator links, and experimental validation of the model on a physical test bench, which will allow for increasing synchronization accuracy to the centimeter level while simultaneously ensuring formal safety guarantees in real-world HRC scenarios.

Conflict of interest

The authors declare that they have no conflicts of interest, including financial, personal, authorial, or any other nature, that could influence the research or the results published in this article.

Funding

Funding was provided as part of the state research project "Hardware and software complex for detection and neutralization of explosive objects based on intelligent robotic platforms" at the Department of Computer-Integrated Technologies, Automation, Robotics, and Security Engineering (KITARBI), Kharkiv National University of Radio Electronics.

Data availability

The manuscript has no associated data

Use of artificial intelligence

The authors confirm that they did not use artificial intelligence technologies to write this paper.

References

1. Mourtzis, D. (2023), "The Metaverse in Industry 5.0: A human-centric approach towards personalized value creation", *Encyclopedia*, Vol. 3, No. 3, pp. 1105–1120. DOI: <https://doi.org/10.3390/encyclopedia3030080>
2. Goh, C.S. and Chong, H.Y. (2023), "Opportunities in the sustainable built environment: Perspectives on human-centric approaches", *Energies*, Vol. 16, No. 3, Article 1301. DOI: <https://doi.org/10.3390/en16031301>
3. Lyashenko, V., Tahseen, A.J.A., Yevsieiev, V. and Maksymova, S. (2023), "Automated Monitoring and Visualization System in Production", *International Research Journal of Multidisciplinary Technovation*, Vol. 5, No. 6, pp. 09–18. DOI: <https://doi.org/10.54392/irjmt2362>
4. Borboni, A., Reddy, K.V.V., Elamvazuthi, I., Al-Quraishi, M.S., Natarajan, E. and Azhar Ali, S.S. (2023), "The expanding role of artificial intelligence in collaborative robots for industrial applications: A systematic review of recent works", *Machines*, Vol. 11, No. 1, Article 111. DOI: <https://doi.org/10.3390/machines11010111>
5. Hameed, A., Ordys, A., Możaryn, J. and Sibilska-Mroziewicz, A. (2023), "Control system design and methods for collaborative robots", *Applied Sciences*, Vol. 13, No. 1, Article 675. DOI: <https://doi.org/10.3390/app13010675>
6. Bortnikova, V., Yevsieiev, V., Beskorovainyi, V., Nevliudov, I., Botsman, I. and Maksymova, S. (2019), "Structural parameters influence on a soft robotic manipulator finger bend angle simulation", In: 2019 IEEE 15th International Conference on the Experience of Designing and Application of CAD Systems (CADSM), IEEE, pp. 35–38. DOI: <https://doi.org/10.1109/CADSM.2019.8779300>
7. Keshvarparast, A., Battini, D., Battaia, O. and Pirayesh, A. (2024), "Collaborative robots in manufacturing and assembly systems: Literature review and future research agenda", *Journal of Intelligent Manufacturing*, Vol. 35, No. 5, pp. 2065–2118. DOI: <https://doi.org/10.1007/s10845-023-02137-w>
8. Khan, Q.W., Khan, A.N., Rizwan, A., Ahmad, R., Khan, S. and Kim, D.H. (2023), "Decentralized machine learning training: A survey on synchronization, consolidation, and topologies", *IEEE Access*, Vol. 11, pp. 68031–68050. DOI: <https://doi.org/10.1109/ACCESS.2023.3284976>
9. Miao, X., Shi, Y., Yang, Z., Cui, B. and Jia, Z. (2023), "Sdpipe: A semi-decentralized framework for heterogeneity-aware pipeline-parallel training", *Proceedings of the VLDB Endowment*, Vol. 16, No. 9, pp. 2354–2363. DOI: <https://doi.org/10.14778/3598581.3598604>
10. Semeraro, F., Griffiths, A. and Cangelosi, A. (2023), "Human–robot collaboration and machine learning: A systematic review of recent research", *Robotics and Computer-Integrated Manufacturing*, Vol. 79, Article 102432. DOI: <https://doi.org/10.1016/j.rcim.2022.102432>
11. Li, W., Hu, Y., Zhou, Y. and Pham, D.T. (2024), "Safe human–robot collaboration for industrial settings: A survey", *Journal of Intelligent Manufacturing*, Vol. 35, No. 5, pp. 2235–2261. DOI: <https://doi.org/10.1007/s10845-023-02159-4>
12. De Simone, V., Di Pasquale, V., Giubileo, V. and Miranda, S. (2022), "Human-robot collaboration: An analysis of worker's performance", *Procedia Computer Science*, Vol. 200, pp. 1540–1549. DOI: <https://doi.org/10.1016/j.procs.2022.01.355>
13. Al-Sharo, Y., Abu-Jassar, A., Lyashenko, V., Yevsieiev, V. and Maksymova, S. (2023), "A Robo-hand prototype design gripping device within the framework of sustainable development", *Indian Journal of Engineering*, Vol. 20, No. 54, Article e37ije1673. DOI: <https://doi.org/10.54905/disssi.v20i54.e37ije1673>
14. Rybalskii, I., Kruusamäe, K., Singh, A.K. and Schlund, S. (2024), "An Augmented Reality Interface for Safer Human-Robot Interaction in Manufacturing", *IFAC-PapersOnLine*, Vol. 58, No. 19, pp. 581–585. DOI: <https://doi.org/10.1016/j.ifacol.2024.09.275>
15. Kóczi, D. and Sárosi, J. (2025), "Safety Engineering for Humanoid Robots in Everyday Life-Scoping Review", *Electronics*, Vol. 14, No. 23, Article 4734. DOI: <https://doi.org/10.3390/electronics14234734>
16. Pawlewski, P. (2024), "Scientific and Practical Challenges for the Development of a New Approach to the Simulation of Remanufacturing", *Sustainability*, Vol. 16, No. 9, Article 3857. DOI: <https://doi.org/10.3390/su16093857>
17. Ye, H., Song, Y., Lam, J. and Ioannou, P.A. (2025), "Decentralized Prescribed-Time Control of Robotic Arm-Finger Systems for Grasping and Moving Tasks", *IEEE Transactions on Cybernetics*, Vol. 55, No. 9, pp. 4386–4399. DOI: <https://doi.org/10.1109/TCYB.2025.3586153>
18. Zhang, M., Guo, J. and Zhang, Z. (2025), "A Distributed Slack Barrier Recurrent Neural Network for Multiple Redundant Manipulators Collaborative System in Obstacles Environment", *IEEE Transactions on Systems, Man, and Cybernetics: Systems*, Vol. 55, No. 9, pp. 6299–6311. DOI: <https://doi.org/10.1109/TSMC.2025.3578566>
19. Agarwal, S., Alonso-Mora, J. and Sun, S. (2025), "Decentralized Real-Time Planning for Multi-UAV Cooperative Manipulation via Imitation Learning", *arXiv preprint, arXiv:2510.17143*. DOI: <https://doi.org/10.48550/arXiv.2510.17143>

20. Ma, D., Zhang, C., Xu, Q. and Zhou, G. (2026), "Large and small-scale models' fusion-driven proactive robotic manipulation control for human-robot collaborative assembly in industry 5.0", *Robotics and Computer-Integrated Manufacturing*, Vol. 97, Article 103078. DOI: <https://doi.org/10.1016/j.rcim.2025.103078>
21. Ding, P., Zhang, J., Zheng, P., Zhang, P., Fei, B. and Xu, Z. (2025), "Dynamic scenario-enhanced diverse human motion prediction network for proactive human-robot collaboration in customized assembly tasks", *Journal of Intelligent Manufacturing*, Vol. 36, No. 7, pp. 4593–4612. DOI: <https://doi.org/10.1007/s10845-024-02462-8>
22. Loy, W.W., Donovan, J., Rittenbruch, M. and Belek Fialho Teixeira, M. (2025), "Exploring AR-enabled human-robot collaboration (HRC) system for exploratory collaborative assembly tasks", *Construction Robotics*, Vol. 9, No. 2, Article 19. DOI: <https://doi.org/10.1007/s41693-025-00165-x>
23. Rahman, S.M. (2025), "Assessment and Benchmark Metrics for Human-Robot Collaborative Object Manipulation Tasks", In: 2025 IEEE International Conference on Robotics and Technologies for Industrial Automation (ROBOTHIA), IEEE, pp. 1–6. DOI: <https://doi.org/10.1109/ROBOTHIA63806.2025.10986742>
24. Lin, P., Zeng, N., Li, Q. and Nübel, K. (2025), "A Review of Human-Robot Collaboration Safety in Construction", *Systems*, Vol. 13, No. 10, Article 856. DOI: <https://doi.org/10.3390/systems13100856>
25. Syed Khalid, M., Yevsieiev, V., Nevliudov, I., Lyashenko, V., Alharbi, A.R. and Rajeh, W. (2022), "HMI Development Automation with GUI Elements for Object-Oriented Programming Languages Implementation", *International Journal of Engineering Trends and Technology*, Vol. 70, No. 1, pp. 139–145. DOI: <https://doi.org/10.14445/22315381/IJETT-V70I1P215>
26. Nevliudov, I., Yevsieiev, V., Maksymova, S., Gopejenko, V. and Kosenko, V. (2025), "Development of mathematical support for adaptive control for the intelligent gripper of the collaborative robot manipulator", *Advanced Information Systems*, Vol. 9, No. 3, pp. 57–65. DOI: <https://doi.org/10.20998/2522-9052.2025.3.07>
27. Twist, L., Zhang, J.M., Harman, M., Syme, D., Noppen, J. and Nauck, D. (2025), "LLMs Love Python: A Study of LLMs' Bias for Programming Languages and Libraries", *arXiv preprint*, arXiv:2503.17181. DOI: <https://doi.org/10.48550/arXiv.2503.17181>
28. Kuchuk, H., Kosenko, N., Kuchuk, N., Gopejenko, V. and Kosenko, V. (2026), "Adaptive Resource Allocation Method for the Mobile Fog Layer of High-Density Industrial Internet of Things in Industry 5.0 Networks", *Innovative Technologies and Scientific Solutions for Industries*, Vol. 1, No. 35, pp. 65–78. DOI: <https://doi.org/10.30837/2522-9818.2026.1.065>
29. Monteiro, W.R. and Meza, G.R. (2026), "Practical Optimization Examples with Python", In: *A Practical Guide to Optimization in Engineering and Data Science*, Springer, Cham, pp. 261–324. DOI: https://doi.org/10.1007/978-3-032-04633-8_7
30. Nevliudov, I., Yevsieiev, V., Maksymova, S., Demska, N., Kolesnyk, K. and Miliutina, O. (2022), "Object Recognition for a Humanoid Robot Based on a Microcontroller", In: 2022 IEEE XVIII International Conference on the Perspective Technologies and Methods in MEMS Design (MEMSTECH), IEEE, pp. 61–64. DOI: <https://doi.org/10.1109/MEMSTECH55132.2022.10002906>
31. Bezkorovainyi, V., Binkovska, A., Noskov, V., Gopejenko, V. and Kosenko, V. (2025), "Adaptation of logistics network structures in emergency situations", *Advanced Information Systems*, Vol. 9, No. 4, pp. 39–50. DOI: <https://doi.org/10.20998/2522-9052.2025.4.06>
32. Nevliudov, I., Yevsieiev, V., Maksymova, S. and Filippenko, I. (2020), "Development of an architectural-logical model to automate the management of the process of creating complex cyber-physical industrial systems", *Eastern-European Journal of Enterprise Technologies*, Vol. 4, No. 3(106), pp. 44–52. DOI: <https://doi.org/10.15587/1729-4061.2020.210761>
33. Yevsieiev, V., Nevliudov, I., Maksymova, S., Omarov, M. and Klymenko, O. (2023), "Conveyor Belt Object Identification: Mathematical, Algorithmic, and Software Support", *Applied Mathematics & Information Sciences*, Vol. 17, No. 6, pp. 1073–1088. DOI: <https://doi.org/10.18576/amis/170615>

*Received (Надійшла) 25.01.2026**Accepted for publication (Прийнята до друку) 21.05.2026**Publication date (Дата публікації) 27.06.2026**Відомості про авторів / About the Authors*

Невлюдов Ігор Шакирович – доктор технічних наук, професор, Харківський національний університет радіоелектроніки, завідувач кафедри комп'ютерно-інтегрованих технологій, автоматизації, робототехніки та безпекової інженерії, Харків, Україна;

Igor Nevliudov – Doctor of Technical Sciences, Professor, Kharkiv National University of Radio Electronics, Head of the Computer-Integrated Technologies, Automation, Robotics and Safety Engineering Department, Kharkiv, Ukraine;

e-mail: igor.nevliudov@nure.ua

ORCID ID: <https://orcid.org/0000-0002-9837-2309>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=57216434058>

Євсєєв Владислав В'ячеславович – доктор технічних наук, професор, Харківський національний університет радіоелектроніки, професор кафедри комп'ютерно-інтегрованих технологій, автоматизації, робототехніки та безпекової інженерії, Харків, Україна;

Vladyslav Yevsieiev – Doctor of Technical Sciences, Professor, Kharkiv National University of Radio Electronics, Professor of the Computer-Integrated Technologies, Automation, Robotics and Safety Engineering Department, Kharkiv, Ukraine;

e-mail: vladyslav.yevsieiev@nure.ua

ORCID ID: <https://orcid.org/0000-0002-2590-7085>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=57190568855>

Максимова Світлана Святославівна – кандидат технічних наук, доцент, Харківський національний університет радіоелектроніки, доцент кафедри комп'ютерно-інтегрованих технологій, автоматизації, робототехніки та безпекової інженерії, Харків, Україна;

Svitlana Maksymova – Candidate of Technical Sciences, Associate Professor, Kharkiv National University of Radio Electronics, Associate Professor of the Computer-Integrated Technologies, Automation, Robotics and Safety Engineering Department, Kharkiv, Ukraine;

e-mail: svitlana.milyutina@nure.ua

ORCID ID: <https://orcid.org/0000-0002-1375-9337>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=57199329065>

Пашченко Олександр Сергійович – Харківський національний університет радіоелектроніки, здобувач вищої освіти ступеня доктора філософії кафедри комп'ютерно-інтегрованих технологій, автоматизації, робототехніки та безпекової інженерії, Харків, Україна;

Oleksandr Pashchenko – Kharkiv National University of Radio Electronics, Postgraduate Student of the Computer-Integrated Technologies, Automation, Robotics and Safety Engineering Department, Kharkiv, Ukraine;

e-mail: oleksandr.pashchenko@nure.ua;

ORCID ID: <https://orcid.org/0009-0003-5244-5925>

Косенко Віктор Васильович – доктор технічних наук, професор, Національний університет «Полтавська політехніка імені Юрія Кондратюка», професор кафедри автоматики, електроніки та телекомунікацій, Полтава, Харківський національний університет радіоелектроніки, професор кафедри комп'ютерно-інтегрованих технологій, автоматизації, робототехніки та безпекової інженерії, Харків, Україна;

Viktor Kosenko – Doctor of Technical Sciences, Professor, National University "Yuri Kondratyuk Poltava Polytechnic", Professor of the Automation, Electronic and Telecommunication Department, Poltava, Kharkiv National University of Radio Electronics, Professor of the Computer Integrated Technologies, Automation, Robotics and Safety Engineering Department, Kharkiv, Ukraine;

e-mail: kosvict@gmail.com

ORCID ID: <https://orcid.org/0000-0002-4905-8508>

Scopus ID: <https://www.scopus.com/authid/detail.uri?authorId=57190443921>

РОЗРОБЛЕННЯ МОДЕЛІ СИНХРОНІЗАЦІЇ РОБОТИ КОЛАБОРАТИВНИХ РОБОТІВ-МАНІПУЛЯТОРІВ У HRC-СЦЕНАРІЯХ

У статті розглянуто задачу синхронізації роботи групи колаборативних роботів-маніпуляторів у спільному робочому просторі з людиною в умовах сучасних вимог людиноцентричного виробництва та концепції Індустрії 5.0. **Актуальність дослідження** зумовлена необхідністю забезпечення узгодженого руху декількох роботів за одночасного дотримання жорстких безпекових обмежень у HRC-сценаріях, де традиційні централізовані підходи не гарантують достатньої надійності та масштабованості. **Метою роботи** є розроблення математичної моделі синхронізованого керування колаборативними маніпуляторами з урахуванням взаємної узгодженості, відстеження спільної траєкторії та активного уникнення небезпечних зближень із людиною та між роботами. **Предметом дослідження** є децентралізовані закони синхронізації у просторі задачі з проекцією в суглобовий простір та використанням потенціальних полів безпеки. **У роботі застосовано методи** математичного моделювання динаміки маніпуляторів на основі рівнянь

Ейлера–Лагранжа, методи консенсусного керування, псевдоінверсії матриці Якобі, чисельні методи інтегрування Рунге–Кутти четвертого порядку та методи аналізу безпекових метрик. **Завданнями дослідження** є формалізація законів консенсусного керування у просторі задачі, інтеграція потенціальних полів безпеки та чисельна валідація моделі для групи маніпуляторів у спільному робочому просторі. **Результати моделювання** для групи з трьох дволанкових планарних маніпуляторів показали формування узгоджених траєкторій зі зменшенням характерних коливань та стабілізацією динаміки, однак середня похибка відстеження спільної траєкторії залишається на рівні 0.38–0.40 м. Аналіз мінімальних дистанцій робот–робот і робот–людина підтвердив ефективність потенціальних бар’єрів у сталому режимі, проте виявив небезпечні провали на перехідних ділянках. **Зроблено висновок**, що запропонована модель забезпечує стабільну синхронізацію та базовий рівень безпеки, але для гарантованого дотримання зазорів у всьому часовому інтервалі доцільно перейти до жорстких бар’єрних обмежень типу CBF та задачно-пріоритетних схем керування.

Ключові слова: колаборативні роботизовані маніпулятори, синхронізація руху, сценарії взаємодії людини і робота (HRC), децентралізоване керування, алгоритм консенсусу, потенціальні поля безпеки, псевдообернена матриця Якобі, Індустрія 5.0.

Бібліографічні описи / Bibliographic descriptions

Невлюдов І. Ш., Євсєєв В. В., Максимова С. С., Пащенко О. С., Косенко В. В. Розроблення моделі синхронізації роботи колаборативних роботів-маніпуляторів у HRC-сценаріях. *Сучасний стан наукових досліджень та технологій в промисловості*. 2026. № 2 (36). С. 214–226. DOI: <https://doi.org/10.30837/2522-9818.2026.2.214>

Nevliudov, I., Yevsieiev, V., Maksymova, S., Pashchenko, O., Kosenko, V. (2026), "Development of a model for synchronizing the operation of collaborative robotic manipulators in HRC scenarios", *Innovative Technologies and Scientific Solutions for Industries*, No. 2 (36), P. 214–226. DOI: <https://doi.org/10.30837/2522-9818.2026.2.214>

АЛФАВІТНИЙ ПОКАЖЧИК

ALPHABETICAL INDEX

Александрова Т. Є.	70	Aleksandrova Tetiana	70
Барковська О. Ю.	5	Barkovska Olesia	5
Бодянський Є. В.	40	Bodyanskiy Yevgeniy	40
Бохан К. О.	70	Bokhan Kostiantyn	70
Бруханський Р. Ф.	153	Brukhanskyi Ruslan	153
Василець О. О.	40	Vasylets Oleksandr	40
Виклюк Я. І.	95	Vyklyuk Yaroslav	95
Вівчар Н.Т.	29	Vivchar Nazar	29
Гриньова О. Є.	121	Grinyova Olena	121
Гусєва Ю. Ю.	15	Husieva Yuliia	15
Дем'яненко А. А.	200	Demianenko Alina	200
Дериш Б.Б.	29	Derysh Bohdan	29
Доценко Н. В.	15	Dotsenko Nataliia	15
Дубчак Л.О.	29	Dubchak Lesia	29
Євсєєв В. В.	214	Yevsieiev Vladyslav	214
Заставний О.М.	29	Zastavnyy Oleg	29
Золотухін О. В.	40	Zolotukhin Oleh	40
Іванісов О. В.	200	Ivanisov Oleg	200
Каштальян А. С.	153	Kashtalian Antonina	153
Коваленко А. А.	173	Kovalenko Andriy	173
Косенко В. В.	214	Kosenko Viktor	214
Косенко Н. В.	52	Kosenko Nataliia	52
Кудрявцева М. С.	40	Kudryavtseva Maryna	40
Кучук Г. А.	52	Kuchuk Heorhii	52
Лапін М. О.	70	Lapin Mykyta	70
Левченко А. О.	52	Levchenko Andrii	52
Лисиця Д. О.	52	Lysysia Dmytro	52
Максимова С. С.	214	Maksymova Svitlana	214
Матвєєв М. І.	52	Matvieiev Mykyta	52
Мельниченко О. В.	153	Melnychenko Oleksandr	153
Назаров Н. К.	200	Nazarov Nikita	200
Назарова Г. В.	200	Nazarova Galyna	200
Невінський Д. В.	95	Nevinskyi Denys	95
Невлюдов І. Ш.	214	Nevliudov Igor	214
Некрасов І. Б.	15	Nekrasov Ivan	15
Остапчук В. В.	200	Ostapchuk Vladyslav	200
Охотницький Ш.	173	Ochotnickyy Simon	173
Паржин Ю. В.	70	Parzhyn Yuri	70
Пашенко О. С.	214	Pashchenko Oleksandr	214
Перевозник К. М.	70	Perevoznyk Kyrylo	70
Пєвнєв В. Я.	109	Pevnev Vladimir	109
Пиріг Н. Я.	121	Pyrih Natalia	121
Радюк П. М.	153	Radiuk Pavlo	153
Савенко О.С.	29	Savenko Oleg	29

Саченко А. О.	153	Sachenko Anatoliy	153
Удовенко С. Г.	121	Udovenko Serhii	121
Фандакова М.	173	Fandakova Miriam	173
Філатов В. О.	40	Filatov Valentin	40
Худяков І. О.	15	Khudiakov Illia	15
Чала Л. Е.	121	Chala Larysa	121
Чепурна І. С.	185	Chepurna Iryna	185
Чумаченко І. В.	15	Chumachenko Igor	15
Юдін О. В.	109	Yudin Oles	109

НАУКОВЕ ВИДАННЯ

SCIENTIFIC PUBLICATION

Сучасний стан наукових досліджень та технологій в промисловості

Innovative technologies and scientific solutions for industries

Щоквартальний науковий журнал

Quarterly scientific journal

№ 2 (36), 2026

No. 2 (36), 2026

Відповідальний секретар журналу *І.Г. Перова*
Відповідальний за випуск *А.А. Коваленко*
Відповідальний за ліцензування *В.В. Косенко*
Редактор *Л.В. Кузьміна*
Комп'ютерна верстка *Л.Ю. Свєтайло*

Responsible secretary of journal *I. Perova*
Responsible for release *A. Kovalenko*
Responsible for licensing *V. Kosenko*
Editor *L. Kuzmina*
Computer layout *L. Svietailo*

Формат 60×84/8. Умов. друк. арк. 26,6.
Тираж 150 прим.

Format 60×84/8. Conventional printed sheets 26,6.
Edition of 150 copies.

Віддруковано з готових оригінал-макетів
в типографії ФОП Андреев К.В.
Єдиний державний реєстр юридичних осіб
та фізичних осіб-підприємців.
Запис №24800170000045020 від 30.05.2003.

Printed from ready-made original layouts
in the printing house
of Individual Entrepreneur Andreev K.V.
Unified State Register of Legal Entities
and Individual Entrepreneurs.
Entry No. 24800170000045020 of 30.05.2003.

61157, Харків, вул. Акад. Богомольця, 9, кв. 50,
тел. +38 (063) 993-62-73
e-mail: ep.zakaz@gmail.com

fl. 50, 9, Acad. Bogomolets Str., Kharkiv, 61157,
тел. +38 (063) 993-62-73
e-mail: ep.zakaz@gmail.com