



Дніпровський національний університет імені Олеся Гончара



Інститут кібернетики ім. В.М. Глушкова НАН України



ІНН «Інститут прикладного системного аналізу»
НТУУ «КПІ ім. І. Сікорського»



Київський національний університет ім. Т. Шевченка



ІТ компанія DataArt



ІТ компанія MalevichStudio ОУ у Естонії

XXIII міжнародна науково-практична конференція

**МАТЕМАТИЧНЕ ТА ПРОГРАМНЕ
ЗАБЕЗПЕЧЕННЯ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ
(МПЗІС-2025)
*ТЕЗИ ДОПОВІДЕЙ***

**MATHEMATICAL SUPPORT AND SOFTWARE
FOR INTELLIGENT SYSTEMS
(MSSIS-2025)
*ABSTRACTS***



**19-21 листопада 2025 року
Дніпро, Україна**



Дніпровський національний університет імені Олеся Гончара



Інститут кібернетики ім. В.М. Глушкова НАН України



ННК «Інститут прикладного системного аналізу»
НТУУ «КПІ ім. І. Сікорського»



Київський національний університет ім. Т. Шевченка



ІТ компанія DataArt



ІТ компанія MalevichStudio ОУ у Естонії

XXIII міжнародна науково-практична конференція

**МАТЕМАТИЧНЕ ТА ПРОГРАМНЕ
ЗАБЕЗПЕЧЕННЯ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ
(МПЗІС-2025)
*ТЕЗИ ДОПОВІДЕЙ***

**MATHEMATICAL SUPPORT AND SOFTWARE
FOR INTELLIGENT SYSTEMS
(MSSIS-2025)
*ABSTRACTS***



**19-21 листопада 2025 року
Дніпро, Україна**

Міжнародний науковий комітет

М. Згуровський	– академік НАН України, Україна
І. Сергієнко	– академік НАН України, Україна
О. Хіміч	– академік НАН України, Україна
А. Чикрій	– академік НАН України, Україна
Ю. Крак	– член-кореспондент НАН України, Україна
Н. Панкратова	– член-кореспондент НАН України, Україна
С. Яковлев	– член-кореспондент НАН України, Україна
V. Deineko	– професор, Англія
Y. Melnikov	– професор, США
O. Blyuss	– професор, Англія
T. Romanova	– професор, Англія
M.Polyakov	– засновник компанії Noosphere Ventures USA, Inc, США

М 34 Математичне та програмне забезпечення інтелектуальних систем (МПЗІС-2025):

Тези доповідей XXIII Міжнародної науково-практичної конференції, Дніпро, 19-21 листопада 2025 р. / Під загальною редакцією О.М. Кісельової. – Дніпро: ДНУ, 2025. – 324 с. – Текст: укр., англ.

Щорічна міжнародна науково-практична конференція «Математичне та програмне забезпечення інтелектуальних систем» (МПЗІС) є актуальним та затребуваним форумом фахівців з прикладної математики, інтелектуальних систем прийняття рішень, системного аналізу, новітніх інформаційних технологій. Конференція демонструє актуальність проблем розробки, створення та впровадження нового покоління систем управління та обробки інформації – інтелектуальних систем, а також тематики автоматизації управління в умовах прискореного розвитку математичної теорії і застосувань інтелектуальних систем і середовищ, їх широкого впровадження в повсякденну практику. Тези конференції публікуються в авторській редакції.

М 34 Mathematical support and software for intelligent systems (MSSIS-2025): Abstracts of the XXIII International scientific and practical conference, Dnipro, November 19-21, 2025 / Under the general editorship of E.M. Kiseleva. – Dnipro: DNU, 2025. – 324 p. – Text: ukrainian, english.

The annual international scientific and practical conference "Mathematical support and software for intelligent systems" is a relevant and popular forum of specialists in applied mathematics, intelligent decision-making systems, system analysis and the latest information technologies. The conference demonstrates the relevance of the problems of development, creation and implementation of a new generation of information management and processing systems - intelligent systems, as well as of the topics of control automation in the context of accelerated development of mathematical theory and applications of intelligent systems and environments, their widespread adoption in everyday practice. Conference abstracts are published in the author's edition.

Оргкомітет:

голова	<u>Кісельова Олена Михайлівна</u> – член-кореспондент НАН України, декан факультету прикладної математики та інформаційних технологій Дніпровського національного університету імені Олеся Гончара, д-р фіз.-мат.наук, професор
вчений секретар	<u>Кузенков Олександр Олександрович</u> – канд.фіз.-мат.наук
члени	О.Г. Байбуз – д-р тех.наук; Н.А. Гук – д-р фіз.-мат.наук; Л.Л.Гарт – д-р фіз.-мат.наук; О.М. Притоманова – д-р фіз.-мат.наук; В.А. Турчина – канд.фіз.-мат.наук; Т.А. Зайцева – канд.тех.наук; Н.В. Балейко – зав.лабораторією систем та методів штучного інтелекту; Н.Є. Яценко – пров.інж.
Адреса Оргкомітету:	Дніпровський національний університет імені Олеся Гончара Кафедра обчислювальної математики та математичної кібернетики пр. Науки, 72, Дніпро, 49010, Україна телефон: +38(067)772-11-51 e-mail: mpzis_dnu@ukr.net URL : mpzis.dnu.dp.ua

AN IMPROVED METHOD OF GRAYSCALE IMAGES CONTRAST ENHANCEMENT BASED ON POWER TRANSFORMATIONS

Akhmetshina L.G., akhmlu1@gmail.com, Yegorov A.A.,
for____students@ukr.net, Nikishyna O.Y. nikishyna_o.dnu.edu.ua
Dnipro National University named by O. Honchar

Introduction. When processing images, it is often necessary to enhance contrast, and the complexity of this task is due to both the differences in the characteristics of the original images and the possible presence of control parameters in the used methods. Therefore, the development of automatic contrast enhancement methods is relevant, an example of which is presented in the work [1].

Statement of the problem. This work proposes an improved method of grayscale images automatic contrast enhancement, which allows achieving higher image quality based on using adaptive coefficients in power transformations.

In the proposed method, as in the original [1], the contrast enhancement is performed by scaling the original image I to the range $[0,1]$ and the next power transformations usage:

$$I_{x,y}^1 = \left(I_{x,y} \right)^{1-I_{x,y} / \left(k_1 + \text{sgn}(I_{x,y} - \bar{I}) (I_{x,y})^{1-I_{x,y}} \right)}, \quad (1)$$

$$I_{x,y}^{out} = \left(I_{x,y}^1 \right)^{1-I_{x,y}^1 / \left(k_2 + (I_{x,y}^1)^{1-I_{x,y}^1} \right)}, \quad (2)$$

where $\bar{I}$ is the average value of the image I , and coefficients k_1 and k_2 are not constant values now (3 and 0.5, correspondingly), and are calculated by the next formulas: $k_1 = \left(1 + (0.5 + I_{x,y} + \bar{I}) / 2 \right)^2$; $k_2 = (I_{x,y}^1 + \bar{I}_1) / (1.5 + I_{x,y}^1 + \bar{I}^1)$, where $\bar{I}^1$ is the average value of the image I^1 .

The experimental results were obtained on example of grayscale images processing, an example of which is shown on Fig. 1a. Fig. 1b and 1c show the results of contrast enhancement using the original and improved methods, respectively (Fig. 1d – 1f show the histograms of the original and obtained

images). Although the differences are not particularly noticeable visually (the proposed algorithm provides a higher level of brightness for the forest and the illuminated part of the field), an analysis of the histogram reveals that the improved method provides a higher level of both brightness and contrast.

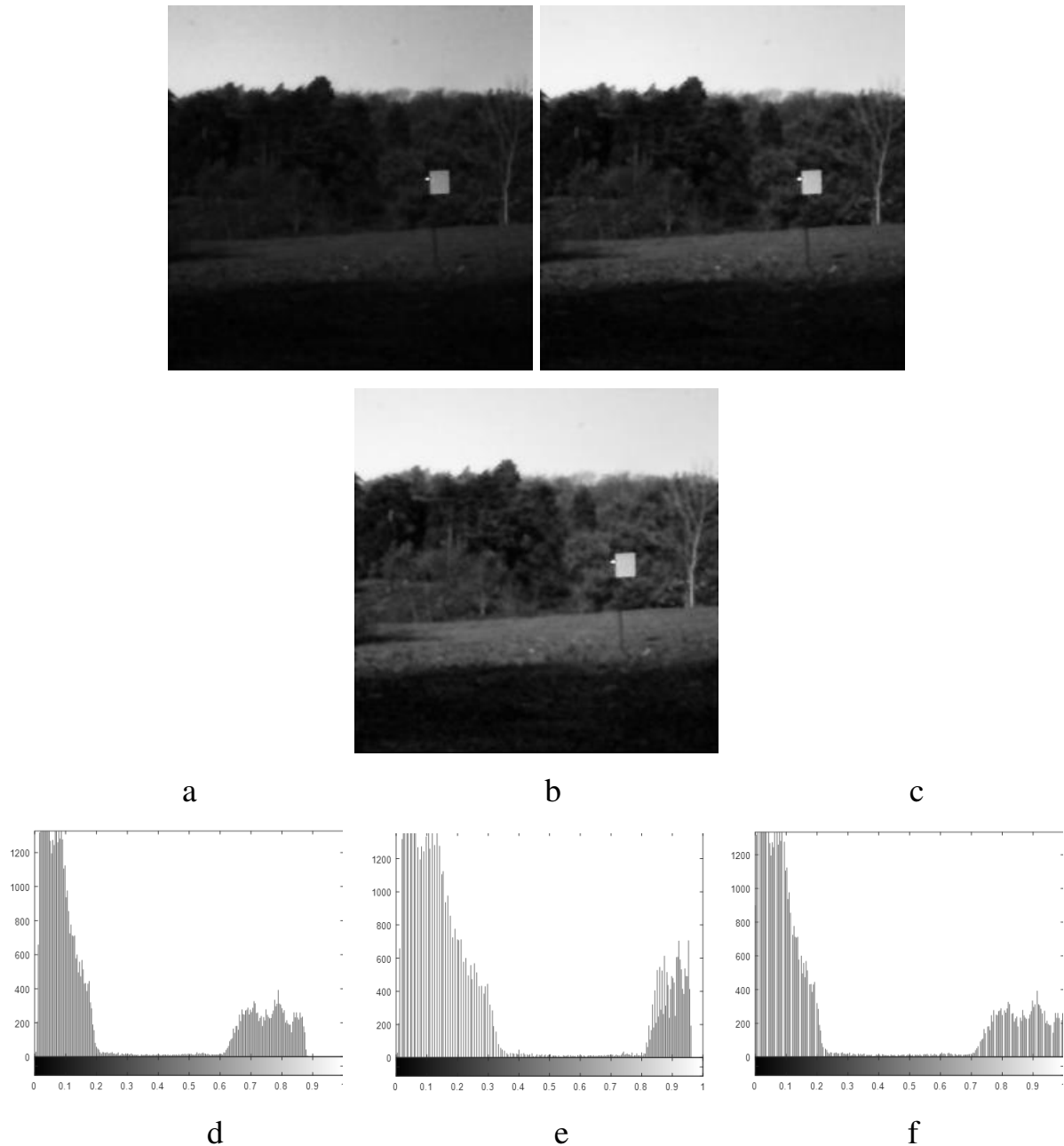


Figure 1 – Processing results: a – source image (256x256); contrast enhancement by: b – original, c – improved methods; d, e, f – histograms of images in Fig. 1a – 1b, respectively

References

1. Lyudmila Akhmetshina, Artyom Yegorov, Olexandra Nikishyna Grayscale Images Automatic Quality Enhancement Based on Power Transformations Usage. *INFORMATION TECHNOLOGIES AND COMPUTER MODELING*: proceedings of the International Scientific Conference c. Ivano-Frankivsk, 2025, May 20–23, 2025. Ivano-Frankivsk, 2025, P. 108–110.

FORECASTING NONLINEAR NON-STATIONARY PROCESSES IN FINANCE

Bidyuk P.I., Tertychnyi R.V., Chuprin D.S., denischuprin9619@gmail.com

National Technical University of Ukraine

«Igor Sikorsky Kyiv Polytechnic Institute», Kyiv, Ukraine (<https://kpi.ua>)

Relevance. Financial processes are a complex field of forecasting with a high level of uncertainty and risk. This is caused by their intricate structure, nonlinearity, and non-stationarity. Traditional methods do not take these features into account. Therefore, the development of information systems that apply modern data analysis methods to improve forecast accuracy is a highly relevant direction.

Objective. To develop an algorithm for modeling and forecasting nonlinear non-stationary financial processes using a selected set of models. Based on the obtained forecasts, to perform a comparative analysis in order to identify the strengths and weaknesses of the chosen models and select the most optimal one.

Methods. A comprehensive approach was applied: data preprocessing, testing for stationarity, seasonality, and heteroscedasticity. For comparison, ARIMA and GARCH models were selected, as well as an LSTM neural network designed to capture nonlinear dependencies. To assess model adequacy, the following criteria were used: DW-statistic, AIC, MSE, MAE, RMSE, MAPE.

Results. For approbation, the time series of AMD stock closing prices (2017–2022) was chosen. The comparative analysis showed that ARIMA adequately describes trends but has limitations under high volatility; GARCH effectively captures volatility clusters but does not always accurately predict the direction of changes; LSTM demonstrated the highest accuracy in reproducing dynamics, minimizing discrepancies with actual data.

Type of model	Adequacy criteria		Forecast assessment		
	R^2	DW	MSE	$RMSE$	MAE
$ARIMA(2,1,2)$	0.8433	1.9820	6.0518	2.4600	1.7882
$GARCH(1,1)$	0.3458	2.1010	0.0019	0.0435	0.0341
$LSTM$	0.8427	1.9978	6.3296	2.5159	1.8653

Conclusions. The use of classical models in combination with neural networks significantly improves the accuracy of financial process forecasting. This approach is a promising tool for risk management and decision-making under uncertainty. Future research will focus on applying a wider range of machine learning methods and automating the analysis.

ESTIMATING TASKS IN SCRUM PROJECTS BASED ON FUZZY LOGIC

Dzhenkova M., dzhenkova.mariya@gmail.com,

Sheveleva A., sheveleva@dnuniversity.edu.ua

Oles Honchar Dnipro National University (www.dnu.dp.ua)

In scrum estimation directly influences planning, so it is crucial for business. The estimation shows task complexity and team's effort to complete the task. Story points are one of the commonly used units and are specific values that are used to estimate effort required to complete tasks. The measurement is based on team experience, workstyle, and competencies.

Since estimation is highly affected by human thinking, it is proposed to apply fuzzy logic to automatically estimate tasks based on previously estimated tasks and input task parameters.

The input object for the system is a task with title, description, type; and list of previously completed tasks with additional story points parameter. The output – is some value from Fibonacci sequence [3], in our case {1, 2, 3, 5, 8, 13}.

First part of the algorithm analyzes input task and historical data. Then calculated parameters such as text complexity, context, past estimation average and sentiment score are processed in fuzzy inference system. We use triangular membership function to fuzzify input, process it with set of fuzzy rules, and use center of gravity method in defuzzification step.

We developed software that implements the described algorithm. Client side uses React and is deployed to Jira [2], server side – .NET 8 [1]. The developed algorithm and software show that automatic system based on fuzzy inference can provide more reliable mechanisms for estimation.

References

1. Nagel, C. Professional C# and .NET 8. Wiley. 2024
2. Podeswa, H. The Agile Guide to Business Analysis and Planning: From Strategic Plan to Continuous Value Delivery. CRC Press. 2021
3. Schwaber, K., & Sutherland, J. The Scrum Guide: The Definitive Guide to Scrum, 2020

IMPLEMENTING DEFAULT PARAMETER VALUES IN GO PROGRAMMING LANGUAGE DIALECT

Forkert P. P., Ivanchenko M. G., fxposter@gmail.com

Oles Honchar Dnipro National University

Modern popular programming languages typically offer the ability to specify default parameter values [1] for function and method definitions. For example, C++, C#, JavaScript, Python, Ruby, and others provide this functionality to their developers. Some languages don't have this feature directly, but allow developers to mimic it. For example, Java achieves this via method overloading. JavaScript did not initially have this feature in the language. Still, since the runtime didn't check the number of provided arguments at the call site and all non-provided arguments were assigned to `undefined`, it was possible to implement similar functionality, and many libraries did so.

Providing similar functionality in Go requires either using the builder pattern or defining multiple functions with different names, which leads to a poor user experience.

This work demonstrates one approach to achieving this functionality in the GoNext language dialect, which transpiles to Go and aims to be its superset [2].

From the GoNext source code perspective, the syntax of this feature looks like this:

```
func compute(a int, b int = 1) {}
```

And the simplest way to transpile it to Go would be to replace all call sites that don't specify `b` with its default value:

```
func use() {  
    compute(10)  
    // transpiles to  
    compute(10, 1)  
}
```

However, this is a very limited approach, that immediately breaks as soon as the default parameter values start to refer to the values from the scope of the function definition:

```
import "fmt"  
import "utils"  
var global int = 10
```

```
func compute(
    a int,
    b int = a + 1, c int = b + 1,
    d func() = func() { fmt.Println(global) },
    g int = global, other string = utils.Other()) {}
```

It is impossible to put `a + 1` or `global` at the call site, because they can be defined in a different file or package and have a completely different scope and imports. Instead, the function presented above is transpiled into two different functions:

```
func compute(
    a int,
    b int, c int,
    d func(),
    g int, other) {}

func compute__default(
    a int,
    bOpt gn.Optional[int], cOpt gn.Optional[int],
    dOpt gn.Optional[func()],
    gOpt gn.Optional[int], otherOpt gn.Optional[string]) {
    var b int = a + 1
    if bOpt.Exists {
        b = bOpt.Value
    }
    var c int = b + 1
    if cOpt.Exists {
        c = cOpt.Value
    }
    var d func() = func() {}
    if dOpt.Exists {
        d = dOpt.Value
    }
    var g int = global
    if gOpt.Exists {
        g = gOpt.Value
    }
    var other string = utils.Other()
    if otherOpt.Exists {
        other = otherOpt.Value
    }
    compute(a, b, c, d, g, other)
}
```

And at the call site all calls that specify all values are transpiled as is and calls that specify only some of the values are replaced with calls to `compute__default`:

```
func use() {
    compute(1, 2, 3, func() {}, 5, "") // transpiled as is
    compute(1, 5)
    // transpiles to
    compute__default(1, gn.Some(5), gn.None(), gn.None(), gn.None(), gn.None())
}
```

REFERENCES

1. Default argument. https://en.wikipedia.org/wiki/Default_argument
2. Forkert P. P., Sydorova M. G. Integrating full-featured enums into Go programming language // Актуальні проблеми автоматизації та інформаційних технологій. – Д: Ліра, 2023. Т.27, С. 3-16.

INTELLIGENT CONTROL OF AIR QUALITY BASED ON MASS BALANCE MODEL AND EXPERT RULES

Huk K., huk_k@365.dnu.edu.ua, Sheveleva A.

Oles Honchar Dnipro National University (www.dnu.dp.ua)

Maintaining proper indoor air quality is a key factor influencing human health, comfort, and productivity. The airtightness of modern buildings and the presence of volatile compounds emitted from finishing materials increase the need for adaptive ventilation systems. Such systems must ensure an optimal balance between air quality and energy efficiency. This study presents a hybrid expert system for indoor microclimate control that integrates physical modeling based on the mass balance of air exchange with logical decision-making rules and optimization methods.

Traditional ventilation systems operate at fixed air exchange rates and do not account for variations in room occupancy or pollution intensity. Advances in sensor technologies and intelligent control methods enable the development of adaptive systems that regulate air supply according to current environmental conditions. Contemporary research identifies several major approaches to air quality control: models based on physical equations of mass and energy balance; data-driven methods, including those employing machine learning; and expert systems that use rule-based and logical reasoning for decision-making.

Each of these approaches has distinct advantages and limitations. Physical models provide transparency and interpretability but require precise parameter identification, whereas data-driven systems depend on the quality of sensor measurements and often function as “black boxes.” The most promising direction is the integration of these methods within hybrid intelligent systems that combine the accuracy of mathematical models with the flexibility of expert rule sets.

The developed expert system integrates a mass-balance air exchange model with several types of expert control rules. The system architecture includes a sensor module for measuring pollutant concentrations, temperature, and humidity; an aerodynamic module that describes air movement through supply and exhaust ducts; a heat and mass transfer subsystem that simulates pollutant distribution; and a control module that regulates fan speed. Based on the modeling results, the system determines the minimum required air exchange rate to maintain pollutant concentrations within permissible limits. This value serves as a baseline constraint for energy optimization and the formulation of expert control rules.

To improve energy efficiency, an optimization problem is formulated to minimize fan power consumption while maintaining normative air quality parameters and adhering to equipment operating limits. Solving this problem determines the optimal air supply rate that sustains indoor air quality with minimal energy expenditure. Several categories of expert rules are introduced: emergency rules that activate maximum ventilation in case of excessive pollution; regulatory rules that ensure compliance with sanitary standards; predictive rules that increase airflow in response to rising pollutant accumulation rates; optimization rules that implement computed energy-efficient control strategies; protective meta-rules that prioritize safety over energy saving.

The use of the mass balance model provides a physically grounded foundation for decision-making, while the rule-based layer enables system adaptability to dynamic conditions and resilience to sensor errors. Experimental simulations using synthetic data demonstrated that the proposed system maintains air quality within regulatory limits while reducing average energy consumption by 15–20% compared to conventional control methods.

ADAPTIVE APPROACH TO THE DEVELOPMENT OF EXPERT SYSTEMS BASED ON WEIGHTED RULES AND PATTERN ANALYSIS OF SYSTEM STATES

Huk N., huk_n@365.dnu.edu.ua, Yehoshkin D., yehoshkin_d@365.dnu.edu.ua

Oles Honchar Dnipro National University (www.dnu.dp.ua)

The efficiency of modern information systems largely depends on their ability to automatically analyze large volumes of data and detect patterns that indicate specific states or deviations in system performance. In this context, expert systems acquire particular importance, as they are capable not only of identifying behavioral patterns of systems but also of generating recommendations for eliminating undesired conditions or optimizing processes. Pattern recognition enables a transition from passive monitoring to active, knowledge-based control, making such systems intelligent tools for decision support.

A review of scientific publications indicates that most current studies focus on individual aspects of expert system design—either on methods for detecting patterns in data or on the development of decision-making mechanisms. However, the problem of integrating these two directions into a unified adaptive model that combines data-derived knowledge with its formalized interpretability remains insufficiently explored. The proposed approach seeks to bridge this gap by combining the analytical power of machine learning with the interpretability of expert models.

The theoretical foundation of this approach is the notion that any process can be described as a sequence of states characterized by specific features. The combination of such features forms a pattern that can be recognized by the system and associated with a corresponding action or recommendation. Formally, each pattern is represented as a set of system state attributes, while the rule base defines how combinations of these attributes lead to conclusions about the system's condition. It is proposed to construct a training dataset from historical records and expert decisions that reflect various system states. Based

on these data, patterns are generated using clustering, anomaly detection, or temporal pattern recognition methods. These patterns are subsequently transformed into formal knowledge units, represented as systems of rules of several types—including production, probabilistic, fuzzy, hierarchical, and adaptive rules. This hybrid approach makes it possible to combine the transparency of expert logic with the flexibility of statistical analysis. An essential component of the method is the validation of generated rules for correctness and consistency, achieved through formal verification or testing on control datasets.

A distinctive feature of the proposed approach is the introduction of weighting coefficients for each rule, which reflect both the degree of confidence in the knowledge source and the statistical reliability of the identified patterns. These weights can be dynamically updated during system operation using gradient learning. In this way, the system acquires the ability to self-learn, which makes it possible to improve the accuracy of recommendations based on new data.

The practical significance of the proposed approach is illustrated by the example of an expert system designed for the analysis of computer memory snapshots. In such a system, the identified patterns may indicate data duplication, type inconsistencies, or inefficient resource utilization. Based on these observations, the system generates recommendations for error correction or memory optimization. For instance, if repeated objects with close memory addresses and identical contents are detected, the system may recommend releasing the duplicate or optimizing the data structure.

It is expected that further development of this approach will contribute to the creation of universal adaptive decision-support systems applicable to a wide range of technical and informational domains.

Acknowledgments. The research was partially funded by the Ministry of Foreign Affairs of the Czech Republic under project 25-PKV-V-UM-004 entitled "Development of Education, Research, and International Cooperation at Oles Honchar Dnipro National University (DNU)" implemented by Charles University and DNU.

MIXED REALITY WORKSPACES: REDEFINING REMOTE/HYBRID PRODUCTIVITY VIA XR HEADSETS

Ivanchenko M., Bondarenko B., bogdan.bondarenko99@gmail.com

Oles Honchar Dnipro National University (www.dnu.dp.ua)

In today's remote and hybrid work era, advanced XR headsets featuring mixed-reality (MR) capabilities are emerging as a transformative platform for team collaboration and productivity. Unlike early headsets such as the Meta Quest 2, which largely immersed users in fully virtual spaces and had a very poor quality of passthrough, modern MR devices allow individuals to stay aware of their physical environment while interacting with digital content. This dual awareness opens new possibilities for how distributed teams work together, bond, and connect.

A core strength of MR workspaces lies in bridging geographical separation. Traditional conferencing often fails to replicate the spontaneous and informal interactions typical of a shared office. More often than not – web-cameras are not even enforced, leading to teams getting more and more stranded, with some colleagues having no idea what their other colleagues look like. MR environments enable participants to interact in a shared virtual space - via avatars, while still seeing their desk or keyboard and keeping physical cues. This preserves spatial orientation, reduces disorientation, and gives a stronger feeling of co-presence. Research shows that VR workplaces have the potential to enable both teamwork and task work in remote knowledge work scenarios [1].

MR also supports social connection and team culture, which often suffers in remote settings. Informal chats, ad-hoc interactions and other moments contribute significantly to team identity and morale. MR headsets can create dedicated virtual lounges, shared project zones, or spatial group games—helping distributed teams feel more aligned and less isolated.

Nevertheless, several challenges must be addressed. Firstly, ergonomics and comfort remain a barrier: current headsets may be on a heavier side and

cause fatigue during extended use. Secondly, hardware and space requirements remain non-trivial: users may need plenty of room for tracking, good lighting, and strong connectivity, conditions not guaranteed in all home environments. Thirdly, workflow and software integration are still catching up: many productivity tools have yet to adapt for MR interactions and organizational culture/training need time to evolve. Price also stays a big deterring factor.

Looking forward, MR workspaces hold significant promise. As headsets become lighter, more affordable and better connected, distributed teams may adopt virtual offices as standard rather than a supplemental tool. The distinction between physical and digital workspaces could blur: workers may start at their physical desk and seamlessly move into shared virtual collaborative environments, then back out. Crucially, MR offers the chance to shift remote work from “grid of faces on a screen” to spatial, interactive and presence-rich collaboration.

In summary, mixed reality workspaces, enabled by the latest XR headsets - offer a compelling evolution for remote and hybrid productivity. They present enhanced presence, richer collaborative layouts and stronger socio-cultural connectivity across distributed teams. While challenges remain, the direction is clear: as MR technology matures, it promises a new paradigm of how people work together remotely.

REFERENCES

1. Aufegger L, Elliott-Deflo N, Nichols T. Workspace and Productivity: Guidelines for Virtual Reality Workplace Design and Optimization. *Applied Sciences*. 2022; 12(15):7393. <https://doi.org/10.3390/app12157393>

SOFTWARE IMPLEMENTATION FOR VISUALIZATION OF OPTIMAL PARTITIONING SET RESULTS ON MAPS

Kiseleva¹ O., Prytomanova² O., Zhurava¹ D.

den.itsmart@gmail.com

¹*Oles Honchar Dnipro National University (www.dnu.dp.ua)*

²*Kyiv National Economic University named after Vadym Hetman (www.kneu.edu.ua)*

The problem of optimal partitioning of sets with fixed centers is widely applied in logistics, urban planning, and resource allocation. The classical mathematical formulation involves dividing a bounded set $\Omega \subset \mathbb{R}^n$ into N non-overlapping subsets $\Omega_1, \dots, \Omega_n$ with predetermined centers $\tau_1, \dots, \tau_n$, minimizing the functional that accounts for distances from points to centers. However, practical application of such algorithms requires specialized software tools capable of processing real geographical data and visualizing results in an accessible format for decision-makers.

The developed software implementation is built on a modular architecture using Python with integration of specialized libraries for geospatial analysis. Shor's r -algorithm for optimal partitioning with configurable metrics (Euclidean distance, transport accessibility via road networks) serves as a key component of the core computational module. The geographic data processing module utilizes GeoPandas and Shapely libraries for working with vector geographic data in standard formats (Shapefile, GeoJSON), supporting coordinate system transformations (WGS84, UTM, Web Mercator). The system includes a population density modeling component that accounts for spatial distribution of demand based on raster data or mathematical models of urban centers. To enhance computational performance, parallel computing mechanisms using multiprocessing are implemented, achieving 5-6x speedup on multi-core processors. The visualization module generates interactive maps using the Folium library with capability to display partitioning zones, service centers, and analytical layers.

The software supports batch processing of multiple facility networks, comparison of different partitioning scenarios, and export of results in standard GIS formats for further analysis in professional geographic information systems. Practical applications include optimization of retail network placement, planning of medical and educational facility locations, and analysis of administrative service accessibility in territorial communities.

OPTIMIZATION OF TECHNOLOGICAL DESIGN TASKS OF COMPLEX SYSTEMS

Kiselyova O.M. ¹, kiseleva47@ukr.net,

Stroieva V.O. ², vikastroeva@ukr.net, **Tarasiuk O.S.** ², sashko98man@gmail.com

¹ *Oles Honchar Dnipro National University*

² *Dniprovsky State Technical University*

An analysis of the current state of research in the field of optimization of technological design problems demonstrates the intensive development of methodological approaches and algorithmic solutions. In particular, work [1] presents an innovative method for selecting the type of components, topological and dimensional optimization for complex structures, based on an integrated approach to multi-criteria design. The mentioned research demonstrates a successful combination of different types of optimization to obtain optimal characteristics of technical systems. In the work [2], a detailed review of the use of optimization methods in automotive and agricultural engineering is made with an emphasis on the mathematical basis of modern optimization methods. This work focuses on the importance of an interdisciplinary approach to solving technical problems and the need to develop a common methodological framework. Of particular note is the work [3], devoted to multilevel informed optimization via decomposed kriging for large design problems under uncertainty. The presented approach demonstrates the effectiveness of decomposition methods for solving multidimensional optimization problems.

The presented work considers the results of research into optimization problems of technological design of complex systems, which were obtained using the methodological tools of the theory of optimal partitioning. Namely, mathematical models of a number of relevant applied design problems were developed. In particular, the problems of designing an irrigation system (intelligent "smart" irrigation system); optimal placement of drone delivery

nodes; optimal placement of environmental monitoring stations; optimal management of agricultural land by modern robots; optimal design of an energy network with renewable energy sources.

The above problems in their mathematical formulation can be reduced to problems of optimal partitioning of sets in various interpretations, according to the developed classification of such problems. And their study became possible thanks to the strictly grounded theory of optimal partitioning of sets and a powerful system of methods and algorithms of such studies, developed on the basis of this theory. Some mathematical models have been numerically investigated as a problem of designing an irrigation system. Others are at the stage of refinement and testing. The results obtained are important from the point of view of practical implementation and can serve as the basis for modeling more complex design systems.

References

1. Li, J., Liu, S., Huang, H., & Chen, S. (2025). An engineering method of component type selection, topology and size optimization for complex structures. *Structural and Multidisciplinary Optimization*, 68(5). DOI: <https://doi.org/10.1007/s00158-025-04015-w> Retrieved from <https://link.springer.com/article/10.1007/s00158-025-04015-w>
2. Zhao, W., Duan, L., Ma, B., Meng, X., Ren, L., Ye, D., & Rui, S. (2025). Applications of Optimization Methods in Automotive and Agricultural Engineering: A Review. *Mathematics*, 13(18), 3018. DOI: <https://doi.org/10.3390/math13183018> Retrieved from <https://www.mdpi.com/2227-7390/13/18/3018>
3. Ampellio, E., Gjorgiev, B., & Sansavini, G. (2025). Multi-level informed optimization via decomposed Kriging for large design problems under uncertainty (Version 1). arXiv. DOI: <https://doi.org/10.48550/ARXIV.2510.07904> Retrieved from <https://arxiv.org/abs/2510.07904v1>

WIRELESS COMPUTER NETWORKS SIMULATION

Krasnoshapka D.V., dimakrasnoshapka@yahoo.com

Zolotko K.Y., zolt66@gmail.com

Oles Honchar Dnipro National University (www.dnu.dp.ua)

Computer network modeling is of great importance in the development process, allowing to estimate the performance and other network parameters in advance.[1]

In addition, modeling reduces development costs and makes it possible to prevent some errors in network operation. Computer network modeling is carried out using such tools as Matlab, GNS3, Opnet Modeler, Cisco Packet Tracer, etc. [2, 3]

The purpose of the work is to simulate a wireless local area network using Cisco Packet Tracer [4], which uses the IEEE 802.11 standard. The network equipment includes a model of a WiFi router of the Home Router-PT-AC type and a model of a Leptop-PT laptop.

The Home Router-PT-AC configuration is performed by sequentially configuring the Internet, Lan and Wireless 2.4G interfaces.

To configure the local network (Network Setup), select the default IP address 192.168.0.1, mask 255.255.255.0, DHCP server enabled, address distribution starts at 192.168.0.100.

Settings on the Wireless tab – Basic Wireless Settings. Select the basic settings: mode (mode) – mixed, network identifier (SSID) – HomeWiFi, channel width (Radio Band) – auto, frequency - 1-2.412HGz, network visibility (SSID Broadcast) – visible (enable).

Settings on the Wireless tab – Basic Wireless Settings. Select the basic settings: mode (mode) – mixed, network identifier (SSID) – HomeWiFi, channel width (Radio Band) – auto, frequency - 1-2.412HGz, network visibility (SSID Broadcast) – visible (enable).

Go to the Wireless Security tab. We select the WPA2 Personal encryption mode, the AES encryption algorithm, the keyword for the selected encryption mode of at least 8 characters - passpass.

In the laptop settings for the Wireless interface, do not forget to set the SSID - HomeWiFi. All other settings remain unchanged. If the settings are made correctly, a dotted line will appear between the Wi-Fi router and the laptop.

Checking the WiFi router settings. Open Laptop0, go to Desktop-Command Prompt and execute the ipconfig command. The figure shows that the laptop received the correct IP address - 192.168.0.100.

Methods such as the Add Simple PDU tool and the ping command showed the full functionality of the network model.

Modeling a wireless local area network allows you to study its behavior for different settings of network equipment. Further research can be conducted to simulate computer internetworks.

References

1. SIRAJ, Saba; GUPTA, A.; BADGUJAR, Rinku. Network simulation tools survey. *International Journal of Advanced Research in Computer and Communication Engineering*, 2012, 1.4: 199-206.
2. RAJARAM, Madhupreetha L., et al. Wireless sensor network simulation frameworks: A tutorial review: MATLAB/Simulink bests the rest. *IEEE Consumer Electronics Magazine*, 2016, 5.2: 63-69.
3. ARVIND, T. A comparative study of various network simulation tools. *International Journal of Computer Science & Engineering Technology (IJCSET)*, 2016, 7.08: 374-378.
4. TRACER, Cisco Packet. Cisco packet tracer. URL: <http://www.cisco.com/web/learning/netacad/coursecatalog/PacketTracer.html>, 2013.

GEOMETRIC DESIGN OF NANOCOMPLEXES: MATHEMATICAL AND COMPUTER MODELLING

Letuta N.¹, Dubinskyi, V.¹, Romanova T.^{1,2}, Gutierrez Luis³

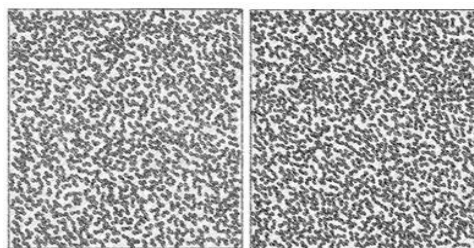
tetiana.romanova@nure.ua

¹*Kharkiv National University of Radio Electronics, Kharkiv, Ukraine*

²*University of Leeds, Leeds, UK*

³*Autonomous University of Nuevo Leon, Monterrey, Mexico*

Smart systems and data-driven services have the potential to answer medical needs in nanomedicine to get faster to the clinic and manufacturing stage. A non-standard packing problem of molecular clusters for the geometric design of nanocomplexes is considered. Each placement object (simplified model of nanocomplex Dox-C₆₀ [1]) is a disconnected set with a spheroidal core and two identical spheroidal components (molecules) that are allowed moving only along the core sphere orbit at the given distance. In addition, allowable distances between each pair of nanoparticles, are given. All objects could be freely moving within the given 3D volume (cuboidal container). The packing problem is aimed at maximizing the number of nanoparticles that could fit the given container, considering distance constraints. The problem is formulated as a mixed integer nonlinear programming model. A fast heuristic algorithm is proposed. Computational results are provided. Figure shows samples of design nanocomplex with 900/1000 clusters found by the proposed algorithm for 247/282 sec. Our findings may contribute to solve molecular packing problems which play a role in synthesis, production as well as formulation and medication of the complexed drug.



[1] R.R. Panchuk, S.V. Prylutska, V.V. Chumakl, N.R. Skorokhyd, L.V. Lehka, M.P. Evstigneev, Y.I. Prylutskey, W. Berger, P. Heffeter, P. Scharff, U. Ritter, R.S. Stoika, Application of C60 Fullerene-Doxorubicin Complex for Tumor Cell Treatment In Vitro and In Vivo, J Biomed Nanotechnol. 11 (2015) 1139–1152.

PROVING SORTING ALGORITHMS IN THE DAFNY VERIFICATION SYSTEM

Levchenko Y.S., evgeniy.levchnko@gmail.com,

Khyzha O.L., alkhizha@gmail.com

Oles Honchar Dnipro National University (www.dnu.dp.ua)

Proving the correctness of mathematical formulas through rigorous reasoning is a fundamental component of mathematics. In programming, program code is the equivalent of mathematical formulas. Typically, verification of code correctness is mainly performed by running it against a finite suite of test cases. However, certain application domains and algorithms require formal (mathematical) proof. This work examines the correctness proofs for specific array sorting algorithms—namely, Insertion Sort and Merge Sort [1].

The well-known theoretical foundations of formal program verification involve transforming a formal specification (written in first-order predicate logic) and the corresponding program code into a first-order predicate using a predicate transformer — a weakest precondition predicate [2].

This transformation and the subsequent proof of the resulting predicate is a complex and laborious task, making program verification largely impractical for real-world applications. This difficulty has stimulated the development of automated formal program verification systems. In this work, the proofs were carried out using the Dafny verification system [3].

For sorting algorithms, a correctness proof must confirm two key properties: that the output array is ordered and the original set of elements is preserved. In Dafny, these properties are expressed as predicates and applied to verify algorithms such as a simple selection sort or a more advanced merge sort.

As shown in Figure 1, the Dafny static verifier confirmed the correctness of the selection sort algorithm. Since a standard implementation of the selection sort algorithm is not included in Dafny's library, the necessary invariants were formulated independently.

```

30 method SelectionSort(a: array<int>)
31   modifies a ensures permutation(a[..], old(a[..]), a.Length) && sorted(a, 0, a.Length)
32 {
33   var N, i := a.Length, 0;
34   while i < N
35     invariant 0 <= i <= N && permutation(a[..], old(a[..]), a.Length) && sorted(a, 0, i) && partitioned(a, i
36     decreases N - i + 1
37   {--
46   }
47 }

```

Figure 1 – Formal verification of the selection sort algorithm

While reference [3] presents formal verifications of the merge sort algorithm, those implementations use linked lists as the primary data structure. This work employs arrays instead, which necessitated additional conditions to ensure correct index handling. Figure 2 illustrates a verified version of this algorithm.

```

24 method MergeSort(a: array<int>, lo: nat, hi: nat) returns (c: array<int>)
25   requires 0 <= lo <= hi <= a.Length
26   ensures sorted(c, 0, c.Length) && multiset(a[lo..hi]) == multiset(c[..])
27   decreases hi - lo
28 {
29   if hi - lo <= 1 {
30     c := new int[hi - lo];
31     if hi - lo == 1 { c[0] := a[lo]; }
32   } else {
33     var mid := (lo + hi) / 2; var left := MergeSort(a, lo, mid); var right := MergeSort(a, mid, hi);
34     assert a[lo..hi] == a[lo..mid] + a[mid..hi];
35     c := Merge(left, right);
36   }
37 }
38
39 method Merge(a: array<int>, b: array<int>) returns (c: array<int>)
40   requires sorted(a, 0, a.Length) && sorted(b, 0, b.Length)
41   ensures sorted(c, 0, c.Length) && multiset(c[..]) == multiset(a[..]) + multiset(b[..])
42 {
43   c := new int[a.Length + b.Length];
44   var i, j, k := 0, 0, 0;
45   while i < a.Length || j < b.Length
46     invariant 0 <= i <= a.Length && 0 <= j <= b.Length && k == i + j && k <= a.Length + b.Length
47     invariant sorted(c, 0, k) && multiset(c[..k]) == multiset(a[..i]) + multiset(b[..j])
48     invariant k == 0 || (i < a.Length && (j == b.Length || a[i] <= b[j]) ==> c[k - 1] <= a[i])
49     invariant k == 0 || (j < b.Length && (i == a.Length || b[j] < a[i]) ==> c[k - 1] <= b[j])
50     decreases a.Length - i + b.Length - j
51   {--
59   }
60   assert a[..i] == a[..] && b[..j] == b[..] && c[..k] == c[..];
61 }

```

Figure 2 – Formal verification of the merge sort algorithm

References

1. Thomas H. Cormen, Charles E. Leiserson, Ronald L. Rivest and Clifford Stein. *Introduction to Algorithms*, Fourth Edition. — The MIT Press, 2022. —1312 p.p
2. Gries, David. *The Science of Programming*. — Springer Science & Business Media, 2012. — 388 p.
3. K. Rustan M. Leino. *Program Proofs*. — The MIT Press (March 7, 2023). — 496 p.

OPTIMIZATION OF LOGISTICS ROUTES USING GRAPH THEORY ALGORITHMS

Lushnya K.K., Yehoshkin D.I.

Oles Gonchar Dnipro National University

Relevance. In the modern context of rapidly growing transportation volumes, the effective organization of logistics processes has become one of the key factors determining the competitiveness of enterprises. One of the main challenges is optimal route planning, which directly affects costs, delivery time, and customer service quality. Traditional methods often fail to account for the complexity of real transport networks and the variability of environmental parameters. The use of graph theory algorithms, particularly Dijkstra's and A*, enables the modeling of logistics systems that consider distance, time, and cost. This opens the door to the development of intelligent decision-support systems capable of automating the routing process and improving the efficiency of logistics operations.

The aim of the study is to develop a web application for logistics route planning that incorporates multi-criteria optimization, graph visualization, and comparison of the efficiency of various routing algorithms (Dijkstra, A*, Greedy, Random Walk).

Main Part. In modern logistics, the issue of route optimization is one of the most relevant, as the efficiency of transportation processes directly affects cost reduction, delivery timeliness, rational resource use, and environmental sustainability. As the number of orders and transport nodes increases, routing tasks become more complex, requiring the use of mathematical methods and algorithms capable of quickly finding optimal solutions under constraints.

One of the most effective approaches is the use of graph theory, which allows modeling a transport-logistics system as a weighted directed graph $G = (V, E, \omega)$, where the vertices V represent logistics points (warehouses, customers, transfer bases), and the edges E represent routes between them with a weight function $\omega : E$, determined by a combination of criteria. The weight of each edge is calculated by the formula: $\omega(e) = \alpha \cdot \text{distance}(e) + \beta \cdot \text{time}(e) + \gamma \cdot \text{cost}(e)$, where α, β, γ are weighting coefficients set by the user according to logistics priorities — minimizing distance, time, or cost.

Such a formulation enables modeling multi-criteria optimization problems, which brings the system closer to real-world operating conditions.

Two classical algorithms were chosen for optimal path search: Dijkstra's algorithm and A*.

Dijkstra's algorithm performs a sequential traversal of all graph vertices, computing the minimum cost path to each vertex. Its advantage lies in guaranteeing a globally optimal solution, though for large numbers of nodes, it has significant time complexity $O(|E| + |V| \log |V|)$ when implemented with a priority queue. The A* (A-star) algorithm uses a heuristic function $h(n)$ that estimates the "distance" from the current vertex to the goal, defined as: $f(n) = g(n) + h(n)$. This allows reducing the number of visited nodes and the average search time. The heuristic is typically the Euclidean or Manhattan distance in geographic coordinates.

To evaluate the efficiency of the algorithms, an experimental study was conducted on test graphs of various scales (from 10 to 100 vertices). Each graph modeled a hypothetical transport network with parameters for distance (km), time (min), and cost (UAH). The comparison criteria included: algorithm execution time, number of visited vertices, deviation from the global minimum (for A* with inexact heuristics), flexibility of adaptation to changes in weighting coefficients.

The results showed that:

- For small networks (up to 30 nodes), Dijkstra's algorithm demonstrated stable results with minimal deviations.
- As the network expanded to 100 nodes, A* reduced average search time by 35–45% without loss of accuracy.
- When coefficients α , β , γ changed, the system dynamically adapted — for example, when time was prioritized ($\beta > \alpha$, γ), routes were selected via highways, while when cost was prioritized (γ largest), routes were chosen through less congested roads.

An interactive web application was developed to serve as a decision-support system. Its interface allows the user to:

- create and edit route graphs (add nodes, connections, weights);
- visualize the network on a map using React-Leaflet;
- select the search algorithm (Dijkstra or A*);
- analyze comparative results in table or chart form.

The system was implemented using React, TypeScript, and Cytoscape.js for graph visualization. The computational logic was developed as a separate module, allowing integration with databases, geolocation APIs, or external transport platforms.

The application has open-source code available at: <https://github.com/lushnakata536/logistics-optimizer>.

Future work aims to expand its functionality through integration with mapping services and the implementation of machine learning algorithms to predict logistics flow dynamics.

During experiments, users could vary the weighting coefficients and observe route changes in real time. This approach enhances the interactivity of decision-making, enables evaluation of different scenarios, and helps select the optimal delivery strategy.

The obtained results confirm the effectiveness of applying graph theory and optimization algorithms in routing tasks. The developed system can serve as a tool for education, planning, and managerial decision support in transport logistics, as well as a foundation for further research in intelligent transportation management systems.

Conclusions. In modern conditions of transport infrastructure development, the effective organization of logistics processes has become a crucial component of business economic stability. The conducted research has shown that the application of graph-theoretical algorithms significantly improves route planning efficiency, reduces time and resource costs, and enhances the quality of managerial decisions.

Route optimization using Dijkstra's and A* algorithms ensures finding the shortest paths in complex networks while considering multiple criteria — distance, time, and transportation cost.

The developed web application has demonstrated its practical usability as a decision-support tool, providing visualization of logistics networks and flexible configuration of optimization parameters. The results confirm the promising combination of mathematical methods and modern information technologies in building intelligent logistics systems.

Reference

1. Dijkstra, E. W. (1959). A Note on Two Problems in Connexion with Graphs. *Numerische Mathematik*, 1, 269–271.
2. Ahuja, R. K., Magnanti, T. L., & Orlin, J. B. (1993). *Network Flows: Theory, Algorithms, and Applications*. Prentice Hall.
3. Cormen, T. H., Leiserson, C. E., Rivest, R. L., & Stein, C. (2009). *Introduction to Algorithms* (3rd ed.). MIT Press.
4. Russell, S., & Norvig, P. (2010). *Artificial Intelligence: A Modern Approach* (3rd ed.). Prentice Hall.
5. Holland, J. H. (1992). *Adaptation in Natural and Artificial Systems*. MIT Press.

**BUILDING A MATHEMATICAL MODEL OF AGENT INTERACTION
IN COLLECTIVE DECISION-MAKING TASKS****Lutsenko O.M.**, a.lutsenko@quantillions.com,**Trofimov O.V.**, atrof2222@gmail.com*Дніпровський національний університет імені Олеся Гончара (www.dnu.dp.ua)*

The process of collective decision-making is one of the central problems of modern multi-agent systems, where several autonomous entities — agents — must coordinate their actions to reach a common or mutually acceptable solution. Such systems are found in distributed robotics, financial markets, social modeling, and computer networks. The purpose of this research is to build a mathematical model describing the dynamics of agent interaction and the emergence of consensus under limited information and rationality.

The model assumes that each agent $i \in N$ possesses an individual utility function $U_i(x)$, where x represents the possible collective decision. Agents communicate through an interaction matrix $W = [w_{ij}]$, which determines the intensity and direction of influence among them. At each iteration, the state of the system is updated according to

$$x_i(t+1) = x_i(t) + a \sum_j w_{ij}(x_j(t) - x_i(t)),$$

where a is a learning rate or degree of openness to external opinions. This formulation corresponds to the DeGroot model of opinion dynamics, later extended to include stochastic perturbations and heterogeneous trust coefficients.

To analyze collective convergence, spectral decomposition of matrix W is performed, where the largest eigenvalue and its corresponding eigenvector determine the steady-state distribution of opinions. The presence of clustering, polarization, or oscillations depends on the topology of W and the presence of stubborn (non-updating) agents. This allows modeling both democratic consensus and conflict-based decision processes.

In financial and engineering applications, each agent can represent a subsystem or trader optimizing a local criterion. Collective learning is achieved via gradient descent over the aggregated utility $F(x) = \sum_i U_i(x)$, where agents exchange partial gradients and iteratively improve coordination. The system stability is estimated through Lyapunov functions, and simulation results demonstrate how varying the influence weights and communication delays affects convergence speed and equilibrium.

To validate the model, several numerical experiments were conducted in Python using libraries such as NumPy, NetworkX, and Matplotlib. Agents were initialized with random initial opinions drawn from a uniform distribution. It was observed that for strongly connected graphs, the group reached consensus exponentially fast, while for sparse or asymmetric structures, multiple subgroups with local consensus were formed. This behavior mirrors real-world social and financial systems, where partial agreement often replaces complete unanimity.

The mathematical model can be extended by incorporating reinforcement learning agents capable of updating their influence coefficients w_{ij} based on previous rewards, thereby creating an adaptive collective decision mechanism. Such hybrid models are promising for distributed optimization, swarm robotics, and market simulations where collective intelligence emerges from individual strategies.

References

1. **Quantillion Research Group** (2025). Multi-Agent Decision Models and Collective Intelligence Framework.
2. **Olfati-Saber R., Fax J., Murray R.** (2007). Consensus and cooperation in networked multi-agent systems.

EXPERIMENTAL EVALUATION OF SEMANTIC SEGMENTATION WITH MULTI-SCALE ADAPTIVE ATTENTION ON THE CAMVID DATASET

Nikolenko K., kateryna.nikolenko@stud.onu.edu.ua,
Moroz V., v.moroz@onu.edu.ua,
Odesa I.I. Mechnikov National University (Ukraine)

Introduction. Semantic segmentation enables pixel-wise classification of road scenes, critical for autonomous driving. Attention mechanisms enhance feature discrimination, but their effectiveness depends on robust integration and evaluation. This study applies a Multi-Scale Adaptive Attention Mechanism (MSAAM) within DeepLabv3+, using the CamVid dataset to assess segmentation accuracy across 11 generalized classes.

Dataset and Preprocessing. The CamVid dataset contains 701 labeled RGB images (960×720), split into training, validation, and test sets. To simplify learning, the original 32 semantic classes were reduced to 11 generalized categories. Data augmentation (flipping, rotation, color jitter) improved generalization. Labels were converted to grayscale masks and loaded via a custom PyTorch Dataset class.

Model Implementation. The model uses ResNet50 as a backbone, removing fully connected layers. MSAAM is inserted between the encoder and classifier, combining:

- **PSA (Pyramid Spatial Attention)** for multi-scale feature pooling.
- **AdaptiveAttention** for spatial and channel refinement.
- **Custom loss function** with variance penalty.

Training was performed over 25 epochs using AdamW optimizer and StepLR scheduler. Evaluation used mean Intersection over Union (mIoU) and visual inspection.

Results and Analysis.

- **Training loss** decreased steadily from 2.24 to 1.14, indicating convergence.
- **Validation loss** fluctuated (spikes at epochs 5 and 13).
- **Best validation mIoU** 0.1221 (epoch 9); final test mIoU: 0.1140.
- **Visual predictions** showed clear segmentation of large objects (road, cars), but poor recognition of small classes (pedestrians, bicyclists).

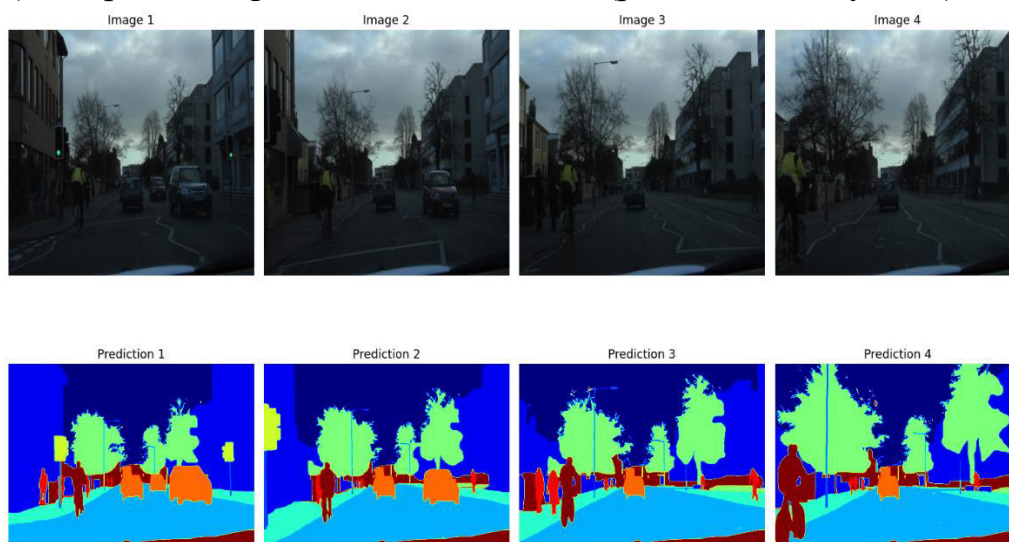


Fig. 1. Semantic segmentation results on CamVid test data: comparison between input images and predicted masks.

Conclusion. The model shows baseline segmentation ability, with attention mechanisms improving recognition of large objects. However, low mIoU and unstable validation loss reveal limited generalization. Future work may explore deeper backbones, longer training, and better attention weighting. CamVid remains a useful benchmark for urban scene segmentation.

References

1. E. Nikolenko, “Semantic Segmentation with Multi-Scale Adaptive Attention Mechanism (MSAAM),” GitHub repository, 2025. [Online]. Available: <https://github.com/EkaterinaNikolenko/semantic-segmentation-msaam>

INTEGRATING ARTIFICIAL INTELLIGENCE INTO ENGLISH LANGUAGE TEACHING

Novikova O.V., novikovaolga19911991@gmail.com

Suima I.P., irinasuima2017@gmail.com

Oles Honchar Dnipro national university

In recent years, the rapid advancement of artificial intelligence (AI) has opened up transformative opportunities in education. English language teaching (ELT) is one of the areas particularly well suited to benefit from AI, given the multiple skills involved (listening, speaking, reading, writing, grammar, pronunciation), the need for feedback, and the diversity of learner profiles [1; 2]. Integrating AI into English language teaching involves much more than simply using fancy tools; it requires thoughtful pedagogical design, careful implementation, ongoing evaluation, and attention to ethical issues.

One of the most compelling advantages of AI in ELT is personalisation: adapting instruction to each learner's pace, level, strengths, and weaknesses. Traditional classroom environments usually have a fixed pace or standardized lessons. AI tools also allow for immediate feedback and automated assessment. Despite the many benefits, there are important challenges and caveats.

In light of these examples, there are several strategies for effective integration of AI in ELT:

First, blend AI tools with human teaching rather than replacing the teacher. Teachers remain essential for scaffolding, motivation, clarifying cultural/pragmatic norms, giving feedback on creative or higher-order skills, and guiding learners in using AI tools critically.

Second, train teachers well—not only in technical operation of AI tools, but in pedagogical use: when to use which tool, how to interpret feedback, how to design tasks that balance AI and human input, how to help learners reflect on results and avoid over-dependence.

Third, ensure alignment with curricular goals and assessment standards. Before adopting AI tools, teachers/schools should map how these tools support or conflict with existing curricula, language levels, assessment rubrics.

Fourth, build learner awareness and digital literacy: help learners understand how AI works, what its limitations are, how to interpret feedback, how to use prompt engineering, how to avoid plagiarism or misuse. Activities like the vibe coding example show that teaching *about* AI is as important as using it.

Fifth, monitor and evaluate usage. Collect data on learner performance, engagement, motivation; compare AI-augmented instruction with more traditional methods; seek qualitative feedback from learners about user experience, anxiety, trust in AI feedback; adjust tools and approaches accordingly.

Sixth, address ethical, privacy, and equity issues: ensure that AI tools protect user data, that consent is obtained (especially for voice recordings), that tools are inclusive of different accents and varieties, that access is equitable.

In conclusion, integrating AI into English language teaching promises substantial benefits: personalization, immediate feedback, scalable practice, greater learner engagement, opportunities for speaking and pronunciation practice outside the class, and more flexible learning. The future of ELT lies in a synergy between teacher expertise, learner agency, and well-designed AI tools.

References

1. Crompton, H. AI and English language teaching: Affordances and challenges. *British Journal of Educational Technology*, 2024, 55(4), 785–802.
2. Du, J. A systematic review of AI-powered chatbots for English as a foreign language learning. *Computers & Education: Artificial Intelligence*, 2024, 5, 100-123.

COMMUNITY ANALYSIS AND MODULARITY ASSESSMENT IN COGNITIVE MODELLING OF COMPLEX SYSTEMS

Pankratova N.D., Musiienko D.I.

natalidmp@gmail.com, dmusiienko@gmail.com

*Institute for Applied System Analysis of the National Technical University of
Ukraine “Igor Sikorsky Kyiv Polytechnic Institute”*

Cognitive mapping has emerged as a powerful methodology for understanding and modeling complex systems through the representation of causal relationships between concepts. In the context of complex system analysis, cognitive maps serve as directed graphs where nodes represent key concepts or factors, and edges denote causal influences between them.

Communities, or modules, are groups of densely interconnected nodes that exhibit stronger internal connections than external ones. The identification of these communities reveals the inherent structural organization of the cognitive model, highlighting clusters of concepts that function as cohesive units within the larger system. This structural analysis provides insights into the system's decomposition, interdependencies, and potential points of intervention.

The application of community detection algorithms to cognitive maps draws from established network science methodologies. Newman [1] introduced fundamental approaches for evaluating community structure through modularity measures, which quantify the strength of division of a network into modules. The modularity Q is calculated as the fraction of edges within communities minus the expected fraction if edges were distributed randomly, providing a scalar value that characterizes the quality of community partitioning.

Modern algorithms such as the Louvain method[2] enable efficient community detection in large-scale networks through iterative optimization of modularity. This approach is particularly relevant for cognitive maps, where the number of concepts and relationships can be substantial.

Applying the Louvain method [2] to the graph of cognitive map “Tunnel 1” [3] yields six communities:

Community 0: ['V1. Condition of tunnel system (1)', 'V7. The scale of the impact of an undesirable event', 'V8. Ability to function', 'V15. Investor', 'V18. Geotechnology of construction']

Community 1: ['V2. Anthropogenic activities', 'V2.1 The Evil Mind: Fighting, Terrorism', 'V2.2. Without malice', 'V4. Natural disasters, weather catastrophes', 'V4.1 Shifts', 'V5. Protection of the object', 'V16.1. Integrity of the ruin system', 'V16. Level of damage', 'V23. Underground vibrations']

Community 2: ['V3. Technogenic events', 'V3.1. Technical', 'V3.2. Technological']

Community 3: ['V24. Atmospheric pollution', 'V10. Environmental consequences', 'V12. Consequences for life', 'V13. The number of injured']

Community 4: ['V9. Time to restore functioning', 'V11. Economic consequences', 'V14. Organizational, technical, etc. capabilities', 'V17. Material damage']

Community 5: ['V19. Capacity of land routes', 'V20. Population', 'V21. Intensity of movement', 'V22. Average speed']

The obtained modularity value of $Q = 0.58510$, indicating a strong community structure that substantially exceeds random expectation. This high modularity suggests that the tunnel system can indeed be understood as a set of distinct yet interrelated functional modules, each representing different aspects of tunnel-system management and risk.

References:

1. Newman, M. E. J. Modularity and community structure in networks. Proceedings of the National Academy of Sciences, 103(23), 2006. P.8577–8582. <https://doi.org/10.1073/pnas.0601602103>
2. Blondel, V. D., Guillaume, J.-L., Lambiotte, R., & Lefebvre, E. Fast unfolding of communities in large networks. Journal of Statistical Mechanics: Theory and Experiment, 2008(10), P10008. <https://doi.org/10.1088/1742-5468/2008/10/p10008>
3. Pankratova N.D., Musiienko D.I. Study of the underground tunnel planning. Cognitive modelling // System research and information technologies. № 1. 2023. P. 37-50. <https://doi.org/10.20535/SRIT.2308-8893.2023.1.03>

DIGITAL TWINS IN DECENTRALIZED UAV SWARM CONTROL

Pankratova N.D., Pankratov V.A., Golinko I.M.

natalidmp@gmail.com, pankratov.volodya@gmail.com, golinko.igor@lil.kpi.ua

Educational and Scientific Complex "Institute for Applied System Analysis

NTUU "Igor Sikorsky Kyiv Polytechnic Institute"

The concept of a digital twin was first introduced in the works of University of Michigan professor M. Grivz, where the key elements are the existence of a physical object, its virtual copy, and automated data exchange between them [1]. In wartime, swarm UAV systems attract particular attention, as they are capable of providing increased stability, efficiency, and speed of task execution compared to individual drones. At the same time, the integration of the CD concept into such systems faces additional challenges: the lack of stable communication with ground infrastructure, the need for cognitive coordination of the swarm, and the need to operate in an electronic warfare environment. This requires modifying the classic definition of CD, providing for the possibility of its operation in digital shadow modes or asynchronous interaction between the ground center and onboard agents.

The purpose of this thesis is to develop a strategy for using digital twins in decentralized swarm control of unmanned aerial vehicles. The digital twin of a swarm of drones acts as a kind of "central brain" that coordinates the actions of all drones. It prepares missions, distributes tasks among each drone, monitors their interaction, and adapts in real time. After completing a mission, the digital twin conducts a detailed analysis, identifies possible improvements and adjustments, and helps optimize and adjust future missions, taking into account the experience of previous operations. It is also important that, in a decentralized control environment, the swarm itself acts autonomously, and the digital twin becomes a strategic tool for analysis and further improvement of efficiency. Swarm systems often do not have a single central node, and each drone operates autonomously, making decisions on the spot. In this case, the digital twin acts

more as a "central planner" during the mission preparation and subsequent analysis stages.

During the preparation phase, the digital twin performs the following actions: prepares and forms the mission; forms an overall strategy; distributes goals and tasks among the swarm drones; does not interfere with the mission execution, as the drones operate autonomously. After the mission is complete, the digital twin performs the following actions: analyzes the results; identifies areas for improvement in the strategy, forecasts and adjusts the mission; updates behavioral strategies during non-mission time.

The architecture of a digital twin for swarm systems involves a modular approach. The main part of the digital twin is located at the ground center, where there are also other modules responsible for software development, mission preparation, output data, etc. The drones themselves operate autonomously in flight and have their own "local" model, which is integrated with the global model of the digital twin. Thus, the architecture of a digital twin for swarm systems is a combination of a central node for mission preparation and results analysis at a ground center and autonomous modules in the form of AI agents on board each UAV. Onboard AI agents provide local adaptation, diagnostics, environment exploration, and cognitive control of drone behavior. In the absence of stable communication with the ground center, the effective functioning of the drone swarm is possible thanks to the distribution of the functions of the digital twin of the ground station and the AI agents of the onboard UAV. As a result, this is a unified system where the ground center coordinates and the UAVs perform tasks independently.

Literatura

1. Grieves, M. Product lifecycle management: the new paradigm for enterprises. *Int. J. Prod.Dev.*, 2, 2005, pp. 71-84. <https://doi.org/10.1504/IJPD.2005.006669>
2. Zgurovsky M.Z., Pankratova N.D., Golinko I.M., Grishyn K.D. Digital twins in AI-controlled navigation tasks for autonomous UAV swarm. *Int. J. System research and information technologies*. №3. 2025. pp.19-32. doi:[10.20535/SRIT.2308-8893.2025.3.02](https://doi.org/10.20535/SRIT.2308-8893.2025.3.02)

METHOD OF CALCULATING PROBABILITY OF STATES OF A SINGLE MASS SERVICE SYSTEM WITH UNRELIABLE SERVICE DEVICES

Poslaiko N.I., poslaikoni@i.ua,

Oles Honchar Dnipro National University

In real queueing systems, while servicing requests, sometimes the serving devices can fail and then recover. If the probability of such situations is sufficiently small, then this factor is usually not taken into account in mathematical models of service systems. Otherwise, the analysis of queueing systems becomes much more complicated, since in addition to the usual characteristics, such as queue length and waiting time for the start of service, it is necessary to take into account the uptime and repair time of devices.

In the works of the Ukrainian mathematician Yezhov I.I. and his student Aliyev T.F., a mathematical model of a mass-service system with an unlimited queue and a queue discipline of “first come, first served” was first considered, in which the service process is represented as a two-dimensional Markov process, one component of which at each moment of time coincides with the number of requests in the system (in the queue and in service), and the other with the number of devices that have failed.

For the probabilities of the service process states, a system of Kolmogorov direct differential equations was derived, which is infinite. Using generating functions, it was reduced to a finite system of differential equations, which, in addition to the generating functions, also included expressions containing part of the sought probabilities. To solve such a system of differential equations, the apparatus of Markov processes homogeneous in the second component was used.

It was shown that the desired probabilities of the states discussed above are found from the system of Volterra integral equations of the second kind.

Based on the results of the above authors, this paper shows how, using numerical methods, it is possible to approximately calculate the probabilities of the states of the auxiliary process (homogeneous in the second component), through which the probabilities of the states of the main service process are expressed, since it is impossible to obtain explicit analytical expressions for them due to local characteristics of the service process (intensity of incoming requests, their service, failure of service devices and their restoration).

The algorithm for approximate calculation of probabilities of service process states consists of a number of stages. First, the system of differential equations for the generating functions of the probabilities of the auxiliary process states is reduced to solving a system of linear algebraic equations by the operational method (using the Laplace transform). After its solution, based on Laurent's theorem, expressions for the probabilities of the states of the auxiliary process are found in the form of curvilinear integrals over a closed loop, which can be reduced to definite integrals by known methods. The latter integrals are functions of the states of the component of the auxiliary process, which is homogeneous in the state space. When the component acquires large absolute values, rapidly oscillating functions are present under the sign of definite integrals. In this case, the use of ordinary quadrature formulas for the approximate calculation of such integrals is impractical. There are various approaches to the approximate calculation of integrals from rapidly oscillating functions. For example, specialized quadrature formulas have been developed (Filon, Levin, and others). Finally, for an approximate solution of the system of integral equations of the second kind, which was discussed above, the method of finite sums can be used, solving it on the interval $(0; t)$ with a certain step.

The state probabilities of a service system are its most informative characteristics. Using them, other important characteristics of the system's operation can be calculated.

1.EXPLAINABLE AI VIA TRANSITION-MATRIX MAPPING FROM DEEP TO INTERPRETABLE FEATURES

Radiuk P.M. ¹, radiukp@khnmu.edu.ua

Barmak O.V. ¹, barmako@khnmu.edu.ua

Krak I.V. ^{2,3}, iurii.krak@knu.ua

¹ *Khmelnyskyi National University*

² *Taras Shevchenko National University of Kyiv*

³ *V.M. Glushkov Institute of Cybernetics*

This work introduces a validated method to enhance the explainability of “black-box” deep learning (DL) models. The approach computes a linear transition matrix that creates an interpretable bridge between a model’s complex internal features and a human-understandable feature space. This target space can be defined by classical machine-learning models or, more powerfully, by domain experts using established criteria. The framework enables both the synthesis of human-inspectable artifacts from model internals and the analytical reconstruction of decision boundaries.

Problem Statement. Given a set of n data samples, let the matrix $(\mathbf{D} \in \mathbb{R}^{n \times d})$ represent the features extracted from a high-performing but uninterpretable “black-box” model, such as a deep neural network, where d is the dimensionality of the deep feature space. Let the matrix $(\mathbf{B} \in \mathbb{R}^{n \times b})$ represent corresponding features from an interpretable space, where b is the dimensionality of a “white-box” or human-understandable feature set. The objective is to find a linear transformation, or transition matrix [1], $\mathbf{T} \in \mathbb{R}^{d \times b}$, that optimally maps the deep feature representations to the interpretable ones. This can be expressed as the approximation problem:

$$\mathbf{D}, \mathbf{T} \approx \mathbf{B}. \quad (1)$$

We estimate $\mathbf{T}$ by minimizing the Frobenius norm of the residual:

$$\mathbf{T}^* = \arg \min_{\mathbf{T} \in \mathbb{R}^{d \times b}} \|\mathbf{DT} - \mathbf{B}\|_F^2. \quad (2)$$

The least-squares solution is given by:

$$\mathbf{T}^* = \mathbf{D}^\dagger \mathbf{B}. \quad (3)$$

If $\mathbf{D}$ has full column rank, this specializes to $\mathbf{T}^* = (\mathbf{D}^T \mathbf{D})^{-1} \mathbf{D}^T \mathbf{B}$.

Methodology. Our approach maps a DL model's latent representations (*the formal model*, matrix $\mathbf{D}$) to a semantically meaningful feature space (*the mental model*, matrix $\mathbf{B}$). The novelty lies in the flexible construction of $\mathbf{B}$, which can be derived from an interpretable ML model or, for greater impact, defined directly by domain experts using established heuristics and clinical guidelines [2].

The computed transition matrix $\mathbf{T}$ enables two key functions:

1) **Synthesis.** For any new deep feature vector $\mathbf{d} \in \mathbb{R}^d$, its interpretable counterpart $\mathbf{b} \in \mathbb{R}^b$ is synthesized using the forward mapping:

$$\mathbf{b}^T = \mathbf{d}^T \mathbf{T}. \quad (4)$$

This allows users to visualize the model's internal representations in understandable terms. In matrix form for a batch $\mathbf{D}_{\text{new}}$, this reads $\mathbf{B} = \mathbf{D}_{\text{new}} \mathbf{T}$.

2) **Analysis.** The matrix $\mathbf{T}$ enables the reconstruction of a decision boundary from the deep space in the interpretable space, supporting inspection of decision geometry and assessment of classification confidence [3]. For linear decision rules $f(\mathbf{d}) = \text{sign}(\mathbf{w}^T \mathbf{d} + b)$, the corresponding interpretable coefficients can be analyzed via the mapping induced by $\mathbf{T}$.

Performance and Validation. The framework's effectiveness was rigorously validated across complementary domains:

– Image synthesis (MNIST) [1]. The method preserved key visual information in reconstructed images, achieving mean SSIM of 0.697 and PSNR of 17.94 dB.

– Text-based stance detection (FNC-1) [1]. Using an ML model to define interpretable space, the approach achieved competitive performance with weighted accuracy of 89.38%.

– Healthcare (ECG, MRI) [2]. For arrhythmia detection from ECG signals and heart disease classification from MRI scans, outputs were mapped to features defined by cardiologists. Results showed strong agreement with expert assessments, achieving Cohen’s Kappa coefficients of approximately 0.89 (ECG) and 0.80 (MRI), indicating reliable alignment with clinical reasoning.

Conclusion. The transition-matrix framework provides a computationally efficient and mathematically principled method for making complex AI models more comprehensible. By creating a linear bridge to an expert-defined or model-defined interpretable space, it renders opaque model internals transparent and auditable. Its strong performance across image, text, and medical signal-processing tasks highlights both versatility and utility.

References

1. Explainable deep learning: a visual analytics approach with transition matrices / P. Radiuk et al. *Mathematics*. 2024. Vol. 12, no. 7. P. 1024. URL: <https://doi.org/10.3390/math12071024> (date of access: 28.08.2025).

2. Toward explainable deep learning in healthcare through transition matrix and user-friendly features / O. Barmak et al. *Frontiers in artificial intelligence*. 2024. Vol. 7. P. 1482141. URL: <https://doi.org/10.3389/frai.2024.1482141> (date of access: 28.08.2025).

3. A comprehensive perspective on explainable AI across the machine learning workflow / G. Paterakis et al. Ithaca, NY, USA : Cornell University, 2025. 97 p. (Preprint. Cornell University ; 2508.11529). URL: <https://doi.org/10.48550/arXiv.2508.11529> (date of access: 23.10.2025).

APPLICATION OF BAYESIAN NETWORKS IN RECOMMENDER SYSTEMS

Rakova K.O., rakova_k@365.dnu.edu.ua,

Sheveleva A.E., shevelevaee@dnu.dp.ua

Oles Honchar Dnipro National University (www.dnu.dp.ua)

Introduction. Recommender systems have become an integral part of intelligent information technologies that personalize user experience across digital platforms. They help analyze large datasets and predict user interests, which significantly improves interaction quality. However, many classical algorithms — such as collaborative or content-based filtering — struggle with the lack of sufficient data, the cold-start issue, and the absence of clear interpretability. To overcome these challenges, Bayesian networks are used as a probabilistic approach that models uncertain relationships between user features and item attributes.

Statement of the problem. This study investigates how Bayesian networks can enhance the accuracy and interpretability of recommender systems. It involves examining Bayesian network principles, developing a prototype to predict user preferences using conditional probabilities, and comparing results with traditional methods. The goal is to show that probabilistic reasoning improves the adaptability and reliability of recommendation engines.

Solution of the problem. The proposed model was implemented in Python, employing the pgmpy library for network creation and pandas for data processing. The experiment utilized a publicly available dataset of user–movie ratings. The designed Bayesian network contained five interconnected variables: UserAge, Genre, Movie, Rating, and Liked. These nodes represent user demographics, content type, and feedback, which together determine the likelihood that a person will enjoy a specific movie. The model calculates

conditional probabilities based on previous ratings and uses them to generate a ranked list of suggested items. For instance, if a user within a certain age group tends to rate comedies highly, the network increases the probability of recommending other comedy films.

Results. After training and evaluation, the developed model achieved a prediction accuracy of about 0.82 while maintaining a response time under five seconds per query. Tests on limited and sparse datasets confirmed the model's stability and robustness compared to conventional collaborative filtering. An important advantage of the Bayesian approach is its explainability — each recommendation can be traced to the probability dependencies within the network, clarifying why certain items appear in the suggested list. The system can therefore combine efficiency with interpretability, which is essential for real-world recommendation services.

Conclusions. The conducted research confirms that Bayesian networks can serve as a strong foundation for building personalized recommendation systems that balance accuracy and transparency. Their probabilistic structure allows modeling complex interactions between user preferences and item features while avoiding the “black-box” nature of deep learning.

Literature

1. Koller D., Friedman N. Probabilistic Graphical Models: Principles and Techniques. MIT Press, 2009.
[https://wcl.cs.rpi.edu/pilots/library/papers/TAGGED/4347-Koller_Friedman%20\(2009\)%20-%20Probabilistic%20graphical%20models.pdf](https://wcl.cs.rpi.edu/pilots/library/papers/TAGGED/4347-Koller_Friedman%20(2009)%20-%20Probabilistic%20graphical%20models.pdf)
2. Ricci F., Rokach L., Shapira B. Recommender Systems Handbook. Springer, 2022. https://doi.org/10.1007/978-0-387-85820-3_1

MODELING DISCRETE PRE-FRACTURE GAPS AND ELECTROMECHANICAL RESPONSE IN DISSIMILAR PIEZOELECTRIC BIMATERIAL SYSTEMS

Shcherbak R., roma.shcherbak1205@gmail.com

Sheveleva A., shevelevaee@dnu.dp.ua,

Oles Honchar Dnipro National University (www.dnu.dp.ua)

Piezoelectric materials exhibit a strong interaction between electrical and mechanical fields, a property that underpins their wide use in devices such as sensors, actuators, and precision control systems. Despite these advantages, piezoelectric ceramics, which are polycrystalline engineered materials, remain mechanically fragile. Their brittleness often leads to the formation of cracks, especially along interfaces.

To study such effects, consider a bimaterial configuration composed of two semi-infinite piezoceramic regions occupying the domains $x_3 > h/2$ and $x_3 < -h/2$. Each region belongs to the symmetry class 6mm, and their polarization axes are both oriented along the x_3 -direction. These two domains are joined by a thin isotropic interlayer of thickness h , characterized by shear modulus μ , Poisson's ratio ν , and yield limit σ_Y . The composite system is exposed, at a far-field distance, to uniform mixed-mode loading conditions

Because the external loading remains constant along the x_2 -axis, the governing boundary-value problem can be simplified to a two-dimensional formulation within the (x_1, x_3) -plane.

It is considered that the interface layer has lower stiffness than the adjacent matrix regions. Consequently, fracture initiation begins through the development of localized pre-fracture areas, denoted by the intervals $[a_1, a]$ and $[b, b_1]$. Because the interface layer is extremely thin and each pre-fracture zone is much shorter than the main crack, the stress components inside

every such region may be regarded as spatially uniform. For the left zone $[a_1, a]$, the stress state is described by the following relations $\sigma_{zz} = \sigma', \tau_{xz} = \tau', \sigma_{xx} = 0$. For the right zone $[b, b_1]$, the corresponding stresses are $\sigma_{zz} = \sigma, \tau_{xz} = \tau, \sigma_{xx} = 0$. The regions $[a_1, a]$ and $[b, b_1]$ are not continuous, they are separated by narrow gaps commonly referred to as crazes.

Assuming that for the crazes the relation between stresses is defined by the equations $f(\sigma, \tau, \sigma_1) = 0$ (right zone), $f(\sigma', \tau', \sigma'_1) = 0$ (left zone) approach [1] was used. The transcendental equations for the determination of the pre-fracture zones' lengths are found and solved numerically for some specific materials and loading. The comparison of the obtained results for number of gaps tending to zero shows good agreement with the case of the continuous pre-fracture zones considered in [1].

[1] Шевельова А.Є., Щербак Р.О. Про зони передруйнування для міжфазної тріщини в п'єзоелектричному матеріалі в полі розтягувальних та зсувних напружень // Проблеми обчислювальної механіки та міцності конструкцій. – 2025. – Вип. 39. – С. 88-100.

OPTIMIZED PACKING OF IRREGULAR OBJECTS

Shekhovtsov S.¹, serhii.shekhovtsov@nure.uaRomanova T.^{1,2,3}, t.romanova@leeds.ac.uk¹ Kharkiv National University of Radio Electronics² *Anatolii Pidhornyi Institute of Power Machines and Systems of the National Academy of Sciences of Ukraine*³ *University of Leeds, UK*

This research addresses the problem of packing two-dimensional irregular objects, defined in a non-Euclidean metric space, into a minimum-area container while satisfying specified placement constraints. Analytical description of non-overlap, containment, and allowable inter-object distances conditions are stated using **Lp norms** for $p \geq 1$ [1,2]. The packing problem is formulated as a **nonlinear programming model**. A corresponding model-based **multistart solution strategy** is proposed. The method consists of three core stages: generation of feasible initial configurations using a fast heuristic algorithm; local enhancement of each configuration, and selection of the best solution among the improved ones. Computational experiments for different values of parameter p are provided, employing the **IPOPT** solver [3] for local optimization and the solver **BARON** [4] to search for global solutions for small-sizes instances. The obtained results demonstrate the efficiency and robustness of the proposed approach.

1. Narici, L., Beckenstein, E. Topological Vector Spaces. Pure and applied mathematics (Second ed.). Boca Raton, FL: CRC Press, 2011.
2. Litvinchev, I., Fischer, A., Romanova, T., Stetsyuk, P. A new class of irregular packing problems reducible to sphere packing in arbitrary norms. Mathematics 12(7), 935, 2024.
3. Wachter, A.; Biegler, L.T. On the implementation of an interior-point filter line-search algorithm for large-scale nonlinear programming. Mathematical Programming. 106, 1, 25-57, 2006.
4. Khajavirad, A. and N. V. Sahinidis, A hybrid LP/NLP paradigm for global optimization relaxations, Mathematical Programming Computation, 10, 383-421, 2018.

The work is supported by the British Academy (Grant #100072)

II. MATHEMATICAL MODELING OF ADAPTIVE LEARNING SYSTEMS

Stroieva V.O., vikaastroeva@ukr.net, **Zhyhaleva S.P.**, sofiazigaleva@ukr.net
Dniprovsky State Technical University

The authors were tasked with developing a modern adaptive system that would be a software product capable of ensuring maximum adaptation of educational content in accordance with the level of knowledge of the person being taught, his individual characteristics and taking into account the current successes that he showed during the educational process. Based on the technology of automated learning, a structure was built in the work, which is the basis for developing a general model of an adaptive distance education system. The advantage of such a model is the use of current results from the student module, based on the analysis of which the further educational trajectory of the student's movement is built. The use of fuzzy logic methods and neural models of presenting educational material allows proportionally dividing educational content and organizing its structured presentation.

A method for assessing the results of test control of the level of learning by the set of knowledge has been developed, which makes it possible to eliminate gaps in the student's knowledge during the study of a new topic by determining the depth of unlearned information by forming blocks of educational material for repeated or in-depth study. The use of the proposed information technology in adaptive distance education systems allows for an individual approach to the organization of programmed learning in accordance with the needs of students, the characteristics of the academic discipline, and learning conditions, and to raise the organization of the educational process within the framework of distance education to a qualitatively higher level.

AI AS A TEACHING ASSISTANT: REDEFINING THE ROLE OF THE ENGLISH LANGUAGE TEACHER

Suima I.P., irinasuima2017@gmail.com

Novikova O.V., novikovaolga19911991@gmail.com

Oles Honchar Dnipro national university

Artificial Intelligence (AI) has rapidly emerged as a transformative force in the field of education, and its impact on English Language Teaching (ELT) is particularly profound. Traditionally, the English language teacher's role has been multifaceted, encompassing functions such as facilitator, mentor, assessor, content creator, and cultural guide. Teachers have historically been responsible not only for delivering linguistic knowledge but also for fostering communicative competence, critical thinking, and intercultural awareness. However, the integration of AI into the educational process is reshaping these responsibilities, introducing new possibilities while also presenting challenges that require careful consideration. One of the most significant contributions of AI in ELT is the automation of administrative and repetitive tasks, which allows teachers to focus on pedagogical and interactive aspects of their work. AI-powered platforms such as Khanmigo, developed by Khan Academy, assist educators in generating lesson plans, exercises, quizzes, and rubrics. These platforms use sophisticated algorithms to align teaching materials with curriculum standards, monitor students' progress, and adjust tasks according to learners' levels. AI also offers innovative opportunities for speaking and listening practice, which are often constrained in large classes due to time limitations. Conversational AI agents and chatbots can simulate real-life dialogues, allowing students to practise interactions such as ordering at a restaurant, conducting interviews, or negotiating in professional contexts. These tools use natural language processing (NLP) to recognize errors, suggest corrections, and provide immediate feedback on pronunciation, fluency, and

vocabulary. A practical example is the use of AI-driven speaking tutors in virtual classrooms, where students engage in role-play scenarios and receive automated feedback on intonation, stress, and rhythm. In practice, successful integration of AI as a teaching assistant in ELT involves blended approaches. For example, pre-class activities may be managed by AI—such as vocabulary drills, grammar exercises, or reading comprehension tasks—while in-class sessions focus on discussions, collaborative projects, and speaking practice facilitated by the teacher. AI can also support ongoing assessment, tracking progress over time and helping educators make data-informed decisions about lesson planning and differentiation.

In conclusion, AI is redefining the role of the English language teacher by acting as an intelligent assistant that automates routine tasks, personalizes learning, facilitates innovative teaching methods, and enhances formative assessment. Rather than replacing teachers, AI provides tools that allow educators to focus on higher-order teaching functions, such as mentoring, cultural guidance, critical thinking facilitation, and emotional support. The role of the teacher evolves from information transmitter to orchestrator of learning experiences, leveraging AI to empower both themselves and their students to achieve optimal linguistic, cognitive, and socio-emotional outcomes.

References

1. Daud, A. Integrating Artificial Intelligence Into English Language Teaching: A Systematic Review. *European Journal of Educational Research*, 2025, 14(3), 405–427.
2. Koç, F.Ş. The use of artificially intelligent chatbots in English language learning: a systematic meta-synthesis. *ReCALL / Cambridge University Press*, 2025, 37(1), 1–24.

DECISION SUPPORT INFORMATION TECHNOLOGY IN VIDEO SURVEILLANCE AND MONITORING TASKS

Sukovenko K., sukovenko@ftf.dnu.edu.ua

Dnipro National University Oles Honchar

Modern video surveillance systems generate vast amounts of data that require intelligent processing to become useful for decision-making. Decision support information technology in this field aims to transform continuous and archived video streams into knowledge that helps operators or automated systems identify threats, react promptly, and improve efficiency. The technology relies on video analytics modules capable of detecting objects, recognizing behavior, and identifying anomalies in real time, such as unattended baggage or movement in restricted areas [1]. The resulting events are structured and prioritized, ensuring that only relevant and high-risk alerts reach human operators.

To enhance performance, these systems often implement multi-level architectures combining edge and cloud computing, which reduce latency and optimize bandwidth use [2]. In addition, cryptographic and blockchain-based approaches are applied to preserve the authenticity of video data, ensuring that recordings remain tamper-proof and reliable as potential evidence [3]. This combination of distributed processing and data integrity mechanisms allows surveillance systems to operate efficiently even under heavy loads, maintaining trust and continuity of monitoring operations.

Besides technical efficiency, decision support technologies must comply with ethical and legal standards. AI-powered surveillance raises concerns about privacy and potential misuse, emphasizing the need for transparent decision-making, access control, and anonymization tools [4]. Therefore, modern intelligent monitoring systems integrate analytics, cybersecurity, and governance mechanisms into a single framework that not only records video but

also converts it into real-time, evidence-based decisions that support safety while respecting human rights.

Integration of these technologies with the Internet of Things and machine learning allows cameras and sensors to operate as intelligent agents that exchange data and adapt to their environment. This enables predictive analysis and automated responses, such as activating alarms or notifying authorities in real time [2], which greatly improves the accuracy and efficiency of surveillance compared to traditional monitoring.

Yet, implementing decision support systems faces challenges. High costs, data protection requirements, and cybersecurity risks demand regular updates and strict control [4]. A balanced approach combining technology, regulation, and human oversight is essential to maintain reliability, safety, and ethical responsibility in modern surveillance systems.

References:

1. Nikouei S. Y., Chen Y., Song S., Xu R., Choi B. Y., Faughnan T. I-SAFE: Instant Suspicious Activity identiFication at the Edge using Fuzzy Decision Making [Electronic resource]. – 2019. – Available at: <https://arxiv.org/abs/1909.05776>
2. Michelin R. A., Maia G., Ramos G. O., Loureiro A. A. F., Villas L. A. Leveraging lightweight blockchain to establish data integrity for surveillance cameras [Electronic resource]. – 2019. – Available at: <https://arxiv.org/abs/1912.11044>
3. Fontes C., Lopes N., Almeida A. I. AI-powered public surveillance systems: why we (might) need to regulate them [Electronic resource] // Technology in Society. – 2023. – Available at: <https://www.sciencedirect.com/science/article/pii/S0160791X22002780>
4. Moepi G. L., Olugbara O. O., Odun-Ayo I. Smart Surveillance Systems: Trends, Challenges and Future Directions [Electronic resource]. – 2024. – Available at: https://www.researchgate.net/publication/391576838_Smart_Surveillance_Systems_Trends_Challenges_and_Future_Directions

AN EXPLAINABLE ARTIFICIAL INTELLIGENCE FRAMEWORK FOR PREDICTING PROSTATE CANCER PROGRESSION IN MEN UNDER ACTIVE SURVEILLANCE

Sushentsev¹ N., Zaikin^{2,3} A., Barrett¹ T., Blyuss⁴ O.

o.blyuss@qmul.ac.uk

*1. Department of Radiology, Addenbrooke's Hospital and University of
Cambridge, Cambridge, UK*

*2. Department of Women's Cancer, Institute for Women's Health, University
College London, London, UK*

3. Department of Mathematics, University College London, London, UK

*4. Wolfson Institute of Population Health, Queen Mary University of London,
London, UK*

Background

Around half of all men diagnosed with prostate cancer (PCa) present with low- or intermediate-risk disease that qualifies them for active surveillance (AS). Despite this conservative approach, many patients eventually show pathological progression over time. This work aimed to construct predictive models capable of identifying histopathological progression in individuals monitored under AS.

Methods

The study included patients with biopsy-verified PCa and at least two years of longitudinal follow-up. Participants were routinely monitored through prostate-specific antigen (PSA) assessments and magnetic resonance imaging (MRI). Between three and six serial MRI scans were available per patient. Progression was defined as an increase in grade group on subsequent targeted biopsy. To forecast progression, we developed machine learning (ML) models that combined radiomic parameters extracted from MRI with clinical indicators. Model interpretability was assessed using SHapley Additive exPlanations (SHAP).

Results

Three predictive frameworks were developed.

1. The **baseline model**, which used radiomic features from the first MRI along with baseline PSA density (PSAd), achieved an AUC of 0.793 (sensitivity = 0.690; specificity = 0.830).
2. The **delta model**, based on temporal changes in radiomic characteristics between the baseline and most recent MRI together with final PSAd, yielded an AUC of 0.913 (sensitivity = 0.793; specificity = 0.936).
3. The **longitudinal model**, incorporating all available MRI-derived radiomic data and PSAd measurements over time, reached an AUC of 0.917 (sensitivity = 0.828; specificity = 0.894).

Conclusion

The models showed strong ability to differentiate patients with and without pathological progression. These findings highlight the promise of ML-driven radiomic analysis for refining risk stratification in PCa and supporting more individualized management strategies for men on active surveillance.

INTELLIGENT MULTI-AGENT SYSTEMS FOR ADAPTIVE MANAGEMENT UNDER UNCERTAINTY

Symonov D., denys.symonov@gmail.com,

Symonov Y., e.symonov@gmail.com, **Zaika B.**, zaikabohdan5@gmail.com,

Kutova I., iippk2014@ukr.net, **Fesenko M.**, clnt_49@ukr.net

*V.M. Glushkov Institute of Cybernetics of the National Academy of Sciences
(NAS) of Ukraine, Kyiv*

Modern complex systems (cyber-physical, socio-technical, economic) operate under dynamic uncertainty, which limits the effectiveness of traditional centralized control due to poor scalability and vulnerability to failures. A promising alternative is multi-agent systems (MAS), where each agent functions as an AI model capable of autonomous decision-making, learning, and interaction. Their decentralized structure removes the single point of failure: even if one agent fails, others sustain collective dynamics. Agents also engage in self-organization, forming coordinated global strategies without centralized supervision. MAS effectively combine deep reinforcement learning, cognitive models, and Bayesian inference, enabling the generation and comparison of alternative scenarios and the real-time adaptation of control. These properties make MAS a universal tool for modeling, forecasting, and managing the dynamics of complex systems under uncertainty.

Adaptive control in multi-agent systems is viewed as the interaction of a set of agents $A = \{a_1, \dots, a_m\}$ with an environment described by the state space $S = \{s_1, \dots, s_n\}$. Each agent follows a policy $\pi_i: S \rightarrow U_i$, which defines its action u_i . The system dynamics are modeled as a stochastic transition $P(s_{t+1}|s_t, u_1, \dots, u_m)$. Since the environment is partially observable, the problem is represented as a POMDP $\langle S, U, T, R, \Omega, O \rangle$, where T is the transition function, R the reward, and Ω the observation space. The objective is to

construct policies $\{\pi_1, \dots, \pi_m\}$ that jointly maximize the system's collective utility under uncertainty.

Modern intelligent agents combine multi-agent reinforcement learning, where the policy π_i is updated according to

$$Q_i(s, u_i) \leftarrow Q_i(s, u_i) + \alpha \left[r + \gamma \max_{u'_i} Q_i(s', u'_i) - Q_i(s, u_i) \right], \quad (1)$$

with deep neural networks used as function approximators, as well as evolutionary algorithms and cognitive models that integrate logical reasoning with Bayesian inference. A system is considered adaptive if agents adjust their policies in response to environmental changes (E_t) and interactions with other agents (I_t):

$$\pi_i^{t+1} = f(\pi_i^t, E_t, I_t). \quad (2)$$

This enables real-time collective control, for example in transportation or energy systems, where robust and coordinated behavior is essential.

Multi-agent intelligent systems are an effective tool for adaptive control of complex systems under uncertainty. They enable the integration of machine learning methods, scenario analysis, and system dynamics. Further progress is associated with the integration of MARL, explainable AI, and cognitive architectures, which will ensure reliable and interpretable agent behavior in critical applications.

References

1. Jia, L., & Pei, Y. (2025). Recent Advances in Multi-Agent Reinforcement Learning for Intelligent Automation and Control of Water Environment Systems. *Machines*, 13(6), 503. <https://doi.org/10.3390/machines13060503>
2. Hong, X., Ayadi, W., Alattas, K.A. et al. Adaptive average arterial pressure control by multi-agent on-policy reinforcement learning. *Sci Rep* 15, 679 (2025). <https://doi.org/10.1038/s41598-024-84791-5>

MATHEMATICAL MODELLING OF COMPLEX NATURAL SYSTEMS: A PROBABILISTIC FIRE SPREAD MODEL WITH ENVIRONMENTAL FACTORS

Trotsenko A.G., trotsenko1998@gmail.com

Kuzenkov O.O., kuzenkov1986@gmail.com

Oles Honchar Dnipro National University (www.dnu.dp.ua)

Complex natural systems—such as ecosystems, climate systems, and biological or socio-economic processes—are characterized by nonlinear dynamics, hierarchical structure, and stochastic interactions between numerous components. Understanding and predicting their evolution is a key challenge in modern applied mathematics, as these systems exhibit phenomena such as self-organization, stability transitions, and chaotic behavior.

This work focuses on the construction and analysis of **nonlinear mathematical models** for complex natural systems, illustrated by the example of a **fire spread process** in a two-dimensional environment. The research presents two models: a **simple deterministic cellular model** and an **improved probabilistic model** that incorporates external environmental factors.

The **deterministic model** assumes that fire spreads automatically to neighboring cells in each time step, forming a symmetric and predictable front. It provides a basic representation of the system's dynamics, serving as a theoretical benchmark for further extensions.

The **probabilistic model**, on the other hand, introduces a stochastic mechanism of ignition based on local environmental conditions. The probability of ignition $p_{i,j}$ for a cell (i, j) depends on parameters such as: wind direction and intensity, humidity, terrain elevation, ambient temperature, and fuel density.

This relationship is expressed as:

$$p_{i,j} = p_0(1 - m)(1 + a_w \max(0, \cos\theta_{ij}))(1 + \alpha_h \Delta h_{ij})(1 + \alpha_T(T_{i,j} - T_0))F_{i,j},$$

where each factor modifies the base probability p_0 according to physical conditions. Ignition occurs if a random number $r \in [0,1]$ satisfies $r < p_{i,j}$, introducing a stochastic component that better represents real-world uncertainty.

The model was **implemented in Python** using the *NumPy* and *Matplotlib* libraries. Numerical simulations were carried out on a 100×100 grid with variable environmental parameters. Visualization results demonstrate that the fire front becomes **anisotropic** under the influence of wind and terrain, while increased humidity significantly reduces the fire propagation rate. Temperature and fuel distribution also strongly affect the ignition pattern.

Simulation results show clear differences between the deterministic and probabilistic models:

- the deterministic model produces a perfectly symmetric, circular front;
- the probabilistic model generates irregular, fragmented fronts and localized fire clusters;
- environmental factors introduce nonlinear and threshold effects that can lead to abrupt transitions between stable and unstable spreading regimes.

The developed model provides a **unified computational framework** for studying nonlinear and stochastic processes in spatially extended natural systems.

Bibliographic references

1. Karafyllidis, I., & Thanailakis, A. (1997). A model for predicting forest fire spreading using cellular automata. *Ecological Modelling*, 99(1), 87–97. [https://doi.org/10.1016/S0304-3800\(96\)01942-4](https://doi.org/10.1016/S0304-3800(96)01942-4)
2. Boychuk, D., Braun, W. J., & Krougly, Z. L. (2009). A stochastic forest fire growth model. *Environmental Modelling & Assessment*, 14(4), 491–502. <https://doi.org/10.1007/s10666-008-9136-6>
3. Masoudian, S., Hosseinpour, R., Rahimi, A., & Khojati, G. (2025). Modelling wildfire propagation using the stochastic level-surface method. *Computers & Chemical Engineering*, 189, 108013. <https://doi.org/10.1016/j.compchemeng.2025.108013>

LINEAR ALGEBRA, COMPUTER METHODS, AND OPTIMIZATION APPARATUS

Tsukanova A. O., ponochka2511@ukr.net,

National Technical University of Ukraine

«Igor Sikorsky Kiev Polytechnic Institute»

Today, with active use of computers in different areas of our life, we have to admit that computer methods give us innovative, truly non-standard, way to look at different problems of classical mathematics. Even the simplest problems of linear algebra are not just routine tasks for university students and can be considered from a different angle via modern methods. Linear algebra is, probably, the most fundamental tool for machine learning, providing indeed powerful framework for representing, analyzing, and manipulating data.

This paper shows some results of comparison between classical method and optimization gradient methods, as well as their variations, for solving arbitrary (not necessarily quadratic) systems of linear algebraic equations. These results have been obtained after testing our own program, written in «Visual Basic for Applications». Namely, we have combined well-known methods of classical linear algebra and the next six optimization methods: gradient descent method with adapted step selection (1), gradient descent method with adapted step correction (2), modified gradient descent method (3), gradient method of the steepest descent (4), and two gradient methods of conjugate gradients: using formulas of Fletcher-Reeves (5) and Polak-Reiber (6).

An essence of these suggested optimization gradient methods is that solving an arbitrary linear system in the next matrix form

$$A\vec{x} = \vec{b}, A = A_{m \times n}, \vec{b} = B_{m \times 1}, \vec{x} = X_{n \times 1},$$

is equivalent to finding minimizer, i.e. such vector $\vec{x}^* \in \mathbf{R}^n$, for which

$$\min_{\vec{x} \in \mathbf{R}^n} f(\vec{x}) = f(\vec{x}^*),$$

for the next quadratic functional, called as residual function,

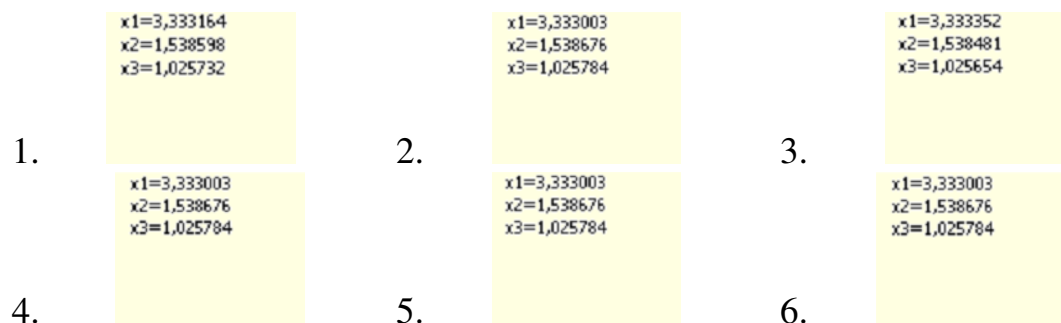
$$f(\vec{x}) = (A\vec{x} - \vec{b}, A\vec{x} - \vec{b}), A \in \mathbf{R}^{m \times n}, \vec{b} \in \mathbf{R}^m, m \leq n.$$

The method is iterative: starting from an arbitrary point $\vec{x}_0$ (called as initial approximation) in the corresponding Euclidean space, we will subsequently visit (by repeating always the same computation) points $\vec{x}_1, \vec{x}_2, \dots$, until we eventually reach a point that is the solution of the system under consideration, or it is so close to the solution that its value will be sufficiently good approximate solution.

The program consists of two parts: the first is based on the Gaussian elimination method and the second is built on using various gradient methods. Screenshots of the program's results, specifically, visual demonstrations of implementation of methods 1 – 6, for the next system of algebraic equations

$$\begin{cases} x_1 + 3x_2 + 2x_3 = 10, \\ 4x_1 + 3x_2 + 2x_3 = 20, \end{cases}$$

are shown below.



We believe our software product can serve as an effective intellectual system for studying the course of linear algebra and related mathematical subjects.

References

1. Axler S. Linear Algebra Done Right (Undergraduate Texts in Mathematics). – Springer, 2015. – 357 p.
2. Hestenes M. R. Conjugate Directions Methods in Optimization. Springer, 1980. 325 p.
3. Shores T. S. Applied Linear Algebra and Matrix Analysis (Undergraduate Texts in Mathematics). – Springer, 2018. – 491 p.

MATHEMATICAL MODELS OF DATA PROCESSING BASED ON THE FUNCTIONAL OF QUASI-EXTENT

Vovk S.M., vovk_s_m@ukr.net

Oles Honchar Dnipro National University (www.dnu.dp.ua)

Over 90 years ago, Edward Charles Titchmarsh, in the chapter X "The theory of measure and the Lebesgue integral" of his book [1, p. 319], gave the following definition of the "extent of a set": "The Riemann integral of such a function ($f(x)=1$ in E , 0 elsewhere) may be called the extent of the set E . *Extent* is thus a generalization of the *length* of an interval. The extent of E , if it exists, is written $e(E)$, so that $e(E) = \int_a^b f(x)dx$ Lebesgue's generalization is in the first place a generalization of extent". He pointed out that some sets have no extent; e.g., a set of all rational values of x in (a, b) has no extent [1].

Using Titchmarsh's definition, one can obtain the definition of the extent of a function as a Riemann integral: $E[f(x)] = \int_a^b \varphi[f(x)]dx$, where $\varphi(x) = \text{sgn}^2(x)$. However, similar to the Titchmarsh definition, the function extent may not exist (e.g., the Dirichlet function has no extent). But, excluding from consideration the cases of such abstract functions (such as Dirichlet functions), for practical issues, one can obtain a strong approach to solving data processing problems, which is based on minimizing the extent of the solution deviation, as well as the extent of the solution itself. Moreover, using the approximation of $\varphi(x)$, one can define the concept of the quasi-extent of the function as a functional [2]: $E^+[f(x)] = \int_a^b \rho[f(x); \alpha, \dots]dx$, where $\rho[f(x); \alpha, \dots]$ is a cost function, which is the Riemann integrable function on (a, b) and which has the property: $\lim_{\alpha \rightarrow \hat{\alpha}, \dots} \rho(x; \alpha, \dots) = \varphi(x)$, where $\hat{\alpha}, \dots$ are some boundary values. The use of the cost function: $\rho_s(x) = k_s \cdot [(1 + |x|^q / \alpha^q)^{\beta/q} - 1]$, where $|x| < \infty$; $k_s = 1 / [(1 + |x_H|^q / \alpha^q)^{\beta/q} - 1]$; $\alpha > 0$; $-\infty < \beta \leq 1$; $0 < q < \infty$;

$\beta < q$; $\rho_s(x_H) = 1$, makes the quasi-extent functional a "universal" data processing tool that allows tuning both to the noise environment and properties of the solution. The condition for its effective use is proper tuning of free parameters α , β , and q .

Data processing mathematical models based on the functional of quasi-extent can be generalized and written as the following minimization problem:

$$\min_{\mathbf{u}} \{ \gamma_1 E_1^+[\mathbf{b} - \mathbf{A}\mathbf{u}] + \gamma_2 E_2^+[\mathbf{T}(\mathbf{u})] \} \quad (1)$$

where $\mathbf{b}$ designates data, $\mathbf{u}$ is the problem solution, $\mathbf{A}$ describes the data formation, E_1^+ and E_2^+ are the quasi-extent functionals of residual and solution, respectively, $\mathbf{T}$ describes the solution property, γ_1 and γ_2 are the trade-off parameters. From (1), one can derive various models for data approximation [2], solving linear inverse problems [3], and others. This report discusses these issues and provides an example of deriving a mathematical model for removing linear structures and noise from photoluminescence spectra, along with the results of its application.

The proposed mathematical models enable the design of measurement information systems that adapt to noise and data properties. Their effective implementation primarily depends on tuning the free and trade-off parameters of the minimization problem; artificial neural networks can facilitate this task.

References

1. Titchmarsh E. C. The theory of functions. New York: Oxford University Press, 1932. 454 p.
2. Vovk S. M. Functional of quasi-extent and its application. *Problems of applied mathematics and mathematical modeling*. 2023. Vol. 23. P. 18–30.
3. Vovk S. M. Using the quasi-extent functional to solve linear inverse problems. *Problems of applied mathematics and mathematical modeling*. 2024. Vol. 24. P. 251–262.

IN RESEARCH OF SIMPLEST CONDITIONS FOR CLUSTERING OF A SET OF AGENTS WITH MEMORY

Yakobson O.M. alexey.yakobson@gmail.com

Nakonechna T.V., naktanya1961@gmail.com

Oles Honchar Dnipro National University

Introduction

For decades, the Prisoner's Dilemma Game has been fundamental in modelling the complexity of cooperative behaviour. Each of two players can cooperate or defect. Payoffs follow $T > R > P > S$, making defection rational and cooperation risky but socially optimal. Modern studies typically focus on memory and spatial structure. For instance, [1] models agents interacting on a network, each optimizing strategy using complete information about neighbours' past actions. In [2], agents on a 2D grid interact with neighbours and remember limited past outcomes, producing spatial clusters of cooperators separated by defectors. This raises the question: under what minimal conditions can such clustering emerge, possibly without predefined spatial structure?

Model

Two models were tested to determine whether clustering can occur. In the first, non-spatial model, N agents were randomly paired each round to play PDG. Each agent had a memory of size M , storing identifiers of previous cooperative opponents. If the opponent was remembered, the agent cooperated with probability $p \approx 1$; otherwise, with probability 0.5. After each round, memory was updated: new cooperators were added if space remained, while defectors were removed. The second model introduced localization: agents occupied positions in $[0,1]$, and interaction probability followed a normal distribution, favouring closer pairs. The goal was to observe whether groups of size roughly $M + 1$ could form and sustain internal cooperation.

Results

Simulations showed distinct outcomes: when $M = 1$, agents formed stable cooperative pairs, representing minimal clustering. However, for $M > 1$, no persistent clusters emerged. With $p = 1$, the system stabilized in fixed states without structure; with $p = 0.9$, it remained dynamic but unorganized. Introducing spatial dependence produced similar results – no clustering occurred for larger memory sizes.

Conclusion

The experiments demonstrate that neither direct reciprocity nor weak spatial localization is sufficient to generate cooperative clustering when agents possess memory beyond a single partner. Further mechanisms – such as strategy evolution, payoff modulation, or complex network effects – may be required to explain the emergence of cooperative communities in PDG populations.

References

1. Della Rossa F., Dercole F., Di Meglio A. Direct reciprocity and model-predictive strategy update. *Games* 2020;11(1):16.
2. Zhixiong Xu, Zhehang Xu, Wei Zhang, Xiao-Pu Han, Fanyuan Meng. Memory-based spatial evolutionary prisoner's dilemma.

PERFORMANCE EVALUATION OF ALGORITHMS FOR FITTING SEGMENTED LINEAR REGRESSION WITH AUTOMATIC BREAKPOINT DETERMINATION

Yuriev M.V., yuriev.m20@cf.dnu.dp.ua

Matsuga O.M., olga.matsuga@gmail.com

Oles Honchar Dnipro National University

Segmented linear regression (SLR) is a model in which the relationship between variables is represented by a set of linear segments connected at specific points called breakpoints. It captures the underlying structure of the data without resorting to complex nonlinear functions.

One of the main challenges in fitting a segmented linear model is the automatic determination of the number and location of breakpoints. In this work, a comparative analysis of four algorithms designed to perform this task automatically was conducted. The comparative analysis considered the following algorithms:

- Adaptive Greedy Segmented Linear Regression (SLRAuto) – a custom-developed algorithm;
- Segmented Linear Regression Thorough (SLRThorough) and Segmented Linear Regression Fast (SLRFast) – algorithms proposed and implemented by V. Stadnik [1];
- Permutation-Based Joinpoint Regression (SLRPermutation) – an algorithm based on a permutation test [2].

To perform the comparative analysis, a desktop application, EvaluationSLRAlgorithms, was developed in C# using Visual Studio. The application enables users to load datasets from a file or generate synthetic datasets, fit a segmented linear regression model using the four specified algorithms, evaluate model quality, and compare the algorithms in terms of accuracy and computational performance.

A series of experiments was conducted using this application (Fig. 1) to compare the accuracy and execution time of the algorithms on datasets with varying structures and noise levels. The goal was to identify the most versatile algorithm capable of providing stable results under different conditions. During the experiments, 500 datasets were generated for each test model described in [2]. For each algorithm, the accuracy of breakpoint number determination and the average execution time were evaluated.

The screenshot shows the 'Comparative test' tab of the application. It displays a table of results for 16 different test models (0-1 to 6-3). Each model is evaluated by four algorithms (1-4) based on accuracy and execution time. To the right of the table, there are configuration panels for each algorithm, including the number of algorithms (set to 4), and specific parameters like tolerance and half length for SLRAuto, SLRThorough-Stz, SLRFast-Stadnik, and SLRPermutation.

	1 Alg Acc (%)	2 Alg Acc (%)	3 Alg Acc (%)	4 Alg Acc (%)	1 Alg Avg Run Time (ms)	2 Alg Avg Run Time (ms)	3 Alg Avg Run Time (ms)	4 Alg Avg Run Time (ms)
0-1	100	0	0	98,8	166,8581828	2,5282212	0,0690826	493,6890294
0-2	100	80	63	99	154,4131292	0,254462	0,0222082	486,7619788
1-1	96,8	0	0	99,8	313,7232528	3,0147008	0,0874114	4181,9359398
1-2	90,8	0,6	0	99,8	274,6701302	1,7783414	0,0496512	2543,9135154
1-3	93,4	99,6	31,8	96,6	325,4763	0,908041	0,0304126	3056,1547924
2-1	88,4	0	0	99,8	247,25174	3,3634754	0,0838756	2823,982445
2-2	91	2,6	0	99	250,2427148	3,4242402	0,1483022	3017,0627176
2-3	84,6	99,6	45	73,4	201,7803258	0,7700494	0,026544	4992,3890628
3-1	85	22,2	35,4	40,4	533,003037	1,1462394	0,0334188	10698,6717538
3-2	85	29,4	76	0	494,4413794	1,106382	0,0328472	10823,5680662
3-3	81,4	13,6	77,6	0	459,4694284	1,1755868	0,0290658	10467,0546332
4-1	58,4	25	11	99,2	344,20676	1,1795418	0,0436754	10617,3393582
4-2	68	0	52	22,4	290,3428702	0,796717	0,028326	10531,4588646
5-1	89,4	50,4	6,6	100	755,3117902	1,4168848	0,0388948	10479,0236402
5-2	90,4	79,4	43,4	100	720,0118502	1,1821384	0,030289	10544,6309088
5-3	93,6	54,2	54	100	592,5180934	1,1296296	0,0367274	11311,2053724
6-1	80,6	0	0	99,4	746,572558	1,9458562	0,067223	10835,3867728
6-2	90,6	27,4	25,4	100	732,543196	1,3212922	0,03732	11249,5784058
6-3	84,8	1,8	35,8	100	603,2059184	0,902375	0,0349298	9622,1606774

Figure 1 – Experiment results

The analysis of the results showed the following. The SLRAuto demonstrated the most consistent breakpoint detection regardless of noise level or model complexity. The SLRThorough and SLRFast required individual parameter tuning for optimal performance; however, the SLRFast was the fastest among all compared algorithms. The SLRPermutation achieved the highest accuracy at the cost of significant computational expense, except in cases where the true breakpoints did not fall within the search grid.

References

1. Stadnik V. Segmented Linear Regression. URL: <https://www.codeproject.com/articles/Segmented-Linear-Regression#comments-section>.
2. Kim H.J., Fay M.P., Feuer E.J., Midthune D.N. Permutation tests for joinpoint regression with applications to cancer rates. Statistics in medicine. 2000. Vol. 19. P. 335–351. [https://doi.org/10.1002/\(SICI\)1097-0258\(20000215\)19:3<335::AID-SIM336>3.0.CO;2-Z](https://doi.org/10.1002/(SICI)1097-0258(20000215)19:3<335::AID-SIM336>3.0.CO;2-Z)

INTELLIGENT SYSTEMS AS A TOOL FOR SOLVING THE PROBLEM OF SUBJECTIVITY IN INPUT DATA

Zora I., Lishchynska L. (ivan.zora@gmail.com, llb@vntu.edu.ua)
Vinnitsia National Technical University

The problem of subjectivity in input data arises in many fields where part of the information is formed based on human judgment or manual observations. Such data often do not correspond to actual conditions and introduce distortions into modeling results. In traditional approaches, data of this type are either averaged or supplemented heuristically, which leads to the loss of important information about the underlying patterns of processes.

The purpose of this study is to explore the potential of intelligent systems to mitigate the lack of objective input data for their use in the mathematical modeling of transport systems. The application of intelligent systems, primarily neural network models, enables the implementation of a new approach with antecedent training on objective and reliable data, followed by the use of learned patterns to reconstruct or refine incomplete and subjective input data.

The mechanism of this approach is based on the “pre-training” stage, during which the network is trained on large datasets of verified information such as vehicle telemetry, GPS traces, ticket transaction statistics, or data from automatic passenger counters. This stage allows the model to form internal connections between features that describe real patterns of transport system behavior. Subsequently, when only partial or subjective data are available (for example, passenger assessments or fragmentary observations along routes), the neural network uses gained knowledge to reconstruct missing values or adjust them toward more probable objective magnitudes [1], [3], [4].

Research in the field of machine learning confirms the effectiveness of neural networks in handling incomplete or distorted data. Studies on learning with missing data show that an appropriately designed network architecture not only restores absent values but also reduces the overall prediction error

compared with classical statistical methods [1], [3]. There are also successful examples of applying recurrent neural networks to transport time series data to compensate for gaps in telemetry and sensor streams [4].

A significant example of practical application is the modeling of passenger flows in urban transport networks. In such cases, there is often a discrepancy between objective sources (sensors, ticketing systems) and subjective ones (passenger surveys, expert evaluations). The use of pre-trained neural networks helps to resolve this gap, improving the accuracy of route load estimation and waiting time prediction. A neural model trained on complete, objective datasets can emulate missing records or adjust subjective estimates in accordance with regularities characteristic of actual transport movement [6], [7].

Among the key advantages of this approach are higher prediction accuracy, model adaptability to changing conditions, reduced need for labor-intensive manual data cleaning, and the unification of information sources. Through the integration of data imputation algorithms into the network training process, the influence of random errors and inconsistencies between different sources of information is reduced [2], [5].

However, the method also has limitations that should be taken into account. These include the risk of overfitting to data obtained from a different context, the potential loss of unique features inherent in subjective observations, and the high computational complexity, which increases the requirements for hardware and financial resources [6], [8].

Thus, the use of intelligent systems to mitigate the subjectivity of input data represents a promising direction in the development of artificial intelligence. Training models on objective datasets enables the formation of an internal “map of regularities” within the system, which can subsequently be used for the evaluation, restoration, or cleansing of incomplete and subjective data. This creates a foundation for improving calculation accuracy, reducing the influence of the human factor, and ensuring the stability of decisions even under conditions of limited reliable observations.

References

- [1] M. Smieja, Ł. Struski, J. Tabor, B. Zieliński, and P. Spurek, “Processing of missing data by neural networks,” *Advances in Neural Information Processing Systems (NeurIPS)*, 2018. [Online]. Available: <https://arxiv.org/abs/1805.07405>
- [2] P. García-Laencina, J.-L. Sancho-Gómez, A. Figueiras-Vidal, and M. Verleysen, “Pattern classification with missing data: A review,” *Neural Computing and Applications*, vol. 19, pp. 263–282, 2010, doi: 10.1007/s00521-009-0295-6.
- [3] T. Emmanuel, T. Maupong, D. Mpoeleng, T. Semong, B. Mphago, O. Tabona, “A survey on missing data handling techniques for machine learning,” *Journal of Big Data*, vol. 8, no. 1, art. 140, 2021, doi: 10.1186/s40537-021-00516-9.
- [4] C. Che, D. Purushotham, K. Cho, D. Sontag, and Y. Liu, “Recurrent neural networks for multivariate time series with missing values,” *Scientific Reports*, vol. 8, art. 6085, 2018, doi: 10.1038/s41598-018-24271-9.
- [5] J. Yi, J. Lee, K. Joon Kim, S. Ju Hwang, E. Yang, “Why not to use zero imputation? Correcting sparsity bias in training neural networks,” *arXiv preprint*, arXiv:1906.00150, 2019.
- [6] Y. Yi, S. Lai, S. Li, J. Dai, W. Wang, J. Wang, “Optimization of missing value imputation for neural networks,” *Information Sciences*, vol. 642, art. 119029, 2023, doi: 10.1016/j.ins.2023.119029.
- [7] L. S. Iyer, “AI enabled applications towards intelligent transportation”, *Transportation Engineering*, vol. 5, 2021, 100083, <https://doi.org/10.1016/j.treng.2021.100083>
- [8] Å. Jevinger, C. Zhao, J.A. Persson, et al., “Artificial intelligence for improving public transport: A mapping study.” *Public Transp* 16, 99–158 (2024). <https://doi.org/10.1007/s12469-023-00334-7>

МАТЕМАТИЧНЕ ТА ПРОГРАМНЕ ЗАБЕЗПЕЧЕННЯ СИСТЕМИ СЕГМЕНТАЦІЇ ДІЛЯНОК ГОЛОВНОГО МОЗКУ ТА АНАЛІЗУ ЇХ АТРОФІЇ ДЛЯ ВИЯВЛЕННЯ ОЗНАК ХВОРОБИ АЛЬЦГЕЙМЕРА

Анісімов О.А., alexanisandr@gmail.com, Білозьоров В.Є.

Дніпровський національний університет імені Олеся Гончара

Хвороба Альцгеймера є найпоширенішою формою деменції та характеризується прогресуючою атрофією мозкових структур. Своєчасне автоматизоване виявлення ознак захворювання за допомогою методів медичної візуалізації має ключове значення для ранньої діагностики та розробки персоналізованих стратегій лікування. У роботі розглянуто використання методів глибинного навчання для сегментації структур головного мозку за даними магнітно-резонансної томографії (МРТ) та аналізу атрофічних змін, пов'язаних із нейродегенеративними процесами.

Метою роботи є створення математичного та програмного забезпечення для автоматизованої сегментації ділянок мозку та оцінки змін тканин, що можуть бути пов'язані з хворобою Альцгеймера. Об'єктом дослідження є процес аналізу медичних зображень, а предметом — методи сегментації та виділення структур головного мозку на МРТ-знімках.

Для розв'язання поставленої задачі було застосовано згорткову нейронну мережу SegResNet, реалізовану в середовищі Visual Studio Code із використанням бібліотек PyTorch Lightning та MONAI. Проведено попередню обробку даних, включно з нормалізацією, обрізанням фону, зміною розміру та просторовими аугментаціями. Як функцію втрат використано комбіновану Dice + Focal Loss, що підвищило точність сегментації при наявності класового дисбалансу. Для оцінки ефективності використано метрики Dice та Intersection-over-Union (IoU).

Тестування моделі виконано на 52 об'ємних МРТ-знімках розміром $64 \times 64 \times 64$. Результати наведено у таблиці 1.

Таблиця 1 — Результати сегментації МРТ-зображень мозку

Метрика	Значення
Dice coefficient	0.9722
IoU	0.8138

Високі значення Dice та IoU підтверджують ефективність обраної архітектури та підходу до навчання моделі. Система демонструє здатність точно сегментувати мозкові структури та відтворювати просторові особливості анатомії. Отримані результати свідчать про потенціал застосування моделі SegResNet для комп'ютерної підтримки медичних рішень та ранньої діагностики хвороби Альцгеймера.

Розроблене рішення може бути інтегроване у медичні системи аналізу МРТ-даних або використане як основа для подальших досліджень у сфері автоматизованої нейровізуалізації.

ПРОГНОЗУВАННЯ ЧАСУ ВИКОНАННЯ ФРАГМЕНТІВ КОДУ

Антонов В. С., azazaka2002@gmail.com, Турчина В. А.

Дніпровський національний університет імені Олеся Гончара

В сучасному світі є багато практичних задач, які можуть бути розв'язані тільки при наявності ефективних алгоритмів. Враховуючи реальні розмірності таких задач про програмне забезпечення, що реалізує відповідні алгоритми бажано б мати деякі відомості отримані апріорі. Це стосується прогнозування часу виконання програм - встановлення застосунків, виконання тривалих операцій, або планування порядку запуску файлів в обчислювальних кластерах. Проте існують задачі, де необхідно знати час виконання не тільки цілої програми, а й її фрагментів. Наприклад, це можуть бути системи м'якого реального часу, автоматичної паралелізації, бази даних, а також компілятори, інтерпретатори, та оптимізатори коду.

Одним із варіантів отримання оцінки часу виконання фрагментів коду є застосування методів аналізу даних до змінних програми, які використовуються цим фрагментом. Оскільки в такому випадку збір даних, їх аналіз та використання, а також сама робота програми виконується одночасно, то важливо, щоб обрані методи аналізу використовували помірний обсяг пам'яті та обчислювальних ресурсів. З цією ціллю було використано комбінований підхід Ridge-регресії та LightGBM-регресії.

Ridge-регресія – це метод, що є вдосконаленою версією лінійної регресії. Він включає L2-регуляризацію, що може зменшити ризик перенавчання. Серед його переваг можна виокремити швидке навчання, прогнозування та використання малого обсягу пам'яті. Проте, головним його недоліком є здатність моделювати лише лінійні зв'язки, тому він застосовується на початку роботи системи, коли зібрано відносно мало даних, і головним пріоритетом є швидкість роботи.

Наступним компонентом для прогнозування часу є LightGBM-регресія. LightGBM це метод градієнтного підсилення, основна ідея якого полягає в побудові ансамблю із декількох слабких моделей, де кожна виправляє помилки попередніх. Він був розроблений та представлений компанією Microsoft в 2017 році [1]. Порівняно із спорідненими альтернативами, цей метод дозволяє використовувати менше обчислювальних ресурсів, але може призводити до перенавчання на малих наборах даних. Крім того, великий набір гіперпараметрів хоч і робить його дуже гнучким, вимагає відповідної експертизи. Метод дозволяє моделювати складні нелінійні зв'язки, тому використовується, лише коли для фрагменту коду зібрано достатньо даних. І якщо цей фрагмент виконується часто, то доцільно підвищити точність оцінки часу перейшовши до LightGBM-регресії.

Систему було протестовано на алгоритмі побудови множини Мандельброта, та отримано наступні результати.

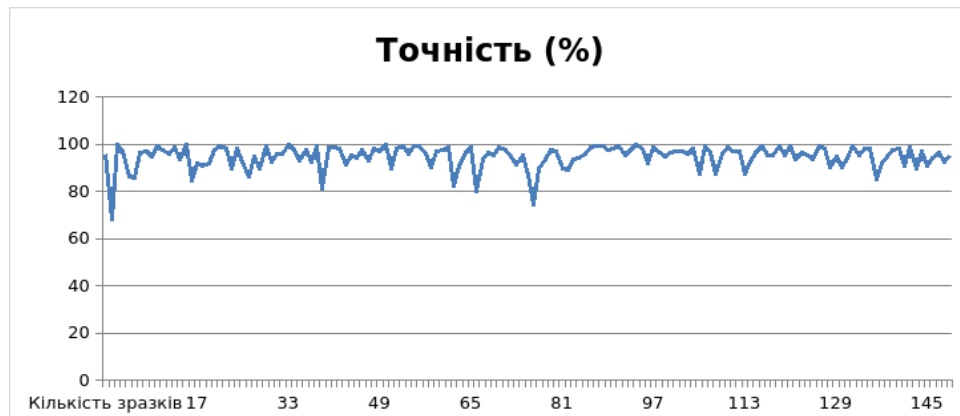


Рисунок 1 – Графічна залежність точності прогнозування від об'єму вибірки

Точність переважно тримається на 90%, із поодинокими випадками зменшення до 80%. Видно, що перехід на LightGBM-регресію позитивно впливає на результат, зменшуючи середню похибку.

Перелік використаних джерел

1. Ke G., Meng Q., Finley T., Wang T., Chen W., Ma W., Ye Q., Liu T.-Y. LightGBM: A Highly Efficient Gradient Boosting Decision Tree. *Advances in Neural Information Processing Systems*. 2017. Vol. 30. P. 3149–3157.

СИНТЕЗ СИСТЕМ КЕРУВАННЯ ПРОСТОРОВИМ РУХОМ ЛІТАЛЬНОГО АПАРАТА

Афанасьєв Д.К., den.4fanasev.2000@gmail.com

Зайцев В.Г., vadymzaytsev65@gmail.com

Дніпровський національний університет імені Олеся Гончара

Стрімкий розвиток безпілотних літальних апаратів (БПЛА) упродовж останніх десятиліть став одним із ключових напрямів сучасної науки і техніки. Зростання ролі БПЛА у цивільній та оборонній галузях супроводжується появою складних завдань з побудови надійних систем керування, здатних забезпечити стабільність польоту в умовах невизначеності, турбулентності, вітрових збурень і технічних обмежень.

Системи управління БПЛА розвивалися поступово — від класичних до інтелектуальних. До найраніших належать PID-регулятори, після були оптимальні методи[5], робастні схеми (наприклад, H_∞ -синтез), нелінійні методи. Серед новітніх тенденцій виокремлюють модельно-прогнозуюче керування (MPC), та інтелектуальні методи, що використовують штучні нейронні мережі, нечітку логіку й алгоритми машинного навчання [4].

На цьому фоні особливу увагу привертає синергетичний підхід до побудови систем керування. Основу цього методу закладено у працях Колеснікова та його школи [1]. Його суть полягає у введенні макрозмінних — узагальнених функцій стану, стабілізація яких забезпечує виконання цілей керування. Такий підхід дозволяє не контролювати кожную координату окремо, а забезпечити самоорганізацію системи навколо заданих інваріантів. У результаті у фазовому просторі формується синергетичний атрактор, який притягує траєкторії незалежно від початкових умов, що гарантує стійкість і точність [1; 2].

У роботі розглянуто моделювання руху ЛА у двовимірній площині. На основі синергетичного підходу побудовано аналітичні вирази керуючих впливів для різних типів траєкторій. Зокрема, передбачається дослідити

три основні режими руху: по прямій, по колу та за законом синусоїди [3]. Для кожної з них синтезовано відповідне керуюче прискорення, яке забезпечить стабільність траєкторії та інваріантність до зовнішніх збурень. Очікується, що результати моделювання дозволять порівняти ефективність законів керування та визначити оптимальні параметри макрозмінних для різних сценаріїв руху. Таким чином, синергетичний підхід поєднує аналітичну суворість і високу ефективність. Він є універсальним інструментом для побудови адаптивних алгоритмів керування, придатних до реалізації в реальних умовах. Його здатність забезпечувати плавне, стійке і робастне керування, що відкриває перспективи для створення автономних систем польоту, що діють у невизначеному середовищі. Використання моделювання у MATLAB дозволяє підтвердити ефективність методу та надати практичні рекомендації для побудови синергетичних систем керування БПЛА [1; 3; 6].

Бібліографічні посилання:

1. Колесников А. А. Синергетична теорія керування. – М.: Энергоатомиздат, 1994.
2. Ahifar A., Ranjbar N. A., Rahmani Z. Finite Time Terminal Synergetic Controller for Nonlinear Helicopter Model. – IJE Transactions B: Applications, 2019.
3. Легкоступ В. В., Маркевич В. Є., Шабан С. А. Методика отримання закону керування із використанням синергетичного підходу для задачі наведення ЛА вздовж гіперболи. – Доклади БГУИР, 2022.
4. Pucci D., Hua M.-D., Hamel T., Morin P., Samson C. Nonlinear Feedback Control of VTOL UAVs. – Aerospace Lab, Issue 8, 2014.
5. Sira-Ramirez H. V., Castro-Linares R., Liceaga-Castro E. A Liouvillian Systems Approach for Trajectory Planning-Based Control of Helicopter Models. – Int. J. of Robust and Nonlinear Control, 2000.
6. Оланд Е. Nonlinear Control of Fixed-Wing Unmanned Aerial Vehicles. – NTNU, Trondheim, 2014.

УЗАГАЛЬНЕНА ЗАДАЧА ПРО ПРИЗНАЧЕННЯ ЯК ІНСТРУМЕНТ ОПТИМІЗАЦІЇ ВИРОБНИЧИХ СИСТЕМ

Балейко А.С., Михальчук Г.Й., artembaleyko@gmail.com

Дніпровський національний університет імені Олеся Гончара

У сучасних умовах високої конкуренції та динамічних змін попиту на продукцію одним із ключових завдань для промислових підприємств є ефективне використання наявних виробничих ресурсів. Виробничі системи мають обмежену кількість обладнання, часові рамки та різні технологічні маршрути виготовлення продукції.

У таких умовах виникає потреба у застосуванні методів оптимізації, які дозволяють визначити найкращий спосіб розподілу ресурсів між завданнями з урахуванням усіх обмежень. Поширеним математичним інструментом моделювання виробничого процесу є узагальнена задача про призначення (General Assignment Problem, GAP), яка розширює класичну задачу про призначення, дозволяючи враховувати обмеження на ресурси та індивідуальні витрати для кожної пари «виконавець – завдання».

Розглянемо виробничу систему підприємства, що виготовляє вироби побутової хімії. Для планового тижня відомий список продуктів і обсяги, які потрібно виготовити. Кожен продукт може бути виготовлений лише на певних виробничих лініях, що поділяються на лінії, де відбувається приготування сировини (bulk lines) та лінії пакування готової продукції (package lines).

Швидкість виготовлення різних продуктів відрізняється, і кожен конвеєр має обмеження у часі роботи. Крім того, особливістю процесу є те, що під час приготування сировини у розмішувач заливається повний об'єм, незалежно від того, скільки продукції потрібно згідно з планом. Це може призводити до перевиробництва, якщо фактичний обсяг партії перевищує планову потребу.

Задача полягає у тому, щоб розподілити роботу між лініями сировини та пакування так, щоб якомога точніше виконати план виробництва, дотримуючись усіх технологічних та часових обмежень.

Задачу планування етапів виробництва можна подати у вигляді узагальненої задачі про призначення. У цьому контексті:

- кожен продукт розглядається як завдання;
- кожна пара (bulk line, package line) – як виконавець;
- функцію витрат (costs) можна визначити як різницю між заданим планом та фактичним виконанням. При цьому, втрати пов'язані з недовиконанням плану повинні вважатись більш значущими;
- ресурси – це доступний час роботи конвеєрів.

Мета задачі – мінімізувати загальні втрати, пов'язані з перевиробництвом або недовиконанням плану, дотримуючись обмежень щодо часу та можливості виробництва кожного продукту на певних лініях. Таким чином, задача з реального виробництва перетворюється на стандартну GAP-модель, яку можна розв'язувати за допомогою існуючих оптимізаційних методів і бібліотек.

Застосування узагальненої задачі про призначення для оптимізації виробничих систем дозволяє формалізувати складні процеси розподілу ресурсів і підвищити ефективність планування. GAP забезпечує можливість урахування обмежень часу, ресурсів і технологічних маршрутів, а її розв'язання з використанням сучасних інструментів, таких як OR-Tools, надає практичні результати для реальних підприємств.

ВИКОРИСТАННЯ ІНФОРМАЦІЙНИХ СИСТЕМ У МЕТОДАХ РОЗПАРАЛЕЛЕННЯ ПРОЦЕСУ ФОРМУВАННЯ ТРИВИМІРНИХ ГРАФІЧНИХ ЗОБРАЖЕНЬ

Бобко О. Л., oleksii.bobko@gmail.com

Ліщинська Л. Б., llb@vntu.edu.ua

Вінницький Національний Технічний Університет

Сучасні комп'ютерні технології тривимірної графіки є однією з найдинамічніших сфер розвитку інформаційних систем. Підвищення реалістичності зображень, розширення функціональних можливостей графічних рушіїв та зростання вимог до швидкодії обумовлюють необхідність оптимізації процесу рендерингу [1].

Метою дослідження є огляд і аналіз існуючих засобів підвищення продуктивності рендерингу шляхом балансування навантаження за допомогою інтелектуальних систем.

Інформаційні системи відіграють ключову роль у забезпеченні взаємодії між обчислювальними вузлами, керуванні потоками даних, контролі навантаження та зборі результатів обчислень. Їх інтеграція у процеси рендерингу забезпечує можливість створення адаптивних, масштабованих і керованих середовищ для високопродуктивної візуалізації тривимірних сцен [2].

Процес формування тривимірного зображення складається з послідовності етапів графічного конвеєра більшість з яких можуть бути реалізовані паралельно, що створює умови для застосування методів паралельного програмування на рівні потоків, пікселів або полігонів [3].

Особливу увагу приділено адаптивним інформаційним системам, які на основі даних про продуктивність вузлів змінюють стратегію розподілу завдань. Такі системи поєднують інструменти збору телеметрії

(Prometheus, Grafana), засоби контейнеризації (Docker, Kubernetes) та інтелектуальні алгоритми оптимізації [4].

Застосування інформаційних систем забезпечує автоматизоване керування ресурсами, зменшує обчислювальні витрати та дозволяє реалізувати масштабовані моделі високопродуктивного рендерингу.

Список використаних джерел

1. Романюк О. Н., Майданюк В. П., Бобко О. Л., Тітова Н. В., Романюк С. О. Використанням GPU для розпаралелення рендерингу // Advanced top technology. Електрон. текст. дані. 2024. № 3. С. 46-48.
2. Романюк О. Н., Завальнюк Є. К., Бобко О. Л. Використання штучного інтелекту для рендерингу тривимірних сцен // Матеріали XVIII Всеукраїнської науково-практичної WEB конференції аспірантів, студентів та молодих вчених «Комп'ютерні інтелектуальні системи та мережі», Кривий Ріг, 25-27 березня 2025 р. Кривий Ріг : Криворізький національний університет, 2025. С. 208-210.
3. Romanyuk O. N., Bobko O. L., Zavalniuk Y. K., Titova N. V., Romanyuk S. O., Stakhov O. Y. Analysis of graphics pipelines. Innovation in der modernen Wissenschaft: Innovative Technologie, Informatik, Sicherheitssysteme, Verkehrsentwicklung, Architektur und Bauwesen, Physik und Mathematik. Monografische Reihe «Europäische Wissenschaft». 2024. Buch 30. Teil 1. Chap. 4. Pp. 71-80.
4. Романюк О. Н., Бобко О. Л. Контейнеризація для розпаралелення рендерингу // Матеріали XIV Міжнародної науково-практичної конференції «Теорія і практика сучасної науки та освіти», м. Львів, 9-10 січня 2025 р. Електрон. текст. дані. Львів : Львівський науковий форум, 2025. С. 63-67.

АНАЛІЗ ЕФЕКТИВНОСТІ АЛГОРИТМІВ ГЛИБОКОГО НАВЧАННЯ З ПІДКРІПЛЕННЯМ У 2D-ІГРОВОМУ СЕРЕДОВИЩІ

Бовкун М.Ф., mikhailbovkun@gmail.com

Мацуга О.М., olga.matsuga@gmail.com

Дніпровський національний університет імені Олеся Гончара

Глибоке навчання з підкріпленням (Deep Reinforcement Learning, DRL) поєднує можливості нейронних мереж та класичного навчання з підкріпленням для створення агентів, здатних приймати оптимальні рішення через взаємодію із середовищем. Одним із ключових завдань під час розробки таких агентів є вибір алгоритму навчання, оскільки різні підходи відрізняються стабільністю, швидкістю збіжності та придатністю до певних типів середовищ. Метою роботи є аналіз основних класів алгоритмів DRL – value-based, policy-based та actor-critic [1] – і порівняння їх ефективності в розробленому 2D-середовищі «*Ghost of Ukraine*», у якому агент навчається керувати винищувачем для знищення ворожих цілей до того, як вони досягають лівого краю екрана.

Алгоритми типу value-based ґрунтуються на апроксимації функції цінності, яка оцінює очікувану сумарну винагороду для кожного стану або пари «стан–дія». Агент обирає дії, що максимізують цю оцінку. Основним представником є Deep Q-Network (DQN), який поєднує класичний Q-learning з нейронними мережами. Модифікації, такі як Double DQN та Dueling DQN, покращують стабільність і зменшують переоцінку значень. У межах середовища «*Ghost of Ukraine*» алгоритм DQN продемонстрував високу ефективність під час навчання агента в дискретному просторі дій, дозволяючи швидко опановувати базові маневри, зокрема ухилення та знищення ворогів.

Policy-based підходи, на відміну від value-based, безпосередньо моделюють політику агента – тобто функцію, що визначає ймовірність вибору певної дії у заданому стані. Оптимізація політики відбувається за

допомогою методів градієнтного підйому. У роботі застосовувався алгоритм Proximal Policy Optimization (PPO), який завдяки обмеженню зміни політики між ітераціями забезпечує стабільніше навчання. Під час експериментів PPO показав здатність до плавної адаптації поведінки агента у змінних умовах середовища, де швидкість руху ворогів і частота появи цілей могли варіюватися. Це дозволило агенту виробити більш узагальнену стратегію дій, стійку до нових сценаріїв.

Алгоритми типу actor-critic поєднують переваги обох попередніх підходів. Вони складаються з двох компонентів: *actor* (актор), який обирає дії, та *critic* (критик), який оцінює їхню ефективність через функцію цінності. Така архітектура дозволяє стабілізувати навчання та підвищити швидкість збіжності. У межах дослідження використовувався алгоритм Advantage Actor-Critic (A2C), який оцінює не лише цінність дій, а й їх перевагу (*advantage*), що сприяє більш точному оновленню політики. Під час тестування A2C забезпечив баланс між стабільністю PPO та швидкістю навчання DQN, демонструючи добру здатність до пристосування в умовах з високою динамікою подій.

Порівняльний аналіз показав, що кожен тип алгоритмів має свої переваги залежно від характеристик середовища. DQN найкраще працює у дискретних та детермінованих сценаріях, PPO – у неперервних і динамічних, тоді як A2C забезпечує компроміс між точністю, стабільністю та швидкістю навчання. Результати експериментів у середовищі «*Ghost of Ukraine*» підтвердили, що правильний вибір типу алгоритму є ключовим для досягнення ефективного навчання агента та стабільної поведінки в умовах складних багатовимірних задач.

Список використаної літератури

1. Morales M. Grokking Deep Reinforcement Learning. Manning, 2020. 472 p.

ЗАСТОСУВАННЯ ГЕНЕТИЧНОГО АЛГОРИТМУ ДЛЯ РОЗВ'ЯЗАННЯ ЗАДАЧІ КОМІВОЯЖЕРА З УРАХУВАННЯМ ЧАСУ ТА ВІДСТАНІ

Богданович Б.М., bogdan548584@gmail.com

Козакова Н.Л., kozakova.natali@gmail.com

Дніпровський національний університет імені Олеся Гончара

Задача комівояжера є класичною комбінаторною задачею оптимізації, яка належить до класу NP-повних. Її мета – знайти найкоротший (або найоптимальніший) маршрут, що проходить через усі задані точки лише один раз і повертається до початкової. У практичних застосуваннях, зокрема у міській логістиці, враховується не лише відстань, а й час пересування, який залежить від трафіку, черг на навантаження чи розвантаження, тощо.

Через обчислювальну складність задачі використання точних методів, таких як алгоритми гілок і меж або повний перебір, є неефективним для великої кількості вузлів. Тому доцільним є застосування евристичних та метаевристичних підходів, серед яких генетичні алгоритми зарекомендували себе як гнучкий і потужний інструмент пошуку наближених оптимальних рішень.

У роботі розроблено програмну реалізацію генетичного алгоритму для розв'язання задачі комівояжера з урахуванням двох критеріїв – часу та відстані. Запропонований підхід має практичне значення у сфері логістики та доставки, зокрема для служб типу Glovo, Bolt Food, Nova Poshta та інших компаній, що оптимізують маршрути в умовах реального часу.

У розробленій системі генетичний алгоритм представлено як популяційний метод еволюційної оптимізації, де кожен можливий маршрут розглядається як особина (хромосома). Основні оператори

алгоритму: селекція (вибір найкращих рішень за комбінованим критерієм часу та відстані), кросовер (обмін частинами маршрутів між батьківськими особинами), мутація (випадкове переставлення точок маршруту для запобігання передчасній конвергенції). Функція пристосованості розраховується як зважена сума часу та довжини маршруту, що дозволяє адаптувати алгоритм під різні сценарії – наприклад, коли важливішим є час доставки або мінімізація відстані. Проведено серію експериментів для оцінки впливу параметрів алгоритму. Встановлено, що підвищення кількості поколінь покращує якість рішення, але збільшує час виконання; оптимальна ймовірність мутації знаходиться в межах 0.05–0.1; комбінування критеріїв часу та відстані дозволяє отримати більш реалістичні маршрути порівняно з одноцільовою оптимізацією.

Розроблений алгоритм продемонстрував стабільну збіжність до оптимального рішення на тестових наборах даних та ефективність у випадках із великою кількістю вузлів. Запропонований підхід є ефективним для задач логістики та доставки, де важливі одночасно швидкість і економічність маршруту.

У подальших дослідженнях доцільно враховувати часові обмеження (доступність точок у певні години); моделювати динамічний трафік, коли ваги ребер змінюються залежно від часу доби; інтегрувати реальні географічні дані (API карт Google або OpenStreetMap) для апроксимації умов реального міста. Такі розширення дозволять зробити модель більш наближеною до реальних логістичних сценаріїв і підвищити її практичну цінність.

ПРО КОМПЛЕКСИ НАВЧАЛЬНИХ ТРЕНАЖЕРІВ ІНЖЕНЕРНОЇ ТЕХНІКИ ДЛЯ ПІДГОТОВКИ ВІЙСЬКОВИХ ФАХІВЦІВ

Богомаз В.М., Богомаз О.В.

wbogomas@i.ua

Український державний університет науки і технологій

В умовах військових дій на території України підготовка військових фахівців має особливе значення для держави. Як відомо, для підготовки таких фахівців запроваджена спеціалізована освіта, яка спрямована на набуття компетентностей у сфері професійної військової діяльності для здобуття відповідних ступенів освіти та рівнів військової освіти.

На основі специфіки кожної освітньої програми у здобувачів освіти формуються відповідні військово-професійні та військово-спеціальні компетентності, які відображають специфіку конкретної предметної сфери професійної діяльності (спеціалізації) військового фахівця.

Для якісної підготовки здобувачів за спеціальністю машинобудування наявна матеріально-технічна база закладів освіти повинна забезпечувати проведення всіх видів та форм навчальних занять для набуття практичних навичок здобувачами освіти із управління підрозділами при застосуванні та обслуговуванні відповідного озброєння та військової техніки. Для цього застосовуються навчально-тренувальні та навчально-лабораторні комплекси. Вони обладнуються тренажерами, розгорнутими зразками озброєння та військової техніки, необхідними для відпрацювання практичних, групових занять з освітніх компонент військово-спеціальної підготовки, а також вдосконалення практичних навичок щодо застосування, експлуатації, технічного обслуговування та відновлення військової техніки.

З метою фортифікаційного обладнання територій сьогодні на кожному об'єкті застосовується комплект інженерної техніки, необхідний для механізації виконання всіх видів робіт.

Для реалізації ефективного освітнього процесу керівників на об'єктах виконання завдань заклади освіти обладнується навчальними тренажерами кожної одиниці техніки з відповідного комплексу (наприклад, тренажери екскаватора, навантажувача, бульдозера, автомобільного крана тощо). Спеціалізоване програмне забезпечення кожного навчального тренажера містить відповідні вправи, які використовуються в реальній обстановці при будівництві фортифікаційних споруд та дає можливість імітувати наявність артилерійських обстрілів, мінно-вибухових загороджень, а також дії безпілотних літальних апаратів при різних погодних умовах, часу доби, зміни типу ґрунту і т.і.

Для формування компетентностей, які необхідні керівнику на даному об'єкті створено програмне забезпечення, яке дає можливість об'єднати роботу різних машин (їх тренажерів) на одному об'єкті та керівнику управляти процесом їх роботи в реальному часі, виходячи із завдань, які ставляться підрозділу на даному об'єкті та умов, які можуть скластися внаслідок виходу з ладу техніки з різних причин. Об'єкти виконання завдань з їх характеристиками (рельєф місцевості, конфігурація та складність побудови необхідних фортифікаційних споруд, інтенсивність обстрілів) формує інтелектуальна система із застосуванням штучного інтелекту, яка враховує реальні об'єкти та умови на них. Також система пропонує різні за складністю завдання, враховуючи досвід проходження попередніх завдань кожним здобувачем та надає рекомендації з правильного їх виконання, виходячи з допущених раніше помилок.

Отже, описаний комплекс тренажерів та його спеціалізоване програмне забезпечення сприяє формуванню компетентностей у здобувачів освіти, необхідних керівнику на об'єктах виконання завдань за призначенням, та в цілому удосконалює підготовку військових фахівців.

ІНТЕРАКТИВНИЙ ВЕБ-РЕСУРС З ВИВЧЕННЯ ІНФОРМАТИКИ В КОНТЕКСТІ НОВОЇ УКРАЇНСЬКОЇ ШКОЛИ

Бодю К., Вовк С., bodyukiara@gmail.com, vovk_s_m@ukr.net

Дніпровський національний університет імені Олеся Гончара (www.dnu.dp.ua)

На даний час в Україні у сфері загальної середньої освіти впроваджується реформа Нової Української Школи (НУШ) [1]. Її сутність полягає в оновленні змісту та підходів до викладання навчального матеріалу з метою набуття учнями основних компетентностей, їх подальшого розвитку й удосконалення, зокрема, у сфері інформаційних технологій. Враховуючи стрімке введення систем штучного інтелекту (ШІ) у різні сфери життєдіяльності людини, реалізація задач реформи НУШ вимагає розробки, опрацювання і широкого впровадження таких систем і у сфері освіти. У поданій роботі пропонується веб-ресурс із інтегрованим ШІ-агентом, що призначений для допомоги учням під час засвоєння навчального матеріалу і розв'язання завдань з інформатики.

Інтерактивний веб-ресурс для вчителів та учнів складається з бази даних та веб-інтерфейсу, який дозволяє взаємодіяти з контентом бази даних у режимі реального часу, а саме виконувати запис завдань у базу даних, зчитувати завдання з бази даних, записувати відповіді на завдання та оцінки, реєструвати користувачів, відображати статистику і т.і. Програмна реалізація веб-інтерфейсу виконана на мові Python із застосуванням фреймворків Flask, Bootstrap, Jinja2, Flask Blueprint та ін. База даних у пілотній версії реалізована на основі технології SQLite. До веб-ресурсу підключений ШІ-агент моделі OpenAI-mini-5, якому надано відповідні інструкції для його роботи. У подальшому, при розгортанні повноцінної версії веб-ресурсу, зазначену модель ШІ планується замінити на більш потужну модель із більшою кількістю токенів, оскільки передбачається робота великої кількості учнів.

У інтерактивному веб-ресурсі реалізовані ролі учня, викладача та адміністрації школи. Адміністрація школи має можливість створювати і керувати класами та персоналом школи. Вчителі мають можливість створювати завдання трьох типів, а саме завдання з відкритою відповіддю, тестові завдання і завдання із програмування. До завдань із програмування у вчителя є можливість увімкнути або вимкнути функцію допомоги ІІІ-агента. У випадку увімкнення ІІІ-агента учень отримує можливість на сторінці виконання завдання запросити у модальному вікні допомогу ІІІ (рис.1).

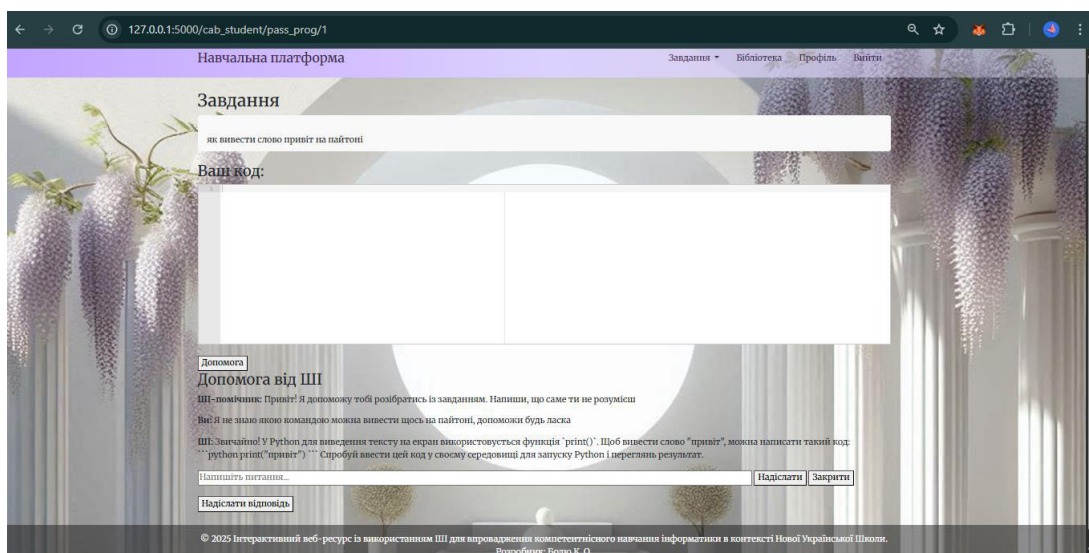


Рис.1. Сторінка учня із запитом допомоги та відповіддю ІІІ

У доповіді наведено опис розробленого функціоналу інтерактивного веб-ресурсу і обговорюються особливості його використання та, зокрема, перспективи його інтеграції із сервісами електронних журналів.

Бібліографічні посилання

1. Міністерство освіти і науки України / Нова українська школа – Електронний ресурс: <https://mon.gov.ua/tag/nova-ukrainska-shkola?&tag=nova-ukrainska-shkola>

ПРО ОПТИМІЗАЦІЮ НАВАНТАЖЕННЯ НА ІНФРАСТРУКТУРУ МЕТОДАМИ МАШИННОГО НАВЧАННЯ ТА АГЕНТНИМИ ТЕХНОЛОГІЯМИ

Божуха Д.І., bozhukha_d@365.dnu.edu.ua

Байбуз О.Г., baibuz_o@365.dnu.edu.ua

Дніпровський національний університет імені Олеся Гончара

Оптимізація ресурсів у хмарному середовищі стає критично важливою через зростання обсягів обчислення, використання CPU, RAM, динамічної зміни навантаження та необхідність забезпечення високої доступності сервісів.

Одним з ключових напрямів розвитку хмарних обчислень є впровадження інтелектуальних агентів, які виконують моніторинг, управління та адаптацію ресурсів хмарного сервісу у режимі реального часу. Такі агенти взаємодіють між собою та N центральними оркестраторами, забезпечуючи автономне масштабування та самовідновлення систем. Поєднання агентних технологій, методів, алгоритмів та підходів машинного навчання (Machine Learning, ML) дозволяє створювати моделі прогнозування навантаження, що забезпечують проактивну оптимізацію — тобто передбачення майбутнього попиту на ресурси та попереднє їх розподілення [1].

Розглянуто сучасні підходи до підвищення ефективності хмарних інфраструктур шляхом використання агентно-орієнтованих систем, технологій аналізу даних та методів прогнозування навантаження.

Використання моделей глибинного навчання (зокрема, LSTM-модуль для запам'ятовування значення як на короткі, так і довгі проміжки часу) та агенти штучного інтелекту надали змогу ефективно оцінювати переваги ресурсів і сховищ, що сприяє зменшенню затримок, оптимізації енергоспоживання та вартості обслуговування. Досліджено алгоритм з

методом розподілу ресурсів на основі агентної системи в хмарному середовищі з використанням методів навчання з підкріпленням [2], що демонструє покращення корисності ресурсів, доступності ресурсів та масштабованості ресурсів порівняно зі статичним розподілом ресурсів.

Експерименти проводилися з модельованими даними при використанні навантажувального тестування API-застосунку через Postman Collection Runner та Apache JMeter у режимі імітації трафіку при різних параметрах тестування (тип запиту, кількість паралельних запитів, час відповіді, відсоток помилок). Метриками оцінювання моделі обрано час відповіді (response time), час відгуку під навантаженням (latency under load), завантаженість CPU/Memory під час роботи, середня доступність сервісів (% uptime). У якості умовного середовища обрано базовий еталон Local (Docker), моделювання на базі JMeter: Simulated AWS, Microsoft Azure, GCP.

Інтеграція агентно-орієнтованих систем DevOps-підходами (наприклад, автоматичним CI/CD у Kubernetes-кластерах) дозволяє реалізувати безперервну адаптацію інфраструктури до змін у навантаженні користувачів або сервісів. Сучасні методи оптимізації хмарних інфраструктур, що базуються на агентному управлінні, аналізі даних і прогнозуванні, забезпечують підвищення продуктивності, зниження витрат і поліпшення якості сервісів. При вирішенні питання щодо оптимізації навантаження на інфраструктуру особливу увагу приділяють практичним результатам, які спрямовані на розвиток колективного навчання агентів (multi-agent reinforcement learning) для повної автономізації управління хмарними середовищами.

Бібліографічні посилання

1. Bodra D and Khairnar S (2025) Machine learning-based cloud resource allocation algorithms: a comprehensive comparative review. *Front. Comput. Sci.* 7:1678976. DOI: 10.3389/fcomp.2025.1678976
2. Shyam, Gopal & Bharti, Priyanka. (2023). Multi-agent Systems for Resource Allocation in Cloud Computing. 1-5. DOI:10.1109/InC457730.2023.10262945.

ЗАСТОСУВАННЯ МЕТОДІВ НЕЧІТКОЇ ЛОГІКИ ДЛЯ ОПТИМІЗАЦІЇ МІСЬКИХ ПЕРЕВЕЗЕНЬ В УМОВАХ НЕВИЗНАЧЕНОСТІ

Бондар О.О., msolesyabondar@gmail.com

Козакова Н.Л., kozakova.natali@gmail.com

Дніпровський національний університет імені Олеся Гончара

У роботі сформульовано та реалізовано підхід до оптимізації міських перевезень в умовах невизначеності, заснований на методах нечіткої логіки. Основною перевагою розробленої моделі є здатність враховувати численні фактори, що важко формалізуються у вигляді точних чисел – зокрема, лінгвістичні оцінки експертів щодо попиту на перевезення, стану дорожньої мережі, погодних умов тощо.

Проблема ефективного управління міськими транспортними потоками є однією з ключових у сучасних урбаністичних системах. Нестабільність попиту, вплив погодних факторів, дорожні події та інші випадкові чинники створюють високий рівень невизначеності, який ускладнює використання традиційних детермінованих методів оптимізації. Методи нечіткої логіки дають змогу формалізувати експертні знання, представлені у вигляді лінгвістичних змінних, і створювати адаптивні системи, що працюють із нечіткими, неточними або неповними даними.

Розроблена система базується на механізмі нечіткого виводу за Мамдані, у якому використовуються трикутні функції належності для входних і вихідних змінних. До входних параметрів моделі віднесено: рівень попиту (низький, середній, високий); стан доріг (добрий, задовільний, поганий); погодні умови (сприятливі, помірні, несприятливі). Вихідною змінною є пріоритетність маршруту, що визначає порядок розподілу обмеженого транспортного ресурсу.

Базу правил сформовано у вигляді лінгвістичних висловлювань типу:

Якщо попит високий і стан доріг задовільний, то пріоритет маршруту високий.

Для реалізації програмного модуля використано мову Python і бібліотеки scikit-fuzzy, tkinter та matplotlib. Система має графічний інтерфейс, що дозволяє користувачу вводити вихідні дані у зручній формі, спостерігати за процесом нечіткого виводу та отримувати візуалізовані результати у вигляді графіків і поверхонь рішень.

Проведені експериментальні дослідження підтвердили ефективність і коректність роботи системи. Модель адекватно реагує на зміну вхідних параметрів, забезпечує логічний розподіл транспортних ресурсів між маршрутами та демонструє адаптивну зміну пріоритетів залежно від умов середовища. Зокрема, у випадках погіршення дорожнього стану або несприятливих погодних умов система автоматично знижує пріоритет маршрутів у цих районах, пропонуючи альтернативні напрями або зменшення інтенсивності перевезень.

Отримані результати можуть бути використані у системах підтримки прийняття рішень у сфері міського транспорту – зокрема, для планування перевезень у години пік, визначення доцільності перенаправлення маршрутів, прогнозування навантаження на окремі ділянки транспортної мережі. Запропоновану методику також можна адаптувати для логістики, диспетчеризації служб таксі, вантажних перевезень та управління потоками у транспортних вузлах.

Запропоновано ефективний підхід до оптимізації міських перевезень на основі методів нечіткої логіки, який дозволяє враховувати невизначеність, суб'єктивні експертні оцінки та динамічні зміни транспортного середовища. Реалізований програмний модуль підтвердив практичну придатність методики для використання у системах підтримки прийняття рішень у сфері міського транспорту.

МАТЕМАТИЧНІ ОСНОВИ ПРОЕКТУВАННЯ АПАРАТНОГО ПРИСКОРЮВАЧА ДЛЯ ЗГОРТКОВОГО ШАРУ НЕЙРОННОЇ МЕРЕЖІ

Бочек М.Ю., Шишканова Г.А., Коротунова О.В.

kolabocek@gmail.com, shganna@ukr.net, okorotunova@gmail.com

Національний університет «Запорізька політехніка»

Спеціалізований блок ІР-ядро для обчислення згортки – це апаратно реалізований модуль, оптимізований для виконання операцій згортки (convolution), які є ключовими в обробці зображень, відео та нейронних мережах, зокрема згорткових нейронних мережах (CNN). ІР-ядро – це готовий до використання модуль цифрової логіки, який можна інтегрувати в більшу апаратну систему. Воно може бути реалізоване на ПЛІС (FPGA) – програмований логічний інтегральний схемі, що дозволяє змінювати конфігурацію апаратури після виготовлення. А також як ASIC (Application-Specific Integrated Circuit) – спеціалізований інтегральний схемі, розроблений для виконання конкретного завдання з максимальною ефективністю.

Операція згортки – це множення і сумування значень пікселів з ядром (фільтром), що ковзає по зображенню. ІР-ядро згортки виконує ці обчислення апаратно, що дозволяє: паралелізувати обчислення, тобто одночасно обробляти кілька пікселів або каналів; зменшити затримки – на відміну від CPU, який виконує інструкції послідовно, FPGA/ASIC можуть виконувати багато операцій одночасно; знизити енергоспоживання, тобто спеціалізовані схеми споживають менше енергії, ніж універсальні процесори. Відомими прикладами є Xilinx DSP48E, де блоки використовуються для реалізації згорток у FPGA та Intel OpenVINO IP Cores – для прискорення нейронних мереж.

Метою цієї роботи є показати як математичний апарат становить основу сучасних проривів у проектуванні апаратного прискорювача для

згорткового шару нейронної мережі. Математика дає точне формальне визначення операції, яку ми хочемо виконати.

Згортка двовимірного вхідного зображення I з фільтром K розміром $t \times t$ описується формулою:

$$O_{m,n} = (I \cdot K)_{m,n} = \sum_{i=0}^{t-1} \sum_{j=0}^{t-1} I(m+i, n+j) \cdot K(i, j),$$

де O – вихідне зображення, m, n – координати у вихідній матриці, i, j – координати для ітерації по фільтру. Ця формула описує проходження фільтра-детектора (наприклад, виділення країв) по всьому зображенню. Кожен елемент фільтра $K(i, j)$ – це "вага", яка навчається нейромережею.

Для більш ефективної реалізації на апаратному забезпеченні згортку часто перетворюють на операцію матричного множення. Цей метод називається `im2col` (image to column). Спочатку ми "розгортаємо" вікна вхідного зображення, з якими взаємодівав фільтр K , в один великий вектор-стовпець.

Потім "розгортаємо" всі фільтри в один великий вектор-рядок (або матрицю, якщо фільтрів багато). В результаті згортка перетворюється на одне велике матричне множення.

$$O = X \cdot W + B,$$

де X – матриця, отримана з вхідного зображення методом `im2col` (кожен її рядок – це "розгорнуте" вікно зображення), W – матриця ваг (кожен її стовпець – це "розгорнутий" один фільтр $t \times t$), B – вектор зміщень (bias).

O – вихідна матриця, яку потім можна перетворити назад у вхідну.

Це важливо, тому що апаратне забезпечення може бути дуже ефективно оптимізоване саме під операцію матричного множення. Таким чином, математика – це мова специфікації, а інтелектуальні системи створюють апаратний інструмент, ідеально заточений під цю специфікацію. Без математики ми не знали б, що будувати, а без апаратних інтелектуальних систем ми не змогли б побудувати це ефективно.

ОПТИМІЗАЦІЯ ВИБОРУ ФУНКЦІЙ НАЛЕЖНОСТІ ТА МЕТОДУ ДЕФУЗЗИФІКАЦІЇ ДЛЯ ПІДВИЩЕННЯ НАДІЙНОСТІ ФАНР-СИСТЕМИ

Вакульчик С.О., vakylchik123@gmail.com

Байбуз О.Г., baibuz_o@365.dnu.edu.ua

Дніпровський національний університет імені Олеся Гончара

Вступ та аналіз досліджень. Гібридний метод Аналітичної Ієрархії на основі Нечіткої Логіки (ФАНР) є потужним інструментом для задач багато-критеріального прийняття рішень (MCDM), ефективно долаючи недолік класичного МАІ — нездатність адекватно обробляти суб'єктивність і невизначеність експертних оцінок [1]. Заміна жорсткої числової шкали на нечіткі лінгвістичні змінні підвищує гнучкість і достовірність ранжування альтернатив [1]. Попередні дослідження, проведені авторами, підтвердили загальні переваги ФАНР над МАІ. Однак, подальше підвищення надійності системи вимагає фокусування на внутрішній оптимізації. Проблема полягає в тому, що вибір форми функцій належності (наприклад, трикутні чи трапецієподібні) та методу дефuzziфікації (наприклад, центроїд чи середній максимум) є джерелом додаткової суб'єктивності з боку розробника, що може впливати на стабільність кінцевих результатів [2]. Метою роботи є дослідження, порівняння та обґрунтування оптимальної конфігурації функціональних елементів ФАНР для мінімізації цієї суб'єктивності та максимізації надійності системи.

Основні положення. Для вирішення поставленої задачі дослідження фокусується на двох ключових етапах ФАНР:

1. **Форма Нечітких Чисел (Функції належності):** Аналізується, чи забезпечують Трапецієподібні Нечіткі Числа (ТрНЧ) вищу гнучкість і точність порівняно з більш поширеними Трикутними Нечіткими

Числами (ТНЧ) [2]. ТрНЧ дозволяють краще моделювати діапазон, у якому експерт абсолютно впевнений у своєму судженні.

2. **Метод Дефузифікації:** Проводиться порівняльний аналіз між методом Центроїда (Center of Area, CoA), який вважається найбільш точним, оскільки враховує всю площу нечіткої ваги, і іншими, менш ресурсоемними методами [3].

Для обґрунтування вибору буде застосовано імітаційне моделювання, яке дозволить оцінити стійкість ранжування альтернатив при різних комбінаціях цих параметрів [4]. Найбільш надійною буде визнана та конфігурація, яка демонструє найменшу варіативність фінальних пріоритетів при невеликих змінах у початкових експертних оцінках.

Висновок: Проведене дослідження дозволить розробити методологію оптимізованого ФАНР, яка мінімізує суб'єктивний вплив на етапах фузифікації та дефузифікації. Очікується, що використання Трапецієподібних Нечітких Чисел у поєднанні з методом Центроїда забезпечить найкращу надійність і стабільність ФАНР-системи. Результати роботи будуть використані для стандартизації та вдосконалення програмної реалізації ФАНР в гібридній системі підтримки прийняття рішень.

Список використаних джерел:

1. ДОСЛІДЖЕННЯ ГІБРИДНОЇ СИСТЕМИ ПІДТРИМКИ ПРИЙНЯТТЯ РІШЕНЬ НА ОСНОВІ АНР ТА НЕЧІТКОЇ ЛОГІКИ ДЛЯ ОБРОБКИ НЕВИЗНАЧЕНОСТІ ЕКСПЕРТНИХ ОЦІНОК / С. О. Вакульчик, О. Г. Байбуз. – Дніпровський національний університет імені Олеся Гончара.
2. Li Y., Liu H. S., Weng W. G. The choice of fuzzy numbers in fuzzy AHP // Expert Systems with Applications. — 2011. — Vol. 38, № 8. — P. 9771–9776.
3. Mendel J. M. Fuzzy sets for systems engineers: A tutorial // Proceedings of the IEEE. — 1995. — Vol. 83, № 3. — P. 345–377.
4. Büyükoçkan G., Feyzioglu O. A study on the stability of the fuzzy AHP with respect to different defuzzification methods // Information Sciences. — 2004. — Vol. 165, № 3-4. — P. 237–254.

ВИКОРИСТАННЯ АУГМЕНТАЦІЇ ДАНИХ ДЛЯ ПІДВИЩЕННЯ СТІЙКОСТІ НЕЙРОННИХ МЕРЕЖ

Васько А.О., Вовк С.М.

aranooke@ukr.net

Дніпровський національний університет імені Олеся Гончара (www.dnu.dp.ua)

У сучасних системах комп'ютерного зору нейронні мережі часто стикаються зі спотвореннями даних — шумом, поворотами, частковим маскуванням об'єктів. Такі зміни можуть суттєво знижувати точність класифікації, і тому підвищення стійкості моделей до спотворень даних є актуальною задачею машинного навчання.

У поданій доповіді обговорюється експеримент з аугментації даних для підвищення стійкості нейронних мереж із використанням попередньо навченої моделі Vision Transformer (ViT-B/16), що базується на механізмі самоуваги (*self-attention*) [1].

Для тестування було відібрано 1000 зображень тварин (зокрема котів і собак) з відкритого набору даних ImageNet, до яких застосовувалися різні типи спотворень: додавання гаусівського шуму, поворот зображення на випадковий кут від 10° до 45° , а також часткове маскування об'єкта чорними або сірими фрагментами. Загальна кількість зображень після аугментації, що підлягала дослідженню, становила 1000 зображень. Точність класифікації вимірювалась для кожного типу спотворення окремо.

Результати показали, що ViT демонструє високу стійкість до глобальних спотворень, зокрема при додаванні шуму або повороті. Точність класифікації без спотворень становила **92.3%**, при шумі — **85.1%**, при повороті — **80.4%**, а при частковому маскуванні — **74.6%**. Зменшення точності при маскуванні свідчить про чутливість моделі до втрати локальної інформації.

Таким чином, Vision Transformer ефективно зберігає узагальнюючу здатність при більшості спотворень, але потребує додаткових підходів для підвищення стійкості до локальних втрат зображення. Перспективним напрямом подальших досліджень є комбінування трансформерних і згорткових архітектур для покращення адаптивності до різних типів спотворень навчальних даних.

На рис. 1 подано приклади спотворених зображень із набору ImageNet, що використовувалися під час експерименту (шум, поворот, маскування).

На рис. 2 наведено графік зміни точності класифікації моделі ViT-B/16 залежно від типу застосованого спотворення.

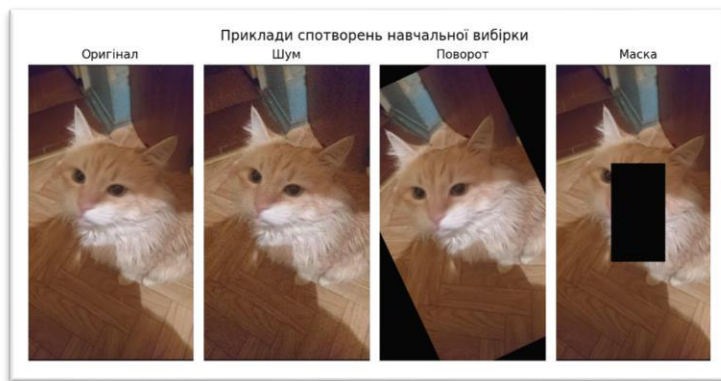


Рис.1 – Приклади аугментованої вибірки: зліва направо – оригінальне зображення кота: з шумом, поворотом та маскою

Бібліографічні посилання

- [1] Khalil, M., Khalil, A. & Ngom, A. (2023). A Comprehensive Study of Vision Transformers in Image Classification Tasks. Department of Computer Science, University of Windsor, Windsor, Ontario, Canada. arXiv:2312.01232.

МОДЕЛЮВАННЯ РОЗВ'ЯЗКУ ІГРОВИХ ЗАДАЧ ПЕРЕХОПЛЕННЯ КЕРОВАНИХ ЦІЛЕЙ

Вишенський В.І., vyshenskyi@ukr.net

Чикрій А.О., g.chikrii@gmail.com

Інститут кібернетики ім. В.М.Глушкова НАНУ

Останнім часом спостерігається зростання інтересу до теоретичного та прикладного аналізу класичних методів перехоплення цілей і ухиляння від переслідування, які раніше вважалися достатньо вивченими. Така актуалізація зумовлена стрімким розвитком високошвидкісних засобів ураження, а також поширенням маневрених, хоча й малошвидкісних, безпілотних літальних апаратів. У зв'язку з цим дослідження різноманітних стратегій перехоплення та втечі, а також їх практичного застосування, набуває нового значення.

У представлених дослідженнях основну увагу зосереджено на методі розв'язувальних функцій [1], який є тісно пов'язаним із першим прямим методом Л. С. Понтрягіна. Зазначений підхід забезпечує ефективне використання сучасного апарату опуклого аналізу, теорії багатозначних відображень та селекторів для обґрунтування ігрових моделей і отримання на їх основі змістовних результатів у широкому класі задач керування в умовах конфлікту. Зокрема, метод дозволяє строго теоретично обґрунтувати класичні закони паралельного переслідування, погонної кривої Ейлера та пропорційної навігації [2, 3].

Для чисельного моделювання розв'язків ігрових задач методом розв'язувальних функцій у даній роботі використано систему GNU Octave. Octave є мовою програмування високого рівня, орієнтованою переважно на наукові та інженерні обчислення. Вона забезпечує засоби для чисельного розв'язання лінійних і нелінійних задач, а також для проведення інших

обчислювальних експериментів, використовуючи синтаксис, здебільшого сумісний із MATLAB.

Проведено чисельне моделювання процесу переслідування за різних сценаріїв поведінки втікача, як у двовимірному, так і тривимірному просторі. У випадку, коли втікач рухається з максимальною швидкістю по прямій під довільним ненульовим кутом до відрізка, що з'єднує початкові положення гравців, метод паралельного переслідування забезпечує менший час упіймання. У цьому випадку геометричне місце точок перехоплення відповідає колу Аполлонія, тоді як для погонної кривої — це радик Паскаля. За умови маневрування втікача ефективнішими можуть виявитися альтернативні стратегії керування переслідувача. Наведено приклади, в яких кожен із розглянутих методів переслідування демонструє перевагу за критерієм мінімального часу упіймання.

В майбутньому планується провести математичне моделювання для задач перехоплення втікача групою переслідувачів.

Посилання

1. Chikrii, A.A. Conflict-controlled processes / A.A. Chikrii. Boston; London; Dordrecht: Kluwer Academ. Publ., 1997. 424 pp.
2. Chikrii A.A., Matichin I.I. Resolving Functions in Parallel and Pure Pursuit. Journal of Automation and Information Sciences, volume 35, issue 11, 2003, pp. 1-6. <https://doi.org/10.1615/JAutomatInfScien.v35.i11.10>
3. Ignatenko, A. P., Chikrij, A. A. On Substantiation of the Proportional Navigation Method in the Simple Pursuit Problem. Journal of Automation and Information Sciences. Volume 36, Issue 1, 2004, pp. 19-27. <https://doi.org/10.1615/JAutomatInfScien.v36.i1.30>

ПІДХІД ДО ВИЯВЛЕННЯ ЦИФРОВОЇ ВТОМИ ЗА ПОВІДОМЛЕННЯМИ ІЗ ВИЗНАЧЕННЯМ СЕГМЕНТІВ СПІЛКУВАННЯ

Віт Р.В., vit.roman.vit@gmail.com, **Мазурець О.В.**, exe.chong@gmail.com

Хмельницький національний університет

Проблематика цифрової втоми вимагає інструментів, здатних інтерпретувати зв'язок між характером цифрової комунікації користувачів і проявами когнітивного виснаження у динамічних середовищах взаємодії [1]. Ускладнення виникають через непряму вираженість емоційних сигналів, неоднорідність стилів спілкування та фрагментарність даних у різних платформах, що робить класичні лексиконні або суто статистичні підходи недостатніми. Додатковим викликом є потреба в пояснюваності результатів для подальшої інтеграції у системи підтримки рішень і добробуту, де прозорість механізмів виявлення має не менше значення, ніж точність. Відповідне рішення повинно поєднувати контекстну чутливість до мовних нюансів, стійкість до шуму й асинхронності комунікацій та можливість агрегувати сигнали на рівні тематичних контекстів, а не окремих повідомлень [2]. У цьому контексті інтерпретована тематична сегментація постає як методологічна основа для прозорого трасування джерел когнітивного навантаження, сумісна з вимогами масштабованості, приватності й відтворюваності у застосунках штучного інтелекту.

Запропоновано метод виявлення цифрової втоми за повідомленнями через сегментацію спілкування, що поєднує контекстні ембединги, кластеризацію та семантичне узагальнення. Повідомлення кодується SentenceTransformer, після чого за допомогою UMAP і HDBSCAN виокремлюються семантично споріднені групи. Ключові дескриптори сегментів визначаються TF-IDF і KeyBERT, а узагальнення здійснюється генеративною моделлю Flan-T5, формуючи пояснювані тематичні профілі, які відображають потенційне когнітивне навантаження, пов'язане з цифровою втомою.

Методологічно підхід забезпечує повний цикл: формування контекстних представлень, виділення комунікативних сегментів, побудову

їхніх семантичних репрезентацій і візуалізацію структури корпусу з подальшою оцінкою зв'язку сегментів із показниками цифрової втоми. Інтеграція UMAP і HDBSCAN у Python-середовищі гарантує топологічну узгодженість ембедингів та стійкість до шуму, тоді як зв'язка TF-IDF та KeyBERT підсилює інтерпретованість за рахунок поєднання статистичної ваги та контекстної подібності лексем. Узагальнювальні описи сегментів дозволяють проєктувати отримані теми на доменні сценарії моніторингу від освітніх і корпоративних комунікацій до соціальних мереж та слугують основою для побудови індикаторів ризику цифрової втоми на рівні тематичних контекстів, а не окремих повідомлень. Концептуально комунікативний сегмент розглядається як одиниця аналізу, що акумулює поведінкові та емоційні патерни і може бути використана як вхідний об'єкт у системах ШІ-моніторингу добробуту користувачів.

Наукова новизна полягає у поєднанні інтерпретованої тематичної сегментації з кластерною картографією комунікацій для операціоналізації феномену цифрової втоми: від ембедингів до пояснюваних сегментів, що підлягають кількісній оцінці зв'язку з втомою. Практична значущість полягає у створенні відтворюваного програмного конвеєра для систем штучного інтелекту, який забезпечує масштабоване групування, інтерпретацію та візуалізацію сегментів із подальшою інтеграцією у модулі моніторингу станів користувачів. Підхід узгоджується з вимогами до математичного й програмного забезпечення систем ШІ та орієнтований на впровадження в автоматизовані засоби аналітики комунікацій.

Перелік посилань:

1. Віт Р. В., Мазурець О. В. Метод виявлення психологічного цифрового перевантаження за аналізом текстових даних нейромережевими моделями глибокого навчання. Науковий журнал «Вісник Херсонського національного технічного університету». 2025. № 2 (93). Т. 2. С. 107–114. URL: <https://doi.org/10.35546/kntu2078-4481.2025.2.2.12> (дата звернення: 30.10.2025).

2. Віт Р. В., Мазурець О. В. Підхід до тематичної класифікації текстової інформації засобами обробки природної мови. Науковий журнал «Наукові праці Донецького національного технічного університету», серія «Проблеми моделювання та автоматизації проєктування». 2025. № 1 (21). С. 94–99. URL: <https://doi.org/10.31474/2074-7888> (дата звернення: 30.10.2025).

РОЗРОБЛЕННЯ ТА ОПТИМІЗАЦІЯ КРОСПЛАТФОРМНОЇ 2D-ГРИ ЖАНРУ ПАСЬЯНС ІЗ СЮЖЕТНОЮ СКЛАДОВОЮ ТА ДОДАТКОВОЮ МІНІ-ГРОЮ «ПОШУК ВІДМІННОСТЕЙ»

Войцехов М.О. voycehov1@gmail.com, **Ізмайлова М.К.**

Дніпровський національний університет імені Олеся Гончара

Сучасна ігрова індустрія пройшла великий шлях від аркадних та 8-бітних ігор до AAA-проектів. Завдяки високій конкуренції та доступності ігор, індустрія кардинально змінила життя великої кількості людей. Особливе місце займають казуальні ігри, серед яких є карткова гра «Пасьянс», яка залишається однією із найпопулярніших ігор протягом багатьох років завдяки простоті правил, низьким вимогам до характеристик пристрою та коротким ігровим сесіям. Розробка кроссплатформної 2D-гри вимагає комплексного підходу до оптимізації, адаптації інтерфейсу до різних типів екрану та систем управління (клавіатура, миша, сенсорні екрани), а також забезпечення стабільної роботи на різних типах пристроїв із різними апаратними характеристиками – від ПК до телефонів. За допомогою ігрового рушія Unity 6 LTS зручно розробляти кроссплатформені ігри. Рушій підтримує понад 25 платформ серед яких: Windows, MacOS, IOS, Android, PS5, WebGL. Рушій містить велику кількість інструментів, що поліпшують розробку 2D гри.

Для формування вимог до кінцевого продукту було досліджено мобільний та комп'ютерний ринок казуальних ігор з механікою пасьянс. Проаналізовано 4 найпопулярніші конкурентні рішення, визначено їх переваги та недоліки. Після розробки основних механік гри важливою частиною залишається оптимізація та адаптація проєкта під різні пристрої. Оптимізація проєкта відбувається в кілька кроків:

1. Оптимізація коду – один з найпродуктивніших способів покращити кількість кадрів у грі. Використання патернів проєктування допомагає оптимізувати код програми та тримати його в чистоті. Замість вбудованого методу Update() використовувати написані корутини. Ці прості способи допоможуть зберегти ресурси пристрою. Наприклад, завдяки патерну Object Pooling при відтворенні ефектів на слабких пристроях не

відбувається гальмування гри, оскільки об'єкти вже створені та не потрібно використовувати ресурсозатратну команду `Instantiate()`.

2. Оптимізація сцен завдяки видаленню всіх об'єктів на сцені, які не використовуються, щоб вивільнити ресурси.

3. Видалення непотрібних файлів перед збіркою проєкту: UI, анімації, скрипти, які не використовуються. Таким чином збережеться пам'ять.

4. Оптимізація UI. Різниця між ПК та смартфонами дуже велика не тільки в плані технічних характеристик апаратного забезпечення, а й екранів. Тому на мобільних пристроях не потрібно використовувати UI високої якості. За допомогою вбудованого в рушій інструменту можна змінити розмір елементу (рис. 1 ліворуч).

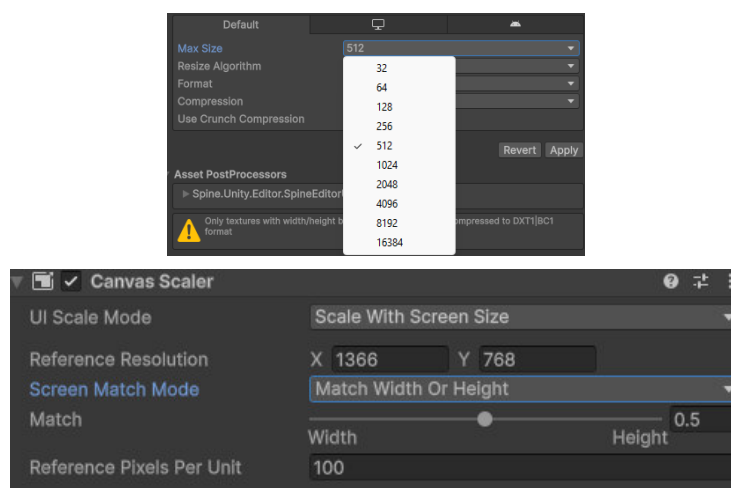


Рисунок 1 – Зміна розширення UI та адаптація гри

Адаптацію під різні пристрої зроблено завдяки вбудованій системі Canvas, що відповідає за відображення UI елементів в грі. В компоненті Canvas Scaler вказано початковий розмір екрану, під який робився проєкт. Screen Match Mode встановлено «Match Width or Height», Match встановлено на «0.5». Таким чином, ми вказуємо, що потрібно розтягувати за висотою та шириною, а параметр 0.5 вказує робити це рівномірно. Коли висота чи ширина буде розтягнуто максимально, то інший параметр не буде розтягуватися до кінця (рис. 1 праворуч).

Розроблений програмний продукт включає в себе основну механіку гри пасьянс, виведення результатів після закінчення гри, серіалізацію та десеріалізацію даних. Також гра адаптована під всі можливі розширення екрану.

ВПЛИВ АНТИВІРУСНИХ ПРОГРАМ НА ПРОДУКТИВНІСТЬ ОПЕРАЦІЙНОЇ СИСТЕМИ

Волк А.В., Саранча В.О., Сірик С.Ф.

volka398@gmail.com, valeriyasarancha@gmail.com, siryk600@gmail.com

Дніпровський національний університет імені Олеся Гончара

Антивірусне програмне забезпечення є ключовим елементом сучасних комп'ютерних систем, що забезпечує захист від широкого спектра шкідливого коду, включно з вірусами, троянськими програмами та шпигунським софтом. Основна функція антивірусів полягає у виявленні, блокуванні та видаленні потенційних загроз. Це дозволяє запобігати пошкодженню системи та втраті важливих даних, що є критично необхідним заходом у сучасному цифровому середовищі, де кількість кіберзагроз постійно збільшується.

Актуальність дослідження впливу антивірусних програм на продуктивність операційної системи зумовлена очевидним конфліктом: попри їхню життєво важливу роль у захисті, вони неминуче створюють додаткове навантаження на апаратні ресурси. Такі процеси, як постійний моніторинг файлової системи, аналіз активних процесів та мережевого трафіку, а також регулярне оновлення баз вірусних сигнатур, споживають значну частину ресурсів процесора та оперативної пам'яті. Це особливо помітно при роботі з великими обсягами даних або на комп'ютерах з обмеженими ресурсами. Таким чином, питання балансу між безпекою та продуктивністю стає центральним.

Позитивним аспектом використання антивірусного ПЗ є забезпечення комплексного захисту. Сучасні антивіруси здатні виявляти та нейтралізувати не тільки відомі, але й нові загрози, ефективно використовуючи поведінковий аналіз поряд із класичними сигнатурними базами. Постійне оновлення дозволяє оперативно реагувати на нові види атак. Крім того, антивірусні програми контролюють безпечність посилань

в інтернеті, перевіряють електронну пошту та файли, що завантажуються. Всі ці функції сприяють збереженню цілісності даних, стабільності роботи системи та безпеки користувача.

Разом із тим, використання антивірусного програмного забезпечення може мати і негативний вплив на продуктивність. Постійна робота у фоновому режимі створює навантаження на центральний процесор і оперативну пам'ять. Це може призводити до помітних затримок у запуску програм, сповільнення часу завантаження операційної системи та зменшення швидкодії під час виконання ресурсомістких завдань. Додатковим фактором є потенційний конфлікт антивірусного ПЗ з іншими програмами, що іноді може викликати збої в роботі.

З метою мінімізації цього впливу існує низка ефективних підходів. Сучасні антивіруси активно використовують оптимізовані алгоритми та технології ресурсоощадження. Окрім цього, користувачам доступне гнучке налаштування: можна задати частоту сканування (наприклад, на нічний час), встановити винятки для перевірених файлів та каталогів, а також вибрати менш ресурсомісткі режими роботи. Регулярне оновлення ПЗ також допомагає зменшити навантаження, оскільки нові версії часто містять оптимізації.

Таким чином, антивірусне програмне забезпечення залишається необхідним засобом захисту комп'ютерних систем. Водночас його вплив на продуктивність вимагає ретельного аналізу та оптимізації. Використання сучасних технологій та правильне налаштування дозволяють досягти раціонального балансу між безпекою та швидкістю, що є ключовим завданням для користувачів.

НЕЙРОМЕРЕЖЕВЕ МОДЕЛЮВАННЯ ДЛЯ ЗАДАЧ РОЗПІЗНАВАННЯ ДОКУМЕНТІВ

Волошин Д.А. imerator.gun@gmail.com, Книш Л. І.

Дніпровський національний університет імені Олеся Гончара

Сучасні інформаційні системи дедалі частіше потребують автоматизованої обробки великої кількості документів, що робить неефективними традиційні алгоритми оптичного розпізнавання тексту OCR через різноманітність форматів, шрифтів і якість вхідних зображень. Тому перспективним напрямом дослідження є застосування нейромережових методів, які забезпечують гнучкість, точність і здатність навчатися на різноманітних даних.

Глибокі згорткові та рекурентні нейронні мережі CNN, LSTM, GRU дозволяють автоматично виділяти ознаки та виконувати послідовне розпізнавання символів. Поширена архітектура CRNN, яка використовується в даному дослідженні, поєднує переваги CNN для просторового аналізу з LSTM для контекстної обробки тексту. Навчання здійснюється за допомогою функції втрат CTC, що забезпечує узгодження послідовностей символів із мітками без потреби точного вирівнювання.

Для підвищення точності в дослідженні були застосовані синтетично згенеровані набори даних і методи аргументації (масштабування, обертання, зміна контрасту та ін.), що покращують узагальнення та стійкість створеної моделі. Інтеграція нейромережової системи у прикладні рішення дозволяє автоматично класифікувати документи, вилучати структуровані дані та формувати результати у форматах JSON або XML.

Отже, запропонований метод побудови нейронної мережі є ефективним інструментом для створення інтелектуальних систем документообігу, здатних суттєво підвищити швидкість, точність та надійність процесів розпізнавання документів.

АВТОМАТИЗАЦІЯ ОЦІНЮВАННЯ ТЕСТІВ РІЗНИХ ТИПІВ В СИСТЕМІ ОФІС 365

Воробійов Я.Ю., yarikomen717@gmail.com

Волошко В. Л., VVL56@i.ua

Дніпровський національний університет імені Олеся Гончара

Сучасний освітній процес широко використовує тестування як метод контролю знань. Серед різновидів тестових завдань часто зустрічаються питання з множинними правильними відповідями та завдання на встановлення відповідностей(логічні пари)[1]. Оцінювання таких завдань вручну є трудомістким і може бути суб'єктивним. Мета роботи – розробити та впровадити програмний інструмент для автоматизованого оцінювання тестів з множинними відповідями і логічними парами у навчанні, що підвищить об'єктивність і ефективність контролю знань в умовах цифровізації освіти.

У рамках роботи розроблено алгоритм і програмний модуль, який автоматично порівнює відповіді студента із заздалегідь підготовленими еталонними відповідями та обчислює підсумковий бал. Для питань із множинними правильними варіантами реалізовано нарахування частинних балів за кожен правильно обрану відповідь, а також передбачено штраф за вибір неправильних опцій. Такий підхід забезпечує справедливе оцінювання навіть у випадку частково правильних відповідей, запобігаючи отриманню незаслужених балів шляхом вгадування. Для завдань на встановлення відповідностей (логічних пар) розроблено функцію перевірки кожної пари: система враховує всі правильно співставленні елементи, нараховуючи бали пропорційно кількості вірних пар. Розроблений програмний модуль може функціонувати як окремий застосунок або бути інтегрований у існуючі системи дистанційного навчання.

Запропонований інструмент було апробовано на серії навчальних тестів, що включали завдання з множинними відповідями та логічними парами. Автоматична система успішно оцінила відповіді студентів; отримані результати збіглися з оцінками викладача при контрольній перевірці, що підтвердило коректність роботи алгоритму. Використання

системи дозволило скоротити час перевірки одного тесту з кількох хвилин до кількох секунд. Крім того, автоматизація усунула людський фактор при виставленні балів та дала змогу надавати студентам результати тестування одразу після завершення тесту.

Серед основних викликів під час реалізації системи було забезпечення універсальності та надійності алгоритмів оцінювання. Необхідно було врахувати різні нестандартні ситуації: коли студент обирає всі можливі варіанти відповідей, не відповідає на питання взагалі або неправильно порівнює всі пари. Для перевірки коректності роботи програми створено низку тестових сценаріїв із граничними випадками. Результати випробувань підтвердили, що алгоритми адекватно обробляють такі ситуації, не нараховують зайвих балів за неправильні дії та вірно обчислюють підсумковий результат.

У перспективі розроблений модуль планується доповнити підтримкою інших типів тестових завдань, в тому числі із застосуванням технологій штучного інтелекту для аналізу текстових відповідей. Також передбачено вдосконалення інтерфейсу та розширення функціоналу системи: зокрема, додавання засобів аналізу поширених помилок студентів і генерування статистичних звітів про результати тестування. Важливим напрямом подальшого розвитку є інтеграція модуля автоматичного оцінювання в існуючі платформи електронного навчання, що дозволить широко застосовувати систему у освітньому процесі.

Висновки. Автоматизація процесу оцінювання тестів із множинними відповідями та логічними парами дозволяє істотно спростити контроль знань і підвищити об'єктивність результатів. Розроблений інструмент підтвердив ефективність, усунув вплив людського фактору та прискорив отримання зворотного зв'язку для студентів. Впровадження таких систем у навчальний процес є важливим кроком до цифрової трансформації освіти, підвищення якості навчання та забезпечення прозорого справедливого оцінювання.

Бібліографічні посилання

1. Толян Ю.В. Методика тестування в системі освіти. Харків: Вид-во ХНПУ ім. Г.С.Сковороди, 2018. – 128 с.

**АНАЛІЗ ЕФЕКТИВНОСТІ МЕТОДІВ-ОБГОРТОК
ПІД ЧАС СТАБІЛЬНОГО ВІДБОРУ ОЗНАК У ЗАДАЧІ
ПРОГНОЗУВАННЯ ОФТАЛЬМОЛОГІЧНИХ ПОКАЗНИКІВ**

Гавриленко М.О., maksym.havrylenko12@gmail.com

Мацуга О.М., olga.matsuga@gmail.com

Дніпровський національний університет імені Олеся Гончара

Якість моделей машинного навчання значною мірою залежить від того, які ознаки включені до моделі. Під час побудови моделей дані зазвичай діляться на навчальну та тестову вибірки, і часто при різних розбиттях на навчальній вибірці відбираються різні ознаки. Це знижує стабільність моделі та ускладнює інтерпретацію результатів, що створює потребу у пошуку стабільної підмножини інформативних ознак.

Метою даної роботи було порівняння методів-обгортки у поєднанні з процедурою стабільного відбору ознак для виділення стійкої підмножини предикторів та підвищення надійності моделі прогнозування офтальмологічних показників.

Аналіз проводився на реальному наборі даних з результатами офтальмологічного обстеження 400 очей 200 пацієнтів, на якому будувалася модель лінійної багатовимірної регресії для передбачення одного з показників. Початкова кількість ознак становила близько 130.

Для відбору ознак було розглянуто чотири методи-обгортки [1]: прямий відбір (forward selection), зворотне вилучення (backward elimination), рекурсивне вилучення (RFE) та покрокова регресія (stepwise regression). Для ідентифікації стійкої підмножини предикторів застосовано процедуру стабільного відбору ознак [2] у поєднанні з цими методами-обгортками.

Процедура стабільного відбору передбачала багаторазове випадкове розділення даних на навчальну та тестову вибірки. На кожній ітерації здійснювався відбір ознак на навчальній вибірці і оцінювалася якість

моделі на тестовій вибірці. Якщо якість перевищувала заданий поріг, результати відбору зберігалися. Після завершення серії ітерацій ознаки, які найчастіше входили до складу високоякісних моделей, формували стабільний набір для фінальної моделі. При цьому метод покрокової регресії автоматично визначав оптимальну кількість ознак на кожній ітерації, тоді як для інших методів-обгортки процедуру виконували для різних цільових розмірів підмножини ознак, щоб знайти оптимальний розмір стабільного набору.

Усі розрахунки проводилися у Google Colaboratory. Для обробки даних, побудови моделей і відбору ознак використовувалися бібліотеки Python: `scikit-learn`, `pandas` і `numpy`. Для покрокової регресії була реалізована власна функція, де додавання й вилучення ознак здійснювалося на основі R^2 .

Результати порівняння ефективності методів-обгортки у поєднанні з процедурою стабільного відбору подано в таблиці 1.

Таблиця 1 – Результати порівняння методів-обгортки

Метод	Значення R^2	Кількість відібрана ознак	Витрачений час
Зворотне вилучення	0.4836	19	10 годин
Прямий відбір	0.4084	11	5 годин
Покроковий відбір	0.405	9	1.5 години
Рекурсивне вилучення	0.3818	13	1 година

Згідно з отриманими результатами, метод зворотного вилучення забезпечив найвищу точність моделі, проте виділив найбільшу кількість ознак і мав найтриваліший час виконання. Метод покрокового відбору виявився найшвидшим, відібравши 9 ознак при незначній втраті якості.

Список використаної літератури

1. Draper N., Smith H. Applied Regression Analysis. Wiley-Interscience, 1998. 736 p.
2. Meinshausen N., Bühlmann P. Stability selection. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*. 2010. Vol. 72. P. 417-473.
<https://doi.org/10.1111/j.1467-9868.2010.00740.x>

РОЗПІЗНАВАННЯ ОБ'ЄКТІВ У ЗОБРАЖЕННЯХ ЗА ДОПОМОГОЮ МОДЕЛІ YOLO

Гайдуков А.С., artem.haidukovv@gmail.com, Степанова Н.І.

Дніпровський національний університет імені Олеся Гончара

Проблема розпізнавання об'єктів у зображеннях є однією з найактуальніших у галузі комп'ютерного зору та штучного інтелекту. Вона привертає значну увагу дослідників і практиків у зв'язку з широким спектром застосувань в автономних транспортних засобах, робототехніці, медицині, економіці, обробці інформації та доповненій реальності. Сучасна наукова спільнота постійно працює над поліпшенням якості розпізнавання, розширенням функціональності моделей і зменшенням їх обчислювальної складності.

Згорткові нейронні мережі є основним інструментом для розпізнавання та класифікації об'єктів на зображеннях і відео. Вони використовують операції згортки, що дозволяє зменшити обсяг інформації, яка використовується в пам'яті, та краще працювати зі зображеннями чи відео з вищою роздільною здатністю.

You Only Look Once (YOLO) – це сучасна система виявлення об'єктів у режимі реального часу. YOLO об'єднує окремі етапи виявлення об'єктів в єдину нейронну мережу. Вона використовує ознаки з усього зображення для прогнозування кожної обмежувальної рамки та прогнозує всі обмежувальні рамки для всіх класів одночасно. Дизайн YOLO забезпечує наскрізне навчання та швидкість роботи в режимі реального часу, зберігаючи при цьому високу середню точність [1]. Основною метою навчання моделі є мінімізація функції втрат для підвищення точності та продуктивності. Функція витрат вказує наскільки помилкова модель в оцінці співвідношення між прогнозованими та фактичними даними.

У даній роботі для розпізнавання об'єктів використовується модель YOLOv8 від Ultralytics [2], що відома своєю простотою використання, швидкістю та точністю. Для задачі виявлення переломів кісток на

рентгенівських знімках та зображеннях МРТ було використано набір даних [3], що містить 1347 зображення у тренувальному наборі та 128 зображення у затверджувальному наборі. Дані поділено на 10 класів та позначено на зображеннях. Для навчання моделі використовувалась претренована модель YOLOv8 Medium, значення внутрішніх параметрів якої допасовувалися впродовж 50 епох. Показники продуктивності моделі зображені на рисунку 1. Значення втрат обмежувальної рамки та класів спадають, значення середньої точності (mAP) зростають, що свідчить про хороше налаштування моделі.

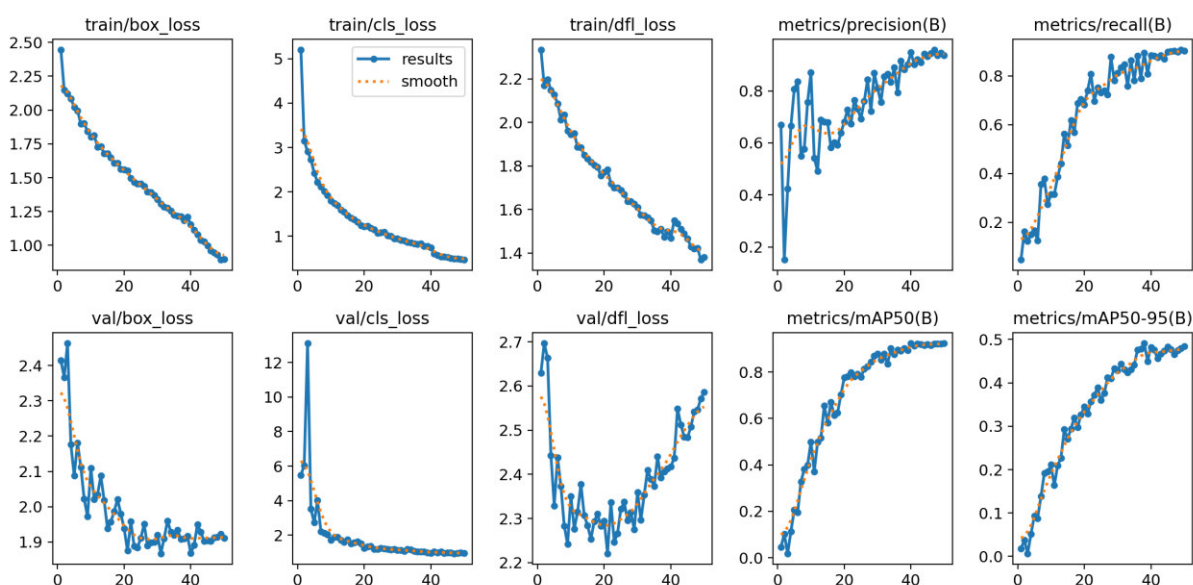


Рисунок 1 – Показники продуктивності натренованої моделі

Отримана модель YOLOv8 показує хороші результати, що демонструє здатність навчити хорошу модель навіть на невеликих наборах даних та за невеликої кількості епох.

Бібліографічні посилання

1. You Only Look Once: Unified, Real-Time Object Detection / J. Redmon та ін. University of Washington. 2016.
2. YOLOv8 by Ultralytics. [Електронний ресурс] – Режим доступу до ресурсу: <https://docs.ultralytics.com/models/yolov8/>
3. Human Bone Fractures [Електронний ресурс] – Режим доступу до ресурсу: <https://www.kaggle.com/datasets/jockeroika/human-bone-fractures-image-dataset>

ШИФРУВАННЯ ТА КОДУВАННЯ ІНФОРМАЦІЇ: МАТЕМАТИЧНІ ОСНОВИ ТА ПРАКТИЧНЕ ЗАСТОСУВАННЯ

Гамбаров М.М., mykyta.gambarov@gmail.com,

Волошко В. Л., VVL56@i.ua

Дніпровський національний університет імені Олеся Гончара

У сучасному інформаційному суспільстві дані є одним з найцінніших ресурсів, а їх захист — ключовим завданням інформатики та кібербезпеки. Зростання обсягів цифрової інформації, розвиток електронних комунікацій і поява нових кіберзагроз обумовлюють актуальність дослідження математичних основ шифрування та кодування інформації.

Метою роботи є комплексне дослідження теоретичних засад і практичних аспектів методів кодування та шифрування інформації, їх математичних основ, застосувань у сучасних інформаційних системах та розробка програмного прикладу на Python.

У роботі розглянуто базові положення теорії інформації, започаткованої Клодом Шенноном, зокрема поняття ентропії, надмірності та кодової відстані. Показано, що ентропія характеризує кількість невизначеності у джерелі повідомлень, надмірність визначає ефективність використання інформаційних каналів, а кодова відстань — здатність кодів виявляти та виправляти помилки. Проаналізовано теорему Шеннона про межу кодування, яка визначає граничну пропускну здатність каналу зв'язку та встановлює межі надійної передачі даних у шумних середовищах.

Досліджено методи кодування інформації, спрямовані на стиснення та забезпечення надійності передачі даних. Розглянуто алгоритми без втрат (Хаффмана, LZW, RLE) та з втратами (JPEG, MP3), а також коди для виявлення і виправлення помилок — коди Хеммінга, CRC, Ріда–Соломона, LDPC, які наближаються до межі Шеннона. Показано принципи побудови

ефективних кодів з урахуванням співвідношення між надмірністю, надійністю та обчислювальною складністю.

Особливу увагу приділено методам шифрування, які забезпечують конфіденційність і цілісність даних. Проаналізовано симетричні (DES, AES) та асиметричні (RSA, ECC) алгоритми, їхні переваги та обмеження. Розглянуто застосування криптографічних хеш-функцій (SHA-256, SHA-3) та електронного цифрового підпису, які забезпечують автентичність і невідомність у цифровому середовищі.

У практичній частині продемонстровано базові принципи шифрування на прикладі реалізації шифру Цезаря мовою Python. Хоча цей метод має історичне значення і не застосовується в реальних системах, він наочно ілюструє сутність криптографічних перетворень.

Результати дослідження узагальнюють сучасні підходи до захисту інформації, демонструючи взаємозв'язок між математичною теорією та прикладною реалізацією криптографічних алгоритмів. Зроблено висновок, що ефективне кодування забезпечує надійність передачі даних, а шифрування — їх конфіденційність і безпеку.

Висновки. Шифрування та кодування, як ключові напрями інформаційної безпеки, ґрунтуються на строгих математичних принципах і забезпечують основу для створення надійних цифрових систем. Їх розвиток залишається одним із пріоритетів сучасної кібербезпеки та інформаційних технологій.

ПРО АПРОКСИМАЦІЙНИЙ ПІДХІД ДО РОЗВ'ЯЗАННЯ ЗАДАЧІ ОПТИМАЛЬНОГО КЕРУВАННЯ РОЗПОДІЛОМ ТЕПЛА У НАПІВОБМЕЖЕНОМУ СТРИЖНІ

Гарт Л.Л., ll_hart@ukr.net

Бугаєнко А.О., buhaienko.ann@gmail.com

Дніпровський національний університет імені Олеся Гончара

Задачі оптимального керування тепловими процесами посідають важливе місце серед прикладних інженерних та промислових задач [1]. Серед них особливу увагу приділяють задачам, що формуються в напівобмежених областях, оскільки вони є більш складними з точки зору системних досліджень. Одним з головних викликів у розв'язанні таких задач є вибір штучної межі, яка дозволяє замінити напівобмежену область обмеженою, а також побудова апроксимаційних алгоритмів, які б забезпечували достатню точність отримуваних наближених розв'язків [2].

У цій роботі розглядається задача оптимального керування процесом термостатики, який відбувається в однорідному напівобмеженому стрижні ($x \in [0, +\infty)$). Необхідно, керуючи температурою внутрішніх джерел тепла, привести температурний режим $u(x)$ у стрижні якомога ближче до заданого режиму $\varphi(x)$. Математичне формування цієї задачі: мінімізувати функціонал

$$J(v) = \int_0^{+\infty} (u(x; v) - \varphi(x))^2 dx, \quad (1)$$

за умов

$$Lu + d(x)u = f(x) + v(x), \quad x \in (0, +\infty); \quad (2)$$

$$u(0) = g_0, \quad u(+\infty) = 0, \quad (3)$$

$$D_0 \leq d(x) \leq D_1, \quad x \in (0, +\infty), \quad (4)$$

де g_0 – задана стала; $\varphi(x)$ – задана функція бажаного розподілу температури у стрижні; $d(x)$ – коефіцієнт теплових втрат; $f(x)$ – задана функція, що описує дію внутрішніх джерел тепла у стрижні; $v \equiv v(x)$ –

шукана функція керування внутрішніми джерелами тепла, $u \equiv u(x) \equiv u(x; v)$ – шукана функція стану (температури), що знаходиться під впливом керування $v(x)$ на проміжку $[0, +\infty)$. Зазначимо, що задача розглядається за відсутності обмежень на керування: $v(x) \in L_2[0, +\infty)$.

Дія диференціального оператора L у рівнянні (2) визначається формулою

$$Lu = -\frac{d}{dx} \left(k(x) \cdot \frac{du}{dx} \right) = -(k(x) \cdot u'(x))', \quad (5)$$

де $k(x)$ – коефіцієнт теплопровідності у точках стрижня, за умови

$$0 < K_0 \leq k(x) \leq K_1, \quad x \in (0, +\infty). \quad (6)$$

Для розв'язання задачі оптимального керування (1)–(6) пропонується застосувати апроксимаційний підхід на основі методу усікання, загальна схема якого стосовно нескінченної системи рівнянь описана у роботах [2, 3] і який дозволяє трансформувати задачу (1)–(6) до відповідної «наближеної» задачі керування, розглядуваної на деякому замкненому скінченному проміжку $[0, l] \subset [0, +\infty)$. Для розв'язання останньої передбачається використати прямий метод [4], оснований на її сіткової апроксимації із подальшим застосуванням методу прогонки для розв'язання різницевого аналогу крайової задачі (2)–(6) та градієнтного методу мінімізації цільового функціонала. Для знаходження значення параметра l , який забезпечив би отримання наближеного розв'язку вихідної задачі (1)–(6) із бажаною точністю пропонується використовувати проєкційно-ітераційний принцип [5, 6], що базується на ітераційному уточненні апроксимації.

Передбачається розробити, програмно реалізувати та дослідити на практичну збіжність наближені алгоритми розв'язання задачі оптимального керування (1)–(6) у напівобмеженій області, оснований як на звичайному методі усікання, так і на проєкційно-ітераційному принципі уточнення апроксимації.

Бібліографічні посилання

1. Zui-Cha Deng, Xin-Rui Yang, Liu Yang. Numerical reconstruction of heat source on a semi-infinite rod. *Journal of Thermophysics and Heat Thansfer*. 2024. Vol. 38 (2): 1 – 11. doi: <https://doi.org/10.2514/1.T6905>.
2. Dang Quang A., Tran Dinh Hung. Method of infinite system of equations for problems in unbounded domains. *Journal of Applied Mathematics*. 2012. Vol. 2012: 1 – 17. doi: 10.1155/2012/584704.
3. Kantorovich L.V., Akilov G.P. Functional analysis, 2nd edition. Pergamon: Elsevier, 1982. 604 pp. <https://doi.org/10.1016/C2013-0-03044-7>.
4. Hart L.L., Buhaienko A.O. On the modification of direct methods for solving optimal control problems of stationary thermal processes. *Challenges and Issues of Modern Science*. 2025. Vol. 4 No.1. doi: <https://doi.org/10.15421/cims.4.314>.
5. Balashova S.D. On solving minimization problems using projection-iterative methods [In Russian]. *Mathematical Models and Computational Methods in Applied Problems*. 1996. C. 99–104. <https://e.surl.li/kdqhsa>.
6. Hart L.L. An explicit projection-iteration method for solving ill-posed operator equations [In Russian]. *Problems of Applied Mathematics and Mathematical Modeling*. 2015. Vol. 15. C. 33–47.

АНАЛІЗ ДИНАМІЧНИХ МОДЕЛЕЙ ОПТИМАЛЬНОГО ІНВЕСТИВАННЯ

Гарт Л.Л., ll_hart@ukr.net; Гарт А.В., av.hart.237@gmail.com;

Петров І.С., petrov.is21@365.dnu.edu.ua

Дніпровський національний університет імені Олеся Гончара (www.dnu.dp.ua)

У сучасному світі прийняття оптимальних управлінських рішень у різних сферах економіки є ключовим фактором ефективного використання ресурсів. Задачі оптимального керування широко застосовуються в моделюванні економічних процесів, таких як оптимізація інвестиційних стратегій, управління активами, фінансове прогнозування тощо.

У роботі розглянута математична модель [1]: мінімізувати функціонал

$$J(u) = \int_{t_0}^T (x(t) - d(t))^2 dt \quad (1)$$

за умов

$$\dot{x}(t) = F(x(t); u(t); r(t); t), \quad t_0 \leq t \leq T; \quad (2)$$

$$x(t_0) = x_0, \quad x(T) = x_T, \quad (3)$$

$$u \equiv u(t) \in U = \{u(t) \in L_2^{(n)}[t_0, T]: u_{\min} \leq u_i(t) \leq u_{\max}, t_0 \leq t \leq T; i = \overline{1, n}\}, \quad (4)$$

де $x(t)$ – шукана вартість інвестиційного портфелю (функція стану); $u(t) = (u_1(t), u_2(t), \dots, u_n(t))$ – шукані частки акцій у інвестиційному портфелі (функція керування); $r(t) = (r_1(t), r_2(t), \dots, r_n(t))$ – вартості акцій n видів у момент часу $t \in [t_0; T]$, які задовольняють умови

$$\dot{r}_i(t) = f_i(r_i(t); t; \alpha_i), \quad t_0 \leq t \leq T; \quad i = \overline{1, n}; \quad (5)$$

$$r_i(t_0) = r_i^{(0)}, \quad i = \overline{1, n}. \quad (6)$$

Тут $d(t)$ – бажаний рівень прибутковості інвестиційного портфеля (цільова траєкторія вартості); $F(x(t); u(t); r(t); t)$ та $f_i(r_i(t); t; \alpha_i)$, $i = \overline{1, n}$, – задані функції, що описують динаміку вартості інвестиційного портфеля та кожної окремої акції в ньому, відповідно; x_0 , x_T – задані початкова та кінцева вартості інвестиційного портфеля; $r_i^{(0)}$ та α_i ($i = \overline{1, n}$) – відповідно,

задані початкові вартості та числові параметри моделі для кожної акції у інвестиційному портфелі. Параметри α_i у фінансових моделях вважають додатними (модель зростання активів), від’ємними (як демпфуючий ефект – наприклад, втрати через податки або ринок) або довільними (з урахуванням властивостей ринку або емпіричних оцінок). У контексті задачі оптимального інвестування, множина $U \subseteq L_2^{(n)}[t_0, T]$ описує допустимі керування, тобто можливі стратегії розподілу активів у портфелі на заданому часовому проміжку $[t_0; T]$; $u_{\min}$, $u_{\max}$ – задані нижня та верхня фінансові межі обсягів вкладень за кожною часткою активу.

Під час виконання досліджень було проведено огляд динамічних моделей прикладних задач інвестування в цінні папери та здійснено формалізацію математичної постановки задачі оптимального керування інвестиційним портфелем. На основі необхідних умов оптимальності була побудована крайова задача принципу максимуму та проаналізована її розв’язуваність. Для числової реалізації прямого ітераційного методу розв’язання, задачу (1)–(6) було попередньо зведено до задачі керування з вільним правим кінцем за допомогою методу штрафу [2]. Розроблено програму в середовищі MATLAB для автоматизованого обчислення оптимального керування та функції стану, що дозволило дослідити вплив вхідних параметрів моделі на критерій якості. Проведено серію тестувань, які демонструють ефективність обраного підходу.

Результати роботи можуть бути застосовані у подальших дослідженнях для аналізу та прогнозування в економічних задачах розподілу інвестиційних ресурсів, а також під час розв’язання прикладних задач оптимального керування системами із зосередженими параметрами.

Бібліографічні посилання

1. Кулян В.Р., Юнькова О.О. Методи оптимального керування в задачах диверсифікації портфеля інвестицій. *Вісник Київського національного університету імені Тараса Шевченка*. 2015, 1(15), 18–22.
2. Hart L.L., Ruzhevych V.O., Hart A.V. Numerical analysis of controlled motion of a quadcopter drone. *Problems of Applied Mathematics and Mathematical Modeling*. 2024, 24, 38–56. doi: 10.15421/322405.

АВТОМАТИЗАЦІЯ ПРОЦЕСІВ ПРИЙНЯТТЯ РІШЕНЬ У БАНКІВСЬКІЙ СФЕРІ: МОЖЛИВОСТІ, РИЗИКИ ТА ПЕРСПЕКТИВИ РОЗВИТКУ

Гладуш А.С., hladush.as@365.dnu.edu.ua

Дніпровський національний університет імені Олеся Гончара

Банківська сфера зараз стикається з необхідністю оперативного прийняття управлінських рішень на основі великих обсягів даних. Традиційні підходи до аналізу інформації вже не забезпечують належної швидкості, точності та ефективності, що зумовлює потребу у впровадженні систем автоматизованого прийняття рішень на базі технологій штучного інтелекту та інтелектуальної аналітики даних. Інтенсивне впровадження штучного інтелекту супроводжується низкою ризиків – від загроз кібербезпеці та витоків персональних даних до проблем алгоритмічної упередженості, відсутності прозорості у прийнятті рішень і недостатнього нормативного регулювання.

Тому дослідження сучасних підходів до автоматизації прийняття рішень у банківській сфері, визначення їхніх можливостей, ризиків та перспектив розвитку є надзвичайно актуальним завданням. Воно має як практичне значення – для підвищення ефективності фінансових установ, так і наукове – для формування нових концепцій безпечного та етичного використання штучного інтелекту у фінансових технологіях.

Інтелектуальний аналіз даних (Data Mining) та технології Big Data стають ключовими інструментами для автоматизації процесів прийняття рішень у банківській діяльності. Застосування алгоритмів машинного навчання, нейронних мереж та моделей прогнозової аналітики дозволяє здійснювати автоматичне виявлення закономірностей у фінансових операціях, прогнозувати ризики, оптимізувати кредитну політику та персоналізувати банківські послуги.

Особливу увагу привертають системи підтримки прийняття рішень (DSS) нового покоління, що інтегрують інтелектуальні аналітичні модулі з корпоративними базами даних. Такі системи забезпечують адаптивність до змін ринку, підвищують рівень безпеки фінансових операцій і сприяють зниженню людського фактору.

Разом з тим, широке впровадження штучного інтелекту у банківській сфері супроводжується низкою ризиків і викликів:

- Алгоритмічна упередженість (bias) може призводити до несправедливих рішень щодо клієнтів, якщо моделі навчаються на нерепрезентативних або некоректних даних.
- Кібербезпека та конфіденційність даних стають критично важливими, адже обробка великих обсягів персональної інформації створює ризики несанкціонованого доступу або витоків.
- Непрозорість алгоритмів («чорна скринька» штучного інтелекту) ускладнює аудит прийнятих рішень і може викликати сумніви у клієнтів чи регуляторів.
- Етичні аспекти застосування штучного інтелекту, зокрема питання автоматизованої відмови у кредитуванні, також потребують нормативного регулювання та підвищення довіри суспільства до таких систем.

Отже, інтелектуальна автоматизація банківських процесів відкриває значні можливості для підвищення ефективності та конкурентоспроможності фінансових установ. Водночас необхідне формування системи контролю, етичних стандартів і механізмів перевірки рішень ШІ, щоб мінімізувати ризики та забезпечити надійність і прозорість використання технологій у банківській сфері.

ENCODER-ONLY АРХІТЕКТУРА TRANSFORMER У ЗАДАЧАХ КЛАСИФІКАЦІЇ ТЕКСТУ

Глушко О.С., hlushko_o@365.dnu.edu.ua,

Золотько К.Є., zolotko_k@365.dnu.edu.ua

Дніпровський національний університет імені Олеся Гончара

Стрімкий розвиток технологій обробки природної мови призвів до появи архітектур, які суттєво перевищують ефективність класичних моделей машинного навчання. Однією з ключових інновацій у цій галузі є архітектура Transformer, що базується на механізмі самоуваги (self-attention) [1]. Encoder-only підхід у трансформерах став основою для створення сучасних моделей обробки тексту, зокрема BERT, RoBERTa, ALBERT та їхніх похідних [2].

Encoder-only архітектури використовуються переважно для задач класифікації, аналізу тональності, визначення інтенцій користувача та інших завдань, де необхідне розуміння контексту без генерації нового тексту. На відміну від повних Transformer-моделей, що містять encoder і decoder, encoder-only підхід оптимізований для вилучення семантичних зв'язків між токенами та створення контекстуальних представлень тексту [3]. Це забезпечує високу ефективність при помірних обчислювальних витратах.

Математично модель ґрунтується на багатоголовому механізмі уваги (Multi-Head Self-Attention), який дозволяє враховувати залежності між усіма елементами послідовності. Були розглянуті залишкові зв'язки та види нормалізації шарів, які забезпечують стабільність навчання. Feed-forward блоки з обраною функцією активації GELU для сприяння формування глибших семантичних представлень [4, 5]. Також було розглянуто використання токена [CLS] як агрегатора контекстної інформації, який дозволяє реалізувати ефективний класифікаційний шар із функцією втрат Binary Cross-Entropy у випадку мультиміткових задач [6].

Encoder-only архітектури мають низку переваг: швидкість навчання, стабільність збіжності, можливість попереднього тренування на великих текстових корпусах і подальшого донавчання для конкретних задач. Водночас вони залишаються вимогливими до ресурсів, що стимулює розвиток легших варіантів – DistilBERT, TinyBERT та інших моделей, оптимізованих для використання у прикладних системах реального часу.

Таким чином, encoder-only підхід є одним з ключових напрямів розвитку сучасного NLP, який поєднує високу точність, масштабованість та гнучкість при вирішенні задач класифікації тексту. Подальші дослідження спрямовані на оптимізацію архітектури під задачу автоматичного тегування тексту, та динамічний вибір кількості вихідних даних в залежності від контекстуального змісту вхідного тексту.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Attention Is All You Need / Ashish Vaswani та ін. 31st Conference on Neural Information Processing Systems (NIPS 2017). 2017. URL: <https://arxiv.org/pdf/1706.03762>
2. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding / J. Devlin та ін. Proceedings of the 2019 Conference of the North. Minneapolis, Minnesota. 2019. <https://doi.org/10.18653/v1/n19-1423>
3. Unsupervised Cross-lingual Representation Learning at Scale / A. Conneau та ін. Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics, м. Online. Stroudsburg, PA, USA, 2020. URL: <https://doi.org/10.18653/v1/2020.acl-main.747>
4. Lee M. Mathematical Analysis and Performance Evaluation of the GELU Activation Function in Deep Learning. Journal of Mathematics. 2023. Т. 2023. С. 1–13. URL: <https://doi.org/10.1155/2023/4229924>
5. An Analysis of State-of-the-art Activation Functions For Supervised Deep Neural Network / A. Nguyen та ін. 2021 International Conference on System Science and Engineering (ICSSE), м. Ho Chi Minh City, Vietnam, 26–28 серп. 2021 р. 2021. URL: <https://doi.org/10.1109/icsse52999.2021.9538437>
6. GeeksforGeeks. Binary Cross Entropy/Log Loss for Binary Classification - GeeksforGeeks. URL: <https://www.geeksforgeeks.org/deep-learning/binary-cross-entropy-log-loss-for-binary-classification>

ВЕРИФІКАЦІЯ ТА ВІЗУАЛІЗАЦІЯ КОНТАКТНИХ РОЗРАХУНКІВ ІЗ ВИКОРИСТАННЯМ *ANSYS*, *MATLAB* ТА *IN-HOUSE* ІНСТРУМЕНТІВ ПОСТПРОЦЕСИНГУ ДАНИХ

Гончаров Я.А., Goncharov.yar.an@gmail.com

Дніпровський національний університет імені Олеся Гончара

Контактна взаємодія абсолютно жорсткого штампа з пружним півпростором належить до класичних, проте складних задач. Пропонується підхід до розв’язання таких задач у якому поєднується скінченно–елементне моделювання з глибоким постпроцесінгом. Перше забезпечило математичну та фізичну постановку задачі, друге – виконало аналіз даних та фільтрацію шумів, що сприяло покращенню верифікації та візуалізації контактних напружень і переміщень, а також допомогло введенню чітких інженерних критеріїв для оцінки і подальшої оптимізації конструкцій.

Для групи задач у середовищі Ansys було побудовано скінченно–елементні моделі контактної взаємодії абсолютно жорстких штампів різної форми з пружним півпростором. Отримано розподіли напружень і переміщень та виконано експорт результатів у вигляді таблиць для подальшої обробки. Було створено MATLAB-інструментарій для завантаження з Ansys, аналізу й візуалізації: автоматичне виявлення ділянок заданого контактного тиску, побудова ізоліній і об’ємних моделей, синхронне відображення напружень/переміщень і пакетна обробка кількох сценаріїв. In-house інструмент розроблений за допомогою MATLAB сформував ланцюг постпроцесінгу даних із Ansys: від прямого імпорту табличних результатів і k-NN-оцінювання кроку вибірки до побудови регулярної сітки та інтерполяції. Розроблені модулі забезпечили коректне виділення зон заданого контактного тиску. Стандартизована візуалізація напружень і переміщень полегшила інтерпретацію результатів. Розроблений програмний модуль може бути використаний для розв’язання аналогічних контактних задач.

ПЛАТФОРМА ДЛЯ ТОРГІВ ЕЛЕКТРОЕНЕРГІЄЮ

Горбачук В.М., Ніколенко Д.І., Осипенко Д.О. VGorbachuk@nas.gov.ua

Інститут кібернетики імені В.М.Глушкова НАН України

Існуючі європейські внутрішньодобові ринки (ВДР) поступово об'єднуються через пан'європейське рішення, відоме як Єдине внутрішньодобове поєднання (Single Intraday Coupling (SIDC)). Модульне інтегроване ІТ-рішення M7 XBID як пан'європейська платформа для SIDC уможливорює безперервний внутрішньодобовий трейдинг (через паросполучення (matching)) за енергетичними контрактами на фізичну поставку електроенергії 24 години на добу, 365 днів на рік по всій Європі: M7 XBID – це сервіс у загальноєвропейському ВДР для електроенергії транскордонних мереж. Докладніше зупинимося на технічних модулях M7 XBID, а також на потоках даних, включаючи висхідні (upstream) та низхідні downstream системи. Платформа M7 XBID перебуває під постійним моніторингом і вдосконалюється з моменту свого створення, щоб зберігати свою здатність обробляти зростаючі обсяги даних. Кількість транзакцій заявок (orders) зростала від близько 0,5 млн./доба з моменту її запуску в червні 2018 р до приблизно 2 млн./доба у січні 2021 р. (2-га хвиля впровадження (2-nd wave Go-live)), 4 млн./доба у січні 2023 р. (3-rd wave Go-live), 8 млн./доба (4-th wave Go-live) у липні 2023 р., 12 млн./доба (5-th wave Go-live) у липні 2024 р.; число транзакцій угод (trades) було майже вдвічі меншим числа заявок. 2-nd wave Go-live стосується 2-гої фази трансформації загальної системи планування ресурсів підприємства (Enterprise Resource Planning (ERP)), яка відбувається після запуску системи за початкових умов.

На відміну від 1-st wave Go-live, яка є початковим впровадженням і запуском (включає в себе введення системи в експлуатацію (getting the system live) та навчання користувачів її основним функціям), 2-nd wave Go-live зосереджується на оптимізації процесів системи, використанні нових даних та їх точності для досягнення кращих вимірюваних ділових результатів, таких як поліпшення управління запасами, підвищення ефективності

ланцюгів постачання, досягнення всіх можливих конкурентних переваг. 2-га фаза включає дії, що вживаються для переходу від базової функціональності ERP до глибшого рівня оптимізації та розширення спроможностей прийняття стратегічних рішень. IT-рішення M7 XVID базується на Спільній книзі заявок (Shared Order Book (SOB)), Модулі менеджменту потужностей (Capacity Management Module (CMM)), Інтеграції моделі зрілості можливостей (Capability Maturity Model Integration (CMMI)), Модулі доставки (Shipping Module (SM)), а також на Центральному модулі довідкових даних (Central Reference Data Module).

Модуль трейдингу (SOB) формують такі блоки: паросполучення (безперервне виконання заявок (паросполучення заявок (ордерів) на купівлю та продаж)), інтерфейс публічних повідомлень (інтерфейс до локальних рішень трейдингу), книга заявок (розрахунки локального перегляду книги заявок), маршрутизація (розрахунки транспортування між зонами доставки; розрахунки матриці хаб-хаб (Hub-to-Hub)), звітність (створення, генерування та документування діяльності стосовно трейдингу і заявок, а також неявне розміщення спроможностей у вигляді XML-звітів).

Модуль CMM/ CMMI формують такі блоки: розміщення спроможностей (менеджмент наявної пропускної здатності передачі; неявне та явне розміщення спроможностей на рівні кордону), інтерфейс TSO (спільний інтерфейс для комунікації з TSOs), інтерфейс явного учасника ринку (Спільний інтерфейс для комунікації з явними сторонами ринку). Ця продуктивна модульна архітектура дає більшу гнучкість, масштабованість і зручність обслуговування, а також адаптивність до змінних вимог [1].

Список використаних джерел

1. Горбачук В., Гавриленко С. Вплив ціноутворення хмарних сервісів на прибуток провайдера, споживчий надлишок і суспільний добробут. *Проблеми програмування*. 2020. № 2–3. С. 237–245.

СИНЕРГЕТИЧНЕ КЕРУВАННЯ АВТОНОМНИМ МОБІЛЬНИМ РОБОТОМ

Грегорович М.Ю., nikitagregorovich@gmail.com

Зайцев В.Г., vadymzaytsev65@gmail.com

Дніпровський національний університет імені Олеся Гончара

Стрімкий розвиток сенсорики, енергоефективних обчислювальних платформ і вбудованих систем створив підґрунтя для широкого впровадження автономних мобільних роботів (AMP) у логістиці, сільському господарстві, медицині, безпеці та рятувальних операціях. Збільшення складності середовищ і вимог до надійності висуває жорсткі умови до систем керування: потрібні стійкість, робастність, енергоефективність і передбачуваність за неповної/за-шумленої інформації та змінних зовнішніх впливів [1; 4; 10].

Еволюція методів керування пройшла шлях від класичних ПД-регуляторів до оптимальних (LQR), робастних (зокрема H_∞), адаптивних та інтелектуальних підходів (нечітка логіка, нейромережі), а також прогнозного керування модельного типу (MPC), яке дозволяє прямо враховувати обмеження, оптимізувати енерговитрати та компенсувати збурення в горизонті прогнозу [1–2; 5]. Водночас MPC потребує значних обчислювальних ресурсів і якісних моделей. На цьому тлі синергетичний підхід вирізняється поєднанням аналітичної керованості, робастності та плавності сигналів. Його сутність — у введенні макрозмінних та стабілізації інваріантів, що формують «синергетичний атрактор» і забезпечують збіжність траєкторій незалежно від початкових умов, зі знизеними піками навантажень на приводи [2; 11].

У роботі узагальнено практичні аспекти побудови архітектур керування AMP на базі ROS та вбудованих платформ (Raspberry Pi, Linux), використання цифрових двійників і стендових макетів для верифікації алгоритмів, а також інтеграції планування шляху (A^* , D^* , RRT) з локальним уникненням перешкод (потенціальні поля, VFH) і

низькорівневими регуляторами [4; 8; 10]. Показано доцільність гібридизації MPC із синергетичними законами для суміщення планування з глобальною стабілізацією через макрозмінні, що підвищує безпеку, точність і енергоефективність у реальному часі [1–2; 5; 8].

Як приклад наведено синтез синергетичного керування для шестивимірної моделі вертольота: цілі визначаються як стабілізація орієнтації та висоти, вводяться макрозмінні для кутів та вертикального руху, здійснюється динамічна декомпозиція і побудова внутрішніх законів керування з подальшим узгодженням на перетинах інваріантних різноманітностей. Підхід забезпечує асимптотичну стійкість і робастність до збурень; термінальні (кінцевочасові) варіанти прискорюють перехід у безпечні режими без надмірних осциляцій [2; 11].

Зроблено висновок про перспективність синергетичного підходу для ресурсно обмежених платформ і його високу сумісність з практичною імплементацією; рекомендовано поступову валідацію (симуляція → стенд → полігон), ітеративне налаштування параметрів та інтеграцію адаптивних механізмів для компенсації невизначеностей моделі та середовища [1–2; 8; 10].

У роботі змодельовано рух АМР у двовимірному середовищі. Спираючись на синергетичний підхід, виведено замкнені формули сигналів керування для трьох базових траєкторій — прямолінійної, кругової та синусоїдальної [3]. Для кожного режиму синтезовано керуюче прискорення, що забезпечує стійке стеження та нечутливість до зовнішніх збурень. Отримані результати дають змогу зіставити ефективність розглянутих законів керування та підібрати оптимальні параметри макрозмінних під різні сценарії руху. Синергетичний підхід поєднує теоретичну строгість і практичну дієвість, слугуючи універсальним інструментом для побудови адаптивних алгоритмів керування, придатних до реальної експлуатації АМР. Завдяки здатності формувати гладке, надійне й робастне керування в умовах невизначеності метод відкриває

перспективи для автономних наземних систем. Імітаційні експерименти в MATLAB підтверджують працездатність підходу та дають практичні рекомендації щодо проєктування синергетичних систем керування АМР [2; 3; 9].

Бібліографічні посилання:

1. Глазунов В. А. *Синергетика і управління складними системами*. Київ: Наукова думка, 2010.
2. Колесников А. А. *Синергетическая теория управления*. Москва: Энергоатомиздат, 1994.
3. Нечипоренко В. А. *Методи керування БПЛА із застосуванням комп'ютерного зору*. Технічні науки, №5, 2024.
4. Pavlov S., Novoselov S. *Керування мобільним автономним роботом з використанням ROS*. Universum: Технічні науки, 2023/2024.
5. Mellinger D., Kumar V. *Minimum Snap Trajectory Generation and Control for Quadrotors*. IEEE ICRA, 2011.
6. Грозмані Я. О. *Метод керування автономним мобільним роботом: кваліфікаційна робота магістра*. ХНУ, 2023.
7. Стеценко К. К. *Розробка системи керування автономного мобільного робота: пояснювальна записка (бакалавр)*. ХНУРЕ, 2024.
8. Поліщук М. М. *Автоматизований синтез мобільних роботів довільної орієнтації...* Автореф. докт. дис., КПІ, 2021.
9. Єфімік Є. О. *Розроблення АСК мобільним роботом з комбінованими рушіями: пояснювальна записка (магістр)*. ХНУРЕ, 2025.
10. Батюк В., Стрембицький М. *Побудова алгоритму керування автономним мобільним роботом*. ТНТУ, 2019.
11. Дзейкало А. А. *Синергетичний синтез системи управління просторовим рухом вертольота*. Матеріали FААCTIS-2015. Суми: СумДУ, 2015.

ЗАДАЧА СИЛЬВЕСТРА ТА ЇЇ УЗАГАЛЬНЕННЯ

Давидов О.О.¹, Давидов О.С.²*Davydov0508@gmail.com*¹*Київський національний університет імені Тараса Шевченка*²*Національний авіаційний університет, Київ*

Задача про найменшу обмежувальну гіперсферу (задача Сильвестра) полягає у знаходженні кулі мінімального радіуса, що містить даний скінченний набір точок у n -вимірному евклідовому просторі. Її узагальнення вимагає побудови найменшої кулі, яка охоплює скінченний набір куль із заданими центрами й радіусами. Такі задачі виникають при розміщенні об'єктів (наприклад, станцій швидкої допомоги або супутників), коли потрібно мінімізувати максимальну відстань до всіх елементів обслуговування. Формально ці задачі зводяться до мінімаксних задач опуклого програмування.

Метою роботи є аналіз обчислювальних можливостей алгоритмів **sylvester1** та **sylvester2**, які побудовані для розв'язання задачі Сильвестра та її узагальнення на основі методу еліпсоїдів [1, 2]. Обидва алгоритми генерують послідовність еліпсоїдів, що містять точку мінімуму відповідної опуклої мінімаксної функції, об'єми яких монотонно зменшуються. Кожна ітерація потребує обчислення субградієнта відповідної функції і дає новий еліпсоїд, центр якого використовується як чергове наближення.

Алгоритм **sylvester1** призначений для розв'язання задачі Сильвестра та полягає в безумовній мінімізації опуклої кусково-квадратичної функції, що характеризує квадрат радіуса кулі. Радіус кулі задається як максимальна відстань від її центра до заданого набору точок. Алгоритм завершує свою роботу, як тільки відхилення квадрату радіуса від мінімального стає меншим за наперед задане значення ε_1 .

Алгоритм **sylvester2** призначений для розв'язання узагальненої задачі Сильвестра та полягає в безумовній мінімізації опуклої

кусково-гладкої функції, що характеризує максимальну із сум відстаней від центра кулі до заданого набору точок та розміщених в них радіусів куль. Алгоритм завершує свою роботу, як тільки відхилення значення функції від мінімального стає меншим за наперед задане значення ε_2 .

Обчислювальні експерименти проводились для обох алгоритмів та було з'ясовано, що зі зменшенням ε_1 та ε_2 збільшення кількості ітерацій відповідає теоретичній оцінці швидкості збіжності методу еліпсоїдів. Для алгоритму **sylvester2** експерименти проводились як для куль однакового радіуса, так і куль нульового радіуса, що відповідає задачі Сильвестра. Результати підтвердили, що **sylvester2** не поступається **sylvester1** за швидкістю збіжності.

Якщо $n = 10\text{--}20$, то алгоритми **sylvester1** та **sylvester2** можна успішно використовувати для знаходження центру та радіуса кулі мінімального радіуса. Вони демонструють високу швидкість та точність на сучасних персональних комп'ютерах, де час розв'язання задач при $n = 20$ складає декілька секунд. Це підтверджують обчислювальні експерименти, які проводились в середовищі GNU Octave 6.2.0 на персональному комп'ютері з процесором AMD Ryzen 5 4500U, 16 ГБ ОЗП під керуванням Windows 10.

Список літератури

1. Стецюк П.І., Хом'як О.М., Давидов О.О. Використання методу еліпсоїдів для задачі Сильвестра та її узагальнення. Кібернетика та комп'ютерні технології. 2024. С. 27-47.
2. Хом'як О.М., Давидов О.О. Метод еліпсоїдів для гіперкулі мінімального радіуса. Матеріали 7-ї Міжнародної наукової конференції "Математичне моделювання, оптимізація та інформаційні технології (ММОТІ-2021)", 15–19 листопада 2021 р. С. 340–342.

МУЛЬТИАГЕНТНІ ПІДХОДИ ДО АВТОМАТИЗОВАНОГО ТЕСТУВАННЯ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ

Деменко І.О., ivandemenkofifa@gmail.com

Інститут кібернетики імені В.М. Глушкова НАН України

Сучасні програмні системи відзначаються високою складністю, розподіленою архітектурою та використанням технологій штучного інтелекту, що ускладнює процес їх тестування. Традиційні автоматизовані методи з фіксованими сценаріями не забезпечують достатньої адаптивності й повноти перевірки. Перспективним напрямом є застосування мультиагентних систем, у яких тестування реалізується як колективна взаємодія автономних агентів, що генерують, виконують і оцінюють тестові сценарії у розподіленому середовищі.

Основна проблема полягає у відсутності динамічного механізму узгодження дій між агентами, здатного забезпечити цілісність і стабільність процесу перевірки складного програмного забезпечення. Пропонується модель мультиагентного тестування, де взаємодія агентів Generator, Validator, Evaluator та Coordinator забезпечує повний цикл створення, виконання й оцінювання тестів у реальному часі.

У запропонованому підході тестування здійснюється через взаємодію автономних агентів: Generator-Agent формує сценарії, Validator-Agent виконує перевірки, Evaluator-Agent обчислює метрики якості, а Coordinator-Agent координує їхню роботу.

Процес відбувається у замкненому циклі, де агенти адаптують свої політики на основі зворотного зв'язку, що забезпечує самонавчання та гнучкість тестування.

Функціонування системи можна описати у вигляді стохастичної динамічної моделі:

$$S(t + 1) = F(S(t), A(t), E(t)),$$

де $S(t)$ – глобальний стан тестової системи, $A(t)$ – вектор дій агентів, $E(t)$ – параметри зовнішнього середовища, а оператор $F(\cdot)$ – визначає оновлення станів у результаті колективної взаємодії.

Стан кожного агента a_i на кроці (t) описується як:

$$x_i(t+1) = f_i(x_i(t), m_i(t), u_i(t), \eta_i(t)),$$

де $f_i(\cdot)$ – функція оновлення пам'яті, статистик або параметрів політики агента; $m_i(t)$ – повідомлення від інших агентів; $u_i(t)$ – локальні дії; $\eta_i(t)$ – стохастичні флуктуації або похибки вимірювань.

Для досягнення узгодженості оцінок агенти виконують процедуру дифузійного консенсусу:

$$x_i^{(k+1)} = (1 - \alpha)x_i^{(k)} + \alpha \sum_{j \in \mathcal{N}_i} w_{ij}x_j^{(k)},$$

де w_{ij} – вагові коефіцієнти комунікаційного графа, α – параметр швидкості узгодження, $\mathcal{N}_i$ – множина сусідів агента i .

Формалізоване подання моделі дає змогу описати інформаційні потоки, умови стабільності та критерії оптимізації дій агентів. Експерименти показали, що мультиагентне тестування підвищує покриття, стабільність і швидкодію перевірок, а децентралізована структура забезпечує стійкість і самонавчання системи. Запропонований підхід створює основу для розроблення адаптивних і надійних систем контролю якості програмного забезпечення.

Список використаних джерел

1. Симонов, Д. І. (2025). Ідентифікація та контроль хаотичних процесів у складних технічних системах. Науковий вісник Ужгородського університету. Серія «Математика і інформатика», 46(1), 273–284.
2. Guo, X., Wang, C., & Liu, L. (2024). Adaptive fault-tolerant control for a class of nonlinear multi-agent systems with multiple unknown time-varying control directions. *Automatisierungstechnik*. <https://doi.org/10.1016/j.automatica.2024.111802>.

РОЗВ'ЯЗУВАННЯ ЗАДАЧ ДРОБОВО-ЛІНІЙНОГО ПРОГРАМУВАННЯ З ТРАНСЦЕНДЕНТНИМИ ФУНКЦІЯМИ

Демиденко В.С., lerademidenko10112004@gmail.com

Волошко В.Л., VVL56@i.ua

Дніпровський національний університет імені Олеся Гончара

Дробово-лінійне програмування є важливим напрямом сучасної оптимізації, що дозволяє враховувати відносні показники ефективності у різних економічних задачах. На відміну від класичного лінійного програмування, яке орієнтується на оптимізацію абсолютних величин, дробово-лінійні моделі спрямовані на пошук найкращих співвідношень між результатами та витратами. Це робить їх ближчими до реальних умов господарювання. Актуальність дробово-лінійного програмування зумовлена його застосуванням у фінансовому аналізі, виробництві, логістиці, маркетингових дослідженнях та інвестиційному плануванні. Застосування відносних показників дає можливість точніше оцінити ефективність і обґрунтовувати управлінські рішення.

Особливістю задач дробово-лінійного програмування є нелінійність цільової функції, яка подається у вигляді частки двох лінійних виразів. Це ускладнює розв'язання та потребує дотримання умов коректності, зокрема невід'ємності знаменника. Для розв'язування таких задач застосовуються спеціальні методи. Найбільш відомий — метод Чарнеса-Купера, який перетворює поіменовану задачу у стандартну задачу лінійного програмування [1]. Він забезпечує високу точність і зручність реалізації, але обмежений у складних моделях. Для задач з двома змінними простим і наочним є графічний метод, який в той же час непридатний для багатовимірних випадків.

Окрему групу становлять задачі з трансцендентними функціями, які виникають у випадку нелінійних і динамічних економічних процесів. Такі

моделі точніше відображають реальність, але їх розв'язання можливе лише за допомогою чисельних алгоритмів. Реалізація аналітичного та градієнтного методів для задачі

$$Z(x_1, x_2) = \frac{2x_1 + 3x_2}{x_1 + x_2 + 1} \rightarrow \max, \quad x_1 \geq 0, \quad x_2 \geq 0$$

дає такі результати. Максимальне значення цільової функції досягається в точці (1;1) і становить аналітичним методом $\max Z(1,1) = 2,97$, градієнтним методом $\max Z(1,1) = 2,94$.

Таким чином, градієнтний метод дозволяє звести дробово-лінійну задачу до задачі лінійного програмування, що спрощує її розв'язання, він є універсальним і дозволяє розв'язувати задачі з трансцендентними функціями, які застосовуються у економічному моделюванні.

Поширеними є задачі максимізації рентабельності: знайти такий розподіл ресурсів (x_1, x_2) , який забезпечує рентабельність виробництва, тобто відношення прибутку $P(x_1, x_2)$ до сукупних витрат $C(x_1, x_2)$

$$Z(x_1, x_2) = \frac{P(x_1, x_2)}{C(x_1, x_2)} \rightarrow \max, \quad x_1 \geq 0, \quad x_2 \geq 0,$$

де $P(x_1, x_2)$ та $C(x_1, x_2)$ можуть бути нелінійними або трансцендентними. Завдання полягає в тому, щоб досягти найбільшої рентабельності, а не просто максимального доходу. Модель дозволяє визначити раціональний розподіл ресурсів при обмежених витратах і стабільних технологічних процесіях.

Бібліографічні посилання

1. Кісельова О.М., Притоманова О.М. Методи оптимізації. Ч.1. Лінійне програмування: навч. посіб. Дніпро: Вид-во ЛПРА.2021. – 168 с.

АНАЛІЗ ПРОДУКТИВНОСТІ ОБРОБКИ ВЕЛИКИХ ДАНИХ В APACHE SPARK ТА POLARS

Дерев'янюк С.О., derevankosergej@gmail.com

Мацуга О.М., olga.matsuga@gmail.com

Дніпровський національний університет імені Олеся Гончара

У сучасних умовах зростання обсягів даних виникає потреба в ефективних інструментах їх обробки, аналізу та зберігання. Платформи Apache Spark [1] та Polars [2] належать до найбільш популярних середовищ для роботи з великими наборами даних, однак ґрунтуються на різних архітектурних підходах – Spark реалізовано на Scala та орієнтовано на розподілені обчислення, тоді як Polars написано на Rust і оптимізовано для обчислень у пам'яті з мінімальною матеріалізацією.

Метою дослідження є порівняльний аналіз продуктивності систем Spark та Polars під час виконання типових операцій обробки даних: агрегації, об'єднання, фільтрації, сортування та усереднення часових рядів; та, звісно, більш складних комбінованих операцій. У роботі також розглянуто гібридний підхід (Spark + Polars), який дає змогу комбінувати розподілену обробку Spark із високошвидкісними локальними обчисленнями Polars, що у випадку складних циклічних алгоритмів з даними може дати приріст у продуктивності.

Методологія експерименту передбачала розробку універсального Python-бенчмарку з підтримкою конфігурацій JSON, логування, CLI-інтерфейсу та можливістю відстеження системних гіперпараметрів: обсягу пам'яті, кількості CPU та розміру датасету. Експерименти проводилися у середовищі Databricks з використанням Docker і Azure Storage Account як основного сховища. Для Polars застосовано оптимізований режим `scan_csv` з lazy-виконанням, для Spark – паралельні DataFrame-операції та Delta Lake.

Отримані результати показали, що Polars забезпечує нижче споживання пам'яті та стабільну криву використання ресурсів, тоді як Spark краще масштабується із зростанням обсягу даних і кількості вузлів. Гібридна комбінація Spark + Polars продемонструвала потенційне скорочення часу виконання до 20-30 % для змішаних сценаріїв (ETL + аналітика).

Результати проведених експериментів включають час виконання, обсяг використаної пам'яті, завантаження CPU, обсяг даних, тип операції тощо. Ці дані зберігаються і формують навчальну вибірку, яку планується використати для побудови моделі класифікації, здатної автоматично визначати, який рушій (Spark, Polars або їх комбінація) буде ефективнішим за заданих умов. Така модель може бути інтегрована у системи автоматизованого управління обчислювальними ресурсами, планувальники аналітичних завдань або пайплайни підготовки даних для тренування моделей машинного навчання. У перспективі це дозволить мінімізувати людське втручання у вибір технології, знизити витрати часу та ресурсів, а також створити основу для розробки самонавчальних аналітичних систем, здатних динамічно адаптувати стратегію обробки під конкретні обчислювальні середовища та цілі дослідження.

Розроблений у роботі бенчмарк спрямований на підвищення ефективності обробки великих даних у системах машинного навчання та аналітики. Він дозволяє не лише визначати оптимальний рушій виконання для кожного окремого сценарію, а й формувати базу знань про їхню поведінку при різних обсягах даних, структурах таблиць та типах трансформацій.

Список літератури

1. Apache Spark. URL: <https://spark.apache.org/> (дата звернення: 29.10.2025)
2. Polars. URL: <https://pola.rs/> (дата звернення: 29.10.2025)

СИСТЕМНИЙ ПІДХІД ДО СЕГМЕНТАЦІЇ КЛІЄНТІВ ЕЛЕКТРОННОЇ КОМЕРЦІЇ НА ОСНОВІ ПОВЕДІНКОВИХ ХАРАКТЕРИСТИК

Діденко Д.К., didenko203@gmail.com

Козакова Н.Л., kozakova.natali@gmail.com

Дніпровський національний університет імені Олеся Гончара

У роботі розглянуто проблему підвищення ефективності аналітичних і маркетингових рішень у сфері електронної комерції шляхом застосування методів кластеризації для сегментації клієнтів. В умовах цифрової трансформації бізнесу обсяги даних про поведінку користувачів зростають експоненційно, що ускладнює використання класичних статистичних методів. Застосування алгоритмів машинного навчання дозволяє автоматизувати процес аналізу клієнтської бази, підвищити точність сегментації та сформувати індивідуалізовані маркетингові стратегії.

Сучасна електронна комерція функціонує в умовах постійного зростання обсягів даних, які компанії отримують від клієнтів: історії покупок, активність у мобільних додатках, перегляди товарів, кліки по рекламі, відгуки, а також поведінкові сигнали із соціальних мереж. Ці дані містять приховану інформацію про потреби, уподобання та рівень лояльності користувачів. Проте їхня висока варіативність і обсяг роблять неможливим ручний або традиційний статистичний аналіз.

У цьому контексті важливим стає використання інтелектуальних систем аналізу, зокрема методів кластеризації, які дозволяють автоматично групувати клієнтів за подібністю їхніх характеристик і поведінки без попереднього маркування. Метою дослідження є розроблення підходу до автоматизованої сегментації клієнтів електронної комерції на основі поведінкових характеристик із використанням методів машинного навчання.

У роботі застосовано алгоритми кластеризації без учителя – K-Means, DBSCAN, Agglomerative Clustering і HDBSCAN. Реалізацію проведено з використанням бібліотек Scikit-learn, HDBSCAN та Pandas. Для попередньої обробки даних виконано нормалізацію ознак, усунення пропусків і зменшення розмірності за допомогою методу головних компонент (PCA).

Особливу увагу приділено використанню HDBSCAN – сучасного щільнісного підходу до кластеризації, який ефективно працює в умовах неоднорідної щільності даних, наявності шуму та складних форм кластерів. На відміну від DBSCAN, що потребує задання глобального параметра ϵ , HDBSCAN автоматично адаптується до локальних змін щільності, зберігаючи стійкість до шумів і здатність виявляти кластери довільної форми. Це робить його особливо корисним у задачах поведінкового аналізу, де сегменти клієнтів можуть мати різну густину та структуру.

Порівняння результатів роботи HDBSCAN з альтернативними методами (K-Means, DBSCAN, Agglomerative) дозволяє глибше зрозуміти структуру даних, визначити оптимальну кількість сегментів і підвищити обґрунтованість вибору методу кластеризації для конкретного бізнес-контексту.

Розроблена система кластеризації дозволила виділити ключові групи клієнтів: постійні покупці (висока частота покупок, стабільний середній чек); цінові мисливці (чутливі до акцій і знижок); пасивні користувачі (низька активність, високий ризик відтоку); преміальні клієнти (високий середній чек, лояльність до бренду).

Застосування HDBSCAN показало переваги у виявленні дрібних, але значущих кластерів, які залишаються непоміченими для класичних алгоритмів. Запропонована методика дозволяє ефективно працювати з великими, неоднорідними наборами даних і враховувати складну топологію простору ознак.

СУЧАСНІ МЕТОДИ КЛАСИФІКАЦІЇ ВЕБСТОРИНОК

Долотов І. О., vadol@ua.fm, Гук Н.А., huk_n@365.dnu.edu.ua

Дніпровський національний університет імені Олеся Гончара (www.dnu.dp.ua)

Одним із ключових завдань у процесі формалізації вебсайту є класифікація сторінок за їхнім типом — наприклад, розпізнавання головної сторінки, сторінки категорії, товару або статті. Така класифікація дозволяє формувати структурні шаблони вебресурсу, що є основою для побудови внутрішньої перелінковки та тематичної кластеризації.

У роботі [1] представлено комплексний огляд основних алгоритмів машинного навчання, що застосовуються для класифікації вебсторінок. На ранніх етапах домінували методи, засновані на частотному аналізі термінів (TF-IDF), однак вони демонстрували обмежену точність через ігнорування структурних характеристик HTML-документів. У статті наголошено на зростанні ролі підходів, що враховують DOM-структуру, оскільки вони забезпечують вищу дискримінативну здатність при аналізі складних вебресурсів.

У [2] проведено порівняльний аналіз методів SVM, Random Forest, Decision Tree та AdaBoost, при цьому показано, що SVM демонструє найвищу точність при комбінованому використанні текстових та DOM-ознак. Крім того, зазначено важливість попередньої векторизації сторінок, що забезпечує стабільність класифікації навіть у випадках змішаної семантики.

У дослідженні [3] автори пропонують комплексну методику прогнозування продуктивності вебсторінки на основі великого набору (59) показників і порівнюють низку класичних алгоритмів класифікації. Автори показують, що попередня обробка й відбір ознак підвищують якість моделей: SVM демонструє найкращу поведінку, тоді як Decision Tree показує найнижчу продуктивність.

У [3] акцентується увага на необхідності обробки, проте немає порівняння цих даних до та після. Для підтвердження різниці ефективності

алгоритмів, а також необхідності обробки даних, було використано ці класичні моделі на тій самій тестовій вибірці без попередньої обробки ознак. У табл.1 наведені результати дослідження на наборі даних вибраними за певними ознаками, а також власний тест алгоритмів, але без попередньої обробки, із використанням метрик: Accuracy, Precision, Recall, F1-score.

Табл.1 Порівняння ефективності методів класифікації вебсторінок

Модель	Точність (%)		Прецизійність (%)		Стійкість (%)		F1-міра (%)	
	З обробкою	Без обробки	З обробкою	Без обробки	З обробкою	Без обробки	З обробкою	Без обробки
SVM	89	87	89	87	89	88	89	87
Decision Tree	69	67	69	68	68	67	69	67
Random Forest	81	80	81	80	81	80	81	80
Ada Boost	77	76	80	78	77	77	77	77

Найбільш ефективними для розв'язання задач класифікації вебсторінок є SVM, а порівняльний аналіз показав що відбір інформативних ознак додатково підвищує точність моделей. В подальшому дослідженні доцільне використання класифікації SVM для формування вектора ознак сторінки для задачі кластеризації вебресурсів.

Список використаних джерел

- 1) Lassri S., Benlahmar E.H., Tragha A. «A review of machine learning algorithms for web page classification.». *2018 IEEE 5th International Congress on Information Science and Technology (CiSt)*. 2018; pp. 220-226. DOI: <https://doi.org/10.1109/CIST.2018.8596420>
- 2) Matošević G., Dobša J., Mladeni D. «Using Machine Learning for Web Page Classification in Search Engine Optimization.». *Future Internet*. 2021; 13(1): 9. DOI: <https://doi.org/10.3390/fi13010009>
- 3) Mohammad G., Antonio M. and Suhail O. «Novel Approach for Evaluating Web Page Performance Based on Machine Learning Algorithms and Optimization Algorithms.». *AI*, 2025; vol. 6, no. 2, p. 19. DOI: <https://doi.org/10.3390/ai6020019>

ПОКРАЩЕННЯ ОЦІНОК ПАРАМЕТРІВ СИГНАЛУ ЗА ДОПОМОГОЮ МЕТОДУ ПУЧКА МАТРИЦЬ ЗАВДЯКИ ПРОПУСКАННЮ ЗАШУМЛЕНИХ ТОЧОК

Дробахін О.О., drobakhino@gmail.com

Олевський О.В., tigozavr@gmail.com

Дніпровський національний університет імені Олеся Гончара

Метод пучка матриць (МПМ) є відомим тим, що він здатен надавати надзвичайно точні оцінки в умовах наявності білого шуму [1], проте наявність імпульсних викидів може призвести до некоректних оцінок і навіть збоїв роботи методу.

Відповідно до [1], МПМ може бути реалізованим шляхом сингулярного розкладання матриці сигналу

$$[S] = \begin{bmatrix} s_0 & \cdots & s_L \\ \vdots & \ddots & \vdots \\ s_{N-L-1} & \cdots & s_{N-1} \end{bmatrix} = [U][\Sigma][V]^H,$$

де матриці сингулярних векторів та матриця сингулярних чисел є урізаними відповідно до порядку моделі сигналу M . Дана матриця перетворюється на дві матриці для побудови пучка матриць шляхом усунення першого та останнього рядків матриці $[V]$ відповідно для матриць $[S_1]$ та $[S_2]$. Після цього проводиться оцінка ненульових власних чисел матриці $[S_1]^+ [S_2]$, що відповідають комплексно спряженим власним числам $[V_1]^+ [V_2]$, де "+" означає псевдоінверсію Мура-Пенроуза. Отримані власні числа z_m^* можуть бути перетворені на значення комплексних частот як

$$\underline{f}_m = \frac{\ln z_m}{2j\pi}.$$

Можна показати, що матриця $[\hat{S}]$ з пропущеними рядками і відповідні їй матриці $[\hat{S}_1]$ та $[\hat{S}_2]$ можуть бути отримані шляхом відкидання відповідних рядків із матриці $[U]$. Оскільки при розрахунку власних чисел значення елементів матриці $[U]$ не впливають на результат у випадку, коли вона має повний ранг, значення ненульових власних чисел $[\hat{S}_1]^+ [\hat{S}_2]$ є

такими самими, що і для випадку повної матриці сигналу. Це означає, що матриця з відкинутими рядками може бути підставлена замість повної матриці.

У якості опціонального базового підходу до ідентифікації шумових викидів було застосовано методи на основі узгодженої фільтрації та скінченних різниць, як описано в [2] для методу Проні.

Під час проведення комп'ютерного експерименту на сигналах типу $s_n = \sum_{m=1}^M \underline{A}_m e^{2\pi j f_m n \Delta t}$ було показано, що застосування МПМ за підходом відкидання точок у напівавтоматичному та автоматичному режимах призводять до покращення оцінки сигналу (рис. 1).

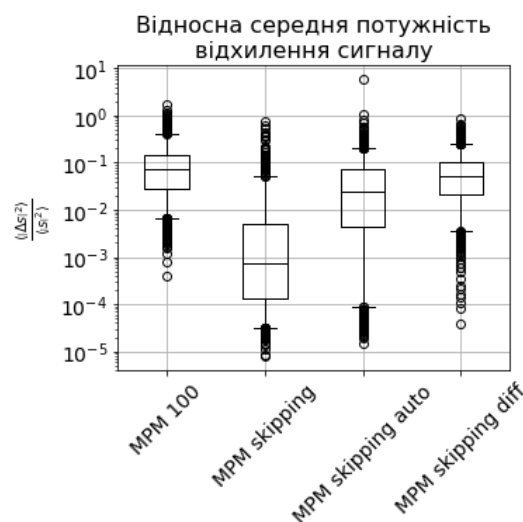


Рисунок 1. Розподіл середньої потужності відхилення сигналу від оригіналу для стандартного підходу та підходів з відкиданням точок.

Бібліографічні посилання

1. Matrix Pencil Method (MPM). *Modern Characterization of Electromagnetic Systems and Its Associated Metrology* / T. K. Sarkar et al. Hoboken, NJ : John Wiley & Sons, Inc., 2021. P. 21–106. ISBN 9781119076469.
2. Дробахін О. О., Олевський О. В. Підвищення стійкості оцінки параметрів методом Проні при наявності імпульсної завади. *Питання прикладної математики і математичного моделювання*. 2023. вип. 23. С. 53–63. URL: <https://doi.org/10.15421/322306>.

ОПТИМАЛЬНИЙ ВИННИЙ МАРШРУТ ВІДЕНЬ – ВЕНЕЦІЯ

Єгер М.Д.¹, Лефтеров О.В.², Стецюк П.І.², Федосєєв О.І.²maksym.yeher@uzhnu.edu.ua¹Ужгородський національний університет²Інститут кібернетики імені В.М. Глушкова НАН України

Задачу сформулюємо наступним чином: побудувати оптимальний маршрут з обов'язковим відвідуванням k винарень на трьох ділянках (відрізках) шляху таким чином, щоб час проходження маршруту був найменший. Перша ділянка Відень – Марібор проходить через дві країни – Австрію та Словенію, друга ділянка – Марібор – Горіція, теж проходить через дві країни – Словенію та Італію. Третя ділянка Горіція – Венеція знаходиться повністю у Італії, і винарні на цій ділянці мають як невеликі, так і значні відхилення від напрямку руху.

Математичні моделі в [1, 2] описують задачу побудови винного маршруту Львів – Вроцлав через одну ділянку. Задача побудови винних маршрутів Відень – Венеція має новий вид обмежень у вигляді ділянок, та кількості вершин обов'язкових для відвідування для кожної ділянки.

Задамо вектор Q_k такий, що $\sum_{k=1,2,3} Q_k = n - 2$, де n – загальна кількість пунктів для відвідування (міста та винарні). Кількість вершин, обов'язкових для відвідування на k -й ділянці, задамо вектором K_k . Позначимо t_{ijk} – час подорожі з пункту i до пункту j на ділянці k , T_{ik} – мінімальний час проведений у i -му пункті k -ї ділянки.

Оптимізаційну задачу для мінімізації часу проходження маршруту та відвідування винарень запишемо наступним чином:

$$t_{\min} = \min \sum_{k=1}^3 \sum_{i=1}^{Q_k} \sum_{j=1, i \neq j}^{Q_k} t_{ijk} \cdot x_{ijk} + \sum_{k=1}^3 \sum_{i=1}^{Q_k} T_{ik} \cdot y_{ik} \quad (1)$$

за обмежень:

$$\sum_{i=2}^{Q_k} x_{1ik} = 1, \quad \sum_{i=1, i \neq Q_k}^{Q_k} x_{iQ_k k} = 1, \quad \forall k = 1, 2, 3, \quad (2)$$

$$\sum_{j=1, j \neq i}^{Q_k} x_{ijk} = y_{ik}, \sum_{j=1, i \neq j}^{Q_k} x_{jik} = y_{ik}, \forall k=1,2,3; \forall i \in 1..Q_k \setminus \{1, Q_k\}, \quad (3)$$

$$\sum_{i=1}^{Q_k} y_{ik} = K_k - 2, \forall k=1,2,3, \quad (4)$$

$$u_{ik} - u_{jk} + K_k \cdot x_{ijk} \leq K_k - 1, \forall k=1,2,3; \forall i, j \in 2..Q_k, i \neq j, \quad (5)$$

$$x_{ijk} = 0 \vee 1; k=1,2,3; i, j \in 1..Q_k, i \neq j, \quad (6)$$

$$y_{ik} = 0 \vee 1; k=1,2,3; i \in 1..Q_k \setminus \{1, Q_k\}, \quad (7)$$

$$u_{ik} \in Z, 1 \leq u_{ik} \leq K_k; k=1,2,3; i \in 2..Q_k, \quad (8)$$

Обмеження (2) означають, що маршрут обов'язково починається в першій і закінчується в останній точці k -ї ділянки. Обмеження (3) гарантують, що вершина дорівнює одиниці, якщо маршрут прийшов і вийшов з i -ї точки k -ї ділянки. Обмеження (4) задають умову того, що кількість вершин включених у маршрут на k -й ділянці дорівнюють заданій кількості. Обмеження (5) забезпечують зв'язність маршруту.

Для розрахунків використовувався NEOS-сервер з солвером Gurobi версії 12.0.3. Модель (1) – (8) описана мовою моделювання AMPL. Тестовий приклад включав 2 проміжних міста (Марібор та Горіція) та 23 винарні (з них 11 на першій ділянці та по 7 винарень на другій і третій).

За допомогою Gurobi отримано оптимальний маршрут для найменшого часу ($t_{min} = 681$). Розглядається також пріоритетність близьких до оптимального маршрутів за рейтингом та кількістю відгуків.

Список літератури

1. Стецюк П.И., Лефтеров А.В., Федосеев А.И. Кратчайший k -вершинный путь. *Компьютерная математика*. 2015. **2**. С. 3–11.
2. Стецюк П.И., Лефтеров А.В., Лиховид А.П., Федосеев А.И. Оптимизационный сервис для выбора винных маршрутов. *Математическое моделирование, оптимизация и информационные технологии: материалы V Международной научной конференции*. 2016. **2**. С. 337–344.

ВДОСКОНАЛЕННЯ АЛГОРИТМУ ДИНАМІЧНОГО РОЗПОДІЛУ РЕСУРСІВ (DARA) НА ОСНОВІ ГЛИБОКОГО НАВЧАННЯ З ПІДКРІПЛЕННЯМ ДЛЯ БАГАТОКРИТЕРІАЛЬНОХ ОПТИМІЗАЦІЙ

Єлі М.Я., max.yar.eli@gmail.com

Байбуз О.Г., baibuz_o@365.dnu.edu.ua

Дніпровський національний університет імені Олеся Гончара

Вступ та аналіз досліджень. Мій попередній гібридний алгоритм DARA успішно вирішував завдання динамічного керування ресурсами, досягаючи високих показників Makespan та Utilization завдяки поєднанню LSTM-прогнозування [3] та генетичних механізмів [4]. Наразі, з огляду на глобальні вимоги до екологічної стійкості, виникла необхідність у створенні багатокритеріальної моделі, що включатиме не лише робочу ефективність, але й мінімізацію енерговитрат [1]. Існуюча оптимізаційна частина DARA не призначена для опрацювання складних, послідовних управлінських рішень, які є ключовими для контролю за живленням фізичного обладнання. Тому я модернізую DARA шляхом інтеграції Глибокого навчання з підкріпленням (DRL), що пропонує автономне та перспективне планування ресурсів [2].

.Концепція DLR-інтеграції та адаптації DARA. Я розробив архітектуру дворівневого планувальника, засновану на розширеному функціоналі DARA. Модулі DARA для LSTM-прогнозування [3] та первинного розміщення на основі генетичних алгоритмів [4] залишаються функціональними, забезпечуючи мінімальний час виконання завдань. Паралельно я вбудовую DRL-агент як спеціалізований елемент для контролю енергоспоживання:

- Обмін інформацією: Прогноз навантаження від модуля DARA передається DRL-агенту як критично важливий вхідний показник для прийняття рішень.

- Багатокритеріальна функція винагороди DRL: Навчання агента відбувається на основі функції винагороди, яка поєднує два протилежні завдання: зниження загальних енерговитрат (зменшення PUE) та запобігання падінню продуктивності (штраф за порушення Quality of Service).
- Спеціалізовані дії DRL: Набір дій агента орієнтований на енергоощадні маніпуляції, включаючи об'єднання навантаження, міграцію віртуальних машин та перемикання режимів живлення фізичних серверів, що значно розширює можливості базового DARA.

Висновок: Моделювання, проведене в середовищі CloudSim, має підтвердити, що оновлений алгоритм DARA-DRL забезпечить значне покращення показника PUE та загального зниження енергоспоживання [1] порівняно з оригінальною версією DARA, зберігаючи при цьому необхідні показники продуктивності. Завдяки інтеграції DRL [2], модернізована система набуває можливості повністю автономного та гнучкого контролю за енергозбереженням. Цей підхід є вагомим внеском у створення багатокритеріальних рішень для хмарних інфраструктур, які оптимально поєднують вимоги до ефективності роботи та екологічної відповідальності.

Список використаних джерел:

1. Beloglazov, A., Abawajy, J., & Buyya, R. (2012). Energy-aware resource allocation heuristics for efficient management of data centers for cloud computing. *Parallel and Distributed Systems, IEEE Transactions on*, Vol. 23, pp. 1656-1668.
2. Manas, V. B., & Misra, A. K. (2021). Reinforcement learning-based resource allocation in cloud computing. *Journal of Network and Computer Applications*, Vol. 176, pp. 102920.
3. IOP Conference Series (2020). Workload Forecasting in Cloud Computing using LSTM Neural Networks. *IOP Conference Series: Materials Science and Engineering*, Vol. 992.
4. Lin, Y. K., Chen, J. G. G. J. W. M. T. H. (2014). A novel load balancing algorithm for cloud computing based on genetic algorithm. *IEEE Transactions on Parallel and Distributed Systems*, Vol. 25, pp. 1234-1245.

ДО ТОЧНОГО ВИЯВЛЕННЯ R-ЗУБЦІВ В ЕЛЕКТРОКАРДІОГРАМАХ НА ОСНОВІ ПОРОГОВОГО МЕТОДУ

Єфремов М.С.¹. (yefremov@knu.ua), Ляшко А.В.¹,
Крак Ю.В.^{1,2}, Стеля О.Б.²

¹Київський національний університет імені Тараса Шевченка

²Інститут кібернетики імені В.М. Глушкова НАН України

Загальний підхід до побудови порогових методів виявлення R-зубців в електрокардіограмах (ЕКГ) складається з трьох кроків[1]: фільтрація або згладжування сигналу, обчислення першої похідної, визначення інтервалів з можливим розташуванням R-зубців та визначення самих зубців. В проведеному дослідженні, хоча і зберігається загальна структура порогового методу, але алгоритмічне наповнення кроків метода значно відрізняється від запропонованих в роботах інших авторів. Так запропоновано новий метод згладжування дискретного сигналу, який ґрунтується на кусково-поліноміальному наближенні ЕКГ. Це дозволило наблизити сигнал кусково-поліноміальною функцією другого порядку неперервною разом зі своєю першою похідною. Інтенсивність шуму в ЕКГ зроблених різними пристроями та в різних умовах може бути різною. Для досягнення необхідної для подальшого аналізу гладкості може знадобитись різна кількість повторень процедури згладжування. В дослідженні на прикладах реальних даних надані рекомендації з кількості необхідних кроків згладжування для сигналів з різною інтенсивністю шумів. На другому кроці знаходиться похідна від квадратичної функції. На відміну від багатьох публікацій, де використовуються різницеві вирази для диференціювання сіткової функції, запропонований підхід дає коректне диференціювання. При обчисленні похідних вибираються точки, де сигнал є спадною функцією. Це є фізіологічно обґрунтованим для пошуку околів R-зубців. На третьому кроці в околах точок, де квадрати похідних є найбільшими виділяються множини точок кількість яких залежить від дискретизації сигналу ЕКГ в

електрокардіографах. Наостанок на кожній множині точок знаходиться максимальне значення сигналу ЕКГ, яке приймається за R-зубець. (рис 1.).

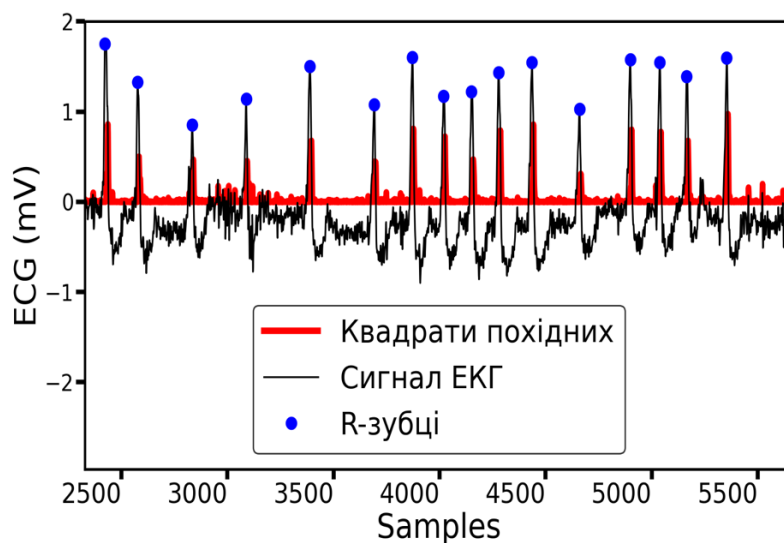


Рис. 1. Знайдені R-зубці та квадрати похідних сигналу ЕКГ

Запропонований метод було протестовано на файлах ЕКГ з бібліотеки MIT-BIH Arrhythmia Database[2], показано точне виявлення R-зубців, що свідчить про високу ефективність такого підходу. Подальші дослідження спрямовані на клінічні випробовування запропонованого методу на реальних даних.

1. Крак Ю.В. Стеля О.Б., Єфремов М.С., Ляшко А.В. Інформаційна технологія обробки даних електрокардіограм для знаходження R-піків. *Доповіді Національної академії наук України*. - 2024. - № 5. - с. 44-52. DOI: 10.15407/dopovidi2024.05.044
2. Moody G. B., Mark R. G. The impact of the MIT-BIH Arrhythmia Database. *EEE Engineering in Medicine and Biology Magazine*, 2001, vol. 20, issue 3, pp 43-50, DOI: 10.1109/51.932724.

ЗАСТОСУВАННЯ ТЕХНОЛОГІЇ DIGITAL TWINS ДЛЯ АНАЛІЗУ НАПРУЖЕНО-ДЕФОРМІВНОГО СТАНУ КОНТАКТУЮЧИХ ТІЛ

Жушман В.В., zhushmanvlad@gmail.com, Зайцева Т.А., ztan2004@ukr.net

Дніпровський національний університет імені Олеся Гончара

Сучасні підходи до проектування та аналізу механічних систем потребують ефективних інструментів для прогнозування поведінки конструкцій під різними типами навантажень. Завдяки Digital Twins стало можливим створення віртуальних моделей контактуючих тіл, які відтворюють реальні фізичні процеси та забезпечують проведення багатоваріантного аналізу без виготовлення фізичних прототипів.

Метою роботи є розробка цифрового двійника для аналізу контактної взаємодії тіл складної форми. Для цього використано програмний комплекс Ansys Workbench, який інтегрує всі етапи моделювання – від побудови геометрії до аналізу результатів та створений власний програмний комплекс.

Розглянуто процес контактної взаємодії тіл різної форми та жорсткості з деформівним півпростором. Виконано серію обчислювальних експериментів для встановлення закономірностей розподілу контактних напружень і переміщень. Проаналізовано характер зміни напруженого стану в залежності від геометричних параметрів контактуючих тіл та умов навантаження.

Результати роботи мають практичне значення для інженерних розрахунків елементів машин та конструкцій, що працюють в умовах контактної взаємодії. Робота над цифровим двійником продовжується, розроблені комп'ютерні моделі обладнано віртуальними датчиками для отримання даних про напружено-деформівний стан у контрольних точках. Для валідації цифрового двійника проведено порівняння показань віртуальних датчиків з відкритими експериментальними даними реальних випробувань.

**УДОСКОНАЛЕННЯ ЕВРИСТИЧНОГО АЛГОРИТМУ УПАКОВКИ
НЕРІВНИХ КРУГІВ В КРУГ МІНІМАЛЬНОГО РАДІУСУ****Задорожний¹ Б.О., Стецюк^{1,2} П.І.**bohzaador@gmail.com¹*Ужгородський національний університет*²*Інститут кібернетики імені В.М. Глушкова НАН України*

У бакалаврській роботі “Задача оптимальної упаковки нерівних кругів та підходи до її розв’язання” було запропоновано евристичний алгоритм для пакування нерівних кругів у круг мінімального радіусу. Нижче представлено його удосконалення, яке забезпечує підвищення швидкодії та точності отриманих розв’язків.

Задача була сформульована наступним чином: задано N кругів з радіусами r_i , таких що $r_i < r_{i+1}, i = 1, 2, \dots, N - 1$. Потрібно знайти їх центри $(x_i; y_i), i = 1, 2, \dots, N$, та радіус R зовнішнього круга із центром у точці $(0; 0)$ такі, що:

- кожний з кругів $i = 1, 2, \dots, N$ розміщується всередині зовнішнього круга (круги можуть торкатися межі зовнішнього круга);
- для будь-якої пари $(i; j)$, де $i, j \in \{1, \dots, N\}$ та $i < j$, круги не перетинаються (їм дозволено торкатися один одного);
- радіус R зовнішнього круга має бути якомога меншим.

У евристичному алгоритмі аналітичні розв’язки підзадач не застосовувалися через чутливість до похибок округлення, що уповільнювало обчислення; у вдосконаленій версії додано корекційні значення до обмежень – малі поправки, які компенсують чисельні похибки, – що підвищує стійкість і точність розв’язків.

Було удосконалено процедуру пошуку R : замість дихотомічної формули, що часто не дозволяє розміщувати або розташувала зі значним вільним простором круги, було використано уточнену формулу обчислення проміжного радіуса $R_M = R_L + (R_H - R_L) / 1.02$, де R_H та R_L

– нижня та верхня межі R . Коефіцієнт 1.02 прискорює знаходження кращого розв’язку.

У таблиці наведено порівняння роботи евристичного алгоритму та його удосконаленого варіанту для набору $r_i = i, i = 1, 2, \dots, N$ при 10^6 ітерацій.

		N				
		10	20	30	40	50
1	R	22.558131	59.949240	108.131506	165.096702	230.410788
	t, c	71.67	274.09	645.80	1196.87	2004.54
2	R	22.000196	58.741419	106.392164	162.282911	225.330514
	t, c	3.94	15.68	39.65	77.66	131.99
3	R	22.000195	58.741416	106.392162	162.282940	225.649145
	t, c	0.62	2.24	5.69	11.48	19.82
	R*	22.000193	58.400567	104.540364	158.954086	220.565400

З таблиці видно, що удосконалений евристичний алгоритм у однопоточному (рядок 2) та десятипоточному (рядок 3) режимах швидші приблизно у 15.5 та 104 разів для ($N = 40$) та 15 та 100 разів для ($N = 50$) за евристичний (рядок 1). При цьому отримано кращі результати, які визначаються вибором параметрів та особливостями розпаралелення алгоритма.

Список літератури

1. Packomania [Електронний ресурс]: <http://www.packomania.com/>
2. Задорожний Б.О., Міца О.В., Стецюк П.І. Про покращення евристичного алгоритму упаковки кругів в круг мінімального радіусу. Cybernetics and Computer Technologies. 2023. 2. С. 32–45.

ПІДВИЩЕННЯ ЯКОСТІ МОДЕЛЕЙ КЛАСИФІКАЦІЇ ТА ПРОГНОЗУВАННЯ ШЛЯХОМ ІМПУТАЦІЇ ПРОПУСКІВ

Земляний О.Д. zemlyanoy.aleksey@gmail.com, **Байбуз О.Г.**
Дніпровський національний університет імені Олеся Гончара

Актуальність проблеми. Пропуски у даних є поширеною проблемою в багатьох прикладних областях, що може суттєво впливати на якість та надійність аналітичних результатів. Наявність відсутніх значень знижує точність алгоритмів машинного навчання та може призводити до некоректних висновків, особливо у завданнях класифікації та прогнозування. Це ускладнює побудову моделей, здатних адекватно відображати реальні процеси, та підвищує ризик упереджених рішень у критичних галузях, таких як охорона здоров'я чи екологічний моніторинг.

Існуючі класичні підходи до імпутації часто не враховують складних взаємозв'язків у даних і можуть бути неефективними при роботі з різнорідними наборами. Тому розробка нових, більш точних, адаптивних і продуктивних підходів до імпутації є актуальною науково-практичною задачею.

Мета дослідження полягає у розробленні та оцінці нових методів імпутування пропусків у даних, що забезпечують підвищену точність і продуктивність у задачах класифікації та прогнозування. Дослідження спрямоване на врахування особливостей даних та знаходження певних патернів у даних, зокрема шляхом застосування методів класифікації, регресії, ентропійного підходу та кореляційного аналізу. Важливою задачею є програмна реалізація запропонованих методів в архітектурі бібліотеки scikit-learn мовою Python для уніфікованого застосування.

Методи та програмна реалізація. У ході дослідження було розроблено нові алгоритми та програмне забезпечення для двох етапів обробки даних:

1. Перетворення якісних ознак на кількісні (зі збереженням пропусків):

- IgnoreNaNLabelEncoder: Алгоритм, що базується на LabelEncoder бібліотеки scikit-learn, який перетворює якісні ознаки

на кількісні, ігноруючи пропуски (NaN) та зберігаючи можливість зворотного перетворення.

- IgnoreNANFrequentEncoder: Аналогічний кодувальник, який додатково враховує частоту появи значень при кодуванні.

2. Імпутування пропусків у даних (розроблені імп'ютери):

- UnifiedClassRegrImputer: Уніфікований метод, що застосовує класифікатор для якісних ознак та регресор – для кількісних. Має модифікації для врахування патернів у даних (наприклад, наявність рядків з одним пропуском) та сортування ознак за кількістю пропусків.

- RegrImputer: Метод, що використовує регресор для всіх типів ознак, але з подальшим коригуванням прогнозованого значення для категоріальних ознак (вибір найближчого зі словника).

- EntropyImputer: Метод на основі ентропійного підходу, що заповнює пропуски значеннями, які мінімізують умовну ентропію ознаки відносно цільової змінної.

- HybridRegrEntropyImputer: Гібридний метод, що поєднує RegrImputer (для кількісних ознак) та EntropyImputer (для якісних).

- CorrelationImputer: Метод, що використовує виявлений кореляційний зв'язок (лінійну або квазілінійну регресію) між двома ознаками для заповнення пропусків.

Усі розроблені методи реалізовано на мові програмування Python в архітектурі бібліотеки scikit-learn, наслідуючи базові класи BaseEstimator і TransformerMixin. Це забезпечує їх сумісність зі стандартними процесами обробки даних, зокрема легку інтеграцію в аналітичні конвеєри (pipeline).

Результати апробації. Ефективність розроблених методів перевірялася на реальних наборах даних: медичних (UCI Heart Disease, Framingham Heart Study) та даних гідрологічного моніторингу (басейн р. Дніпро). Проведено чотири типи тестів:

- Тест 1 (Точність імпутації та швидкодія): На даних зі штучно внесеними пропусками імп'ютер RegrImputer показав один з найкращих результатів за точністю, значно випереджаючи інші методи за швидкодією.

EntropyImputer поступався за точністю та швидкістю, однак його гібридна версія HybridRegrEntropyImputer виявилася конкурентною. CorrelationImputer продемонстрував найкращу точність на гідрологічних даних, де був присутній сильний кореляційний зв'язок.

– Тест 2 (Вплив на якість класифікації): Застосування розроблених методів імпутації суттєво покращило точність (accuracy) та повноту (recall) моделей класифікації (Random Forest). Наприклад, для датасету UCI Heart Disease точність зросла з 83% (базова модель на неповних даних) до 91% (після імпутації методами UnifiedCRSortedDesc та RegrWeight).

– Тест 3 (Вплив на якість прогнозування): На гідрологічних даних імпутація пропусків покращила якість моделей прогнозування часових рядів (XGBoost, Prophet, ARIMA) порівняно з базовими моделями, де пропуски видалялись. Найкраще себе показали імпутери, що враховують залежності між ознаками (CorrelationImputer та UnifiedClassRegrImputer).

– Тест 4 (Аналіз ентропії): Підтверджено, що запропоновані методи (особливо EntropyImputer) сприяють зменшенню умовної ентропії ознак відносно цільового класу, що свідчить про зменшення невизначеності в даних.

Висновки. У роботі розроблено інформаційну технологію імпутування пропусків у даних, що складається з набору нових алгоритмів та програмних класів. Створені методи є адаптивними, враховують структурні особливості даних (кількісні/якісні), патерни пропусків та залежності між ознаками.

Практичне значення результатів полягає у реалізації цих методів в архітектурі scikit-learn, що дозволяє легко інтегрувати їх у стандартні конвеєри (pipeline) машинного навчання. Апробація на реальних даних довела, що застосування розроблених імпутерів дозволяє підвищити точність імпутації та, як наслідок, покращити якість аналітичних моделей класифікації і прогнозування.

МОДЕЛІ І АЛГОРИТМИ СТИСКАННЯ ЗОБРАЖЕНЬ У СИСТЕМІ ОБМІНУ ФАЙЛАМИ

Золотко К. Є., zolotko_kkt@dsu.dp.ua,

Шаповаленко Н. Р., nikitashapvalenko@gmail.com

Дніпровський національний університет імені Олеся Гончара

Стрімке зростання обсягів візуального контенту, що передається, обробляється та зберігається у хмарних середовищах, породжує потребу в ефективніших та гнучкіших методах стиснення графічних файлів. Якщо раніше оптимізація зображень обмежувалась простим вибором формату (наприклад, JPEG чи PNG), то сучасна реальність вимагає динамічної адаптації під клієнтський пристрій, пропускну здатність мережі та архітектурні особливості інфраструктури. Особливої актуальності набуває автоматичне створення прев'ю (thumbnails) — компактних зображень, що швидко завантажуються та використовуються для попереднього перегляду, наприклад, у галереях, списках файлів або месенджерах.

В сучасних системах обміну файлами та веб-платформах питання оптимального стиснення зображень має вирішальне значення для ефективного використання мережевих ресурсів та покращення користувацького досвіду. Особливо це стосується генерації мініатюр — «попереднього перегляду» зображень, що за статистикою може становити до 25 % від загального обсягу сховища та до 40 % навантаження на CDN [1]. Зменшення обсягу цих прев'ю навіть на кілька кілобайт суттєво впливає на продуктивність і вартість хмарних сервісів, особливо при масштабах SaaS-платформ із десятками мільйонів активних користувачів.

Вибір форматів стиснення також суттєво впливає на результати. Формат AVIF, який підтримує покращене стиснення на 25–35 % у порівнянні з WebP [2], стає все більш популярним, однак повільніше кодується. Використання режиму `--cpu-used=8` у `libaom-av1` дозволяє значно пришвидшити кодування при незначній втраті якості, що робить

AVIF прийнятним для серверної генерації прев'ю. При цьому WebP слугує як надійний fallback для пристроїв, які не підтримують AVIF (наприклад, старі версії iOS) [3]. Застосування трирівневої ієрархії (AVIF → WebP → JPEG) забезпечує максимальну сумісність із сучасними браузерами і водночас зберігає високу ефективність передачі.

Таким чином, практика провідних компаній підтверджує ефективність використання серверлес-архітектур для генерації оптимізованих прев'ю зображень у хмарі, із пріоритетом новітніх форматів стиснення та механізмів fallback. Це дозволяє значно знизити навантаження на CDN і прискорити доставку контенту, що є критичним для масштабованих систем обміну файлами та маркетплейсів.

Бібліографічні посилання:

1. What is a content delivery network (CDN)? | How do CDNs work?:
<https://www.cloudflare.com/en-gb/learning/cdn/what-is-a-cdn/>
2. Alliance for Open Media, «AVIF Compression Benefits», 2020:
<https://aomedia.org/specifications/avif/>
3. Can I Use, Browser support for AVIF and WebP:
<https://caniuse.com/avif>

ЕЕГ У ПЛОЩИНІ ХАОТИЧНОЇ ДИНАМІКИ

Інкін О. А., Білозьоров В. Є.

Дніпровський національний університет імені Олеся Гончара

В останні 25 років моделі хаотичної динаміки почали широко застосовуватися для вивчення електроенцефалограми (ЕЕГ). Передусім, основною причиною цього охоплення є те, що дані моделі природно описують поведінку великих нелінійних систем зі зворотними зв'язками, затримками та шумом, саме таких, як кортикальні мережі. У реальних даних це проявляється чутливістю до початкових умов і різким переформатуванням режимів за незначних змін керівних параметрів. Кількісно ці властивості фіксують мірами найбільшого показника Ляпунова (LLE) та спорідненими індексами складності. Зокрема, при зміні глибини анестезії LLE ЕЕГ суттєво зменшується або зростає залежно від препарату та стану, що узгоджено з уявленням про віддалення системи від межі хаосу [1]. Останні огляди підсумовують дедалі ширші застосування хаотичних і нелінійних методів до когнітивних і клінічних задач, від моніторингу станів до розладів ЦНС [2].

Щоби ця мова хаосу була не лише дескриптивною, а й прикладною, її доцільно поєднувати з генеративними нейронними моделями мас (NMM). У канонічній моделі Янсена–Ріта (JR) локальна кортикальна колонка задається кількома взаємодіючими популяціями з нелінійним перетворенням «мембранний потенціал - частота спайків», і вже на цьому рівні модель здатна продукувати ЕЕГ-подібні ритми. Важливо, що під біологічно правдоподібним періодичним драйвом JR відтворює квазіперіодику та хаос у межах реалістичних параметрів - тобто складність тут не припускається, а виникає як наслідок взаємодії збудження, гальмування, зворотних зв'язків і зовнішнього ритмічного впливу [3]. На рівні кількох взаємодіючих колонок діаграми біфуркацій демонструють багатий репертуар режимів, чутливих до сили зв'язку та фону вхідного потоку. Коли ж таку динаміку масштабують до цілісного мозкового рівня, додаючи реалістичні з'єднувальні матриці та розподілені аксональні затримки, NMM здатні відтворювати просторово-часові патерни, ближчі до ЕЕГ. Сучасні інструменти роблять ці симуляції і точними, і

обчислювально керованими, а відповідні платформи забезпечують практичну екосистему для таких експериментів [4].

Для епілепсії розширена JR-модель Вендлінга–Шовеля вводить окремі гальмівні контури (повільний/швидкий) і тим самим відтворює перехідні траєкторії «інтеріктал - іктал», які добре описуються мовою таких систем. Аналіз показує наявність специфічних сценаріїв при вході в напад і пояснює, як малі зсуви у балансі збудження/гальмування змінюють топологію фазового простору з появою пароксизмів. У цьому сенсі хаотична динаміка не імітує ЕЕГ, а надає робочий словник для механістичного обґрунтування переходів між режимами і зворотної інженерії параметрів і до персоналізації [5].

Підсумовуючи, хаотична динаміка підсилює моделювання ЕЕГ саме там, де важливі переходи між станами, чутливість до впливів і багатство тимчасових масштабів. У поєднанні з нейронними моделями мас (від локальних JR до мережових і мозкових із реалістичними затримками) вона забезпечує узгоджений каркас, придатний і для базового розуміння, і для клінічного застосування. Життєздатність підходу забезпечують відтворюваність результатів, їхня незалежна валідація на різних наборах ЕЕГ та узгодженість між модельними передбаченнями й емпіричними даними [6].

[1] Hayashi, K. (2024). Chaotic nature of the electroencephalogram during shallow and deep anesthesia: From analysis of the Lyapunov exponent. *Neuroscience*.

[2] Kargarnovin, S., Hernandez, C., Farahani, F. V., & Karwowski, W. (2023). Evidence of Chaos in Electroencephalogram Signatures of Human Performance: A Systematic Review. *Brain Sciences*, 13(5), 813.

[3] Spiegler, A., Knösche, T. R., Schwab, K., Haueisen, J., & Atay, F. M. (2011). Modeling Brain Resonance Phenomena Using a Neural Mass Model. *PLoS Computational Biology*, 7(12).

[4] Sanz Leon, P., Knock, S. A., Woodman, M. M., Domide, L., Mersmann, J., McIntosh, A. R., & Jirsa, V. (2013b). The Virtual Brain: A simulator of primate brain network dynamics. *Frontiers in Neuroinformatics*, 7.

[5] Köksal Ersöz, E., Modolo, J., Bartolomei, F., & Wendling, F. (2020). Neural mass modeling of slow-fast dynamics of seizure initiation and abortion. *PLOS Computational Biology*, 16(11).

[6] Belozyorov V. Y., Inkin O. A. (2023). Systems of singular differential equations as the basis for neural network modeling of chaotic processes. *Journal of Optimization, Differential Equations and Their Applications*, 31(2), 24.

ТЕХНОЛОГІЇ АВТОНОМНОЇ НАВІГАЦІЇ ДЛЯ ЛЮДЕЙ ІЗ ВАДАМИ ЗОРУ

Каманцев А.С., kamantsev_a@365.dnu.edu.ua

Гарл Л.Л., ll_hart@ukr.net

Дніпровський національний університет імені Олеся Гончара (www.dnu.dp.ua)

Задача навігація для людей із вадами зору є комплексною, тому підходи до її розв'язання суттєво відрізняються як апаратним так і програмним забезпеченням. Рішення починаються від низько технологічних, таких як тактильна тростина, і закінчуються високотехнологічними, такими як портативні GPS-пристрої [1].

У роботах [2 – 8] розглянуті різні системи навігації та описані технології, які можна згрупувати наступним чином:

1. **Камери:** 1-2 камери і обробка даних на пристрої; численні IP-камери та віддалене обчислення; RGB-D камери (створюють карту глибини).
2. **Роботизовані системи:** розумна тактильні палиця; робот, що аналізує фізичні властивості об'єктів (тиск, текстура, динамічна деформація і т.д.); робот-поводир; розумні черевики із вбудованим ультразвуковим датчиком, що дозволяють уникати перешкод.
3. **Датчики:** інерційний вимірювальний пристрій (IMU); ультразвукові; bluetooth, інфрачервоні; радари; SONAR; NFC; виявлення та вимірювання дальності світла (LIDAR); радіочастотна ідентифікація (RFID); індикатор сили прийнятого сигналу (RSSI).
4. **Глобальні системи позиціонування:** GPS; Galileo; BeiDou; GLONASS; GSM; RTK; QZSS.
5. **Бібліотеки:** Google Project Tango Development Kit; ARUCO.

6. **Способи взаємодії із користувачем:** аудіо введення/виведення; 3-D звук; тактильний (вібрація: через стопи, долоні і т.д.).
7. **Архітектури:** обчислення на пристрої; клієнт-серверна; комунікація із навколишніми пристроями; mesh-network системи.
8. **Алгоритми машинного навчання:** класифікація; сегментація; текстовий опис зображень; розпізнавання шаблонів (та їх руху у кадрі); виявлення обличь; розпізнавання тексту (OCR); розпізнавання спеціальних міток (зокрема QR-кодів); visual simultaneous localization and mapping (VSLAM); continuously adaptive mean-shift (CAMShift); D*; SIFT; SURF; HOG detector; FAST; SVM; Hough transform; архітектура MobileNet-V2 (у якості backbone для моделі виявлення об'єктів).
9. **Інші технології:** відео зв'язок із асистентами; “Соціальні сенсори” - аналіз постів у соціальних мережах і виявлення небезпек уздовж запланованого маршруту; використання доповненої реальності для навчання людей із порушеннями зору мобільності; штучне око; тріангуляція - локалізація на основі кількох ідентифікаторів та тріангуляції; інтернет речей (IoT); “тактильні індикатори поверхні для ходіння”- довгі прутки з плоским верхом, розміщені на підлозі, допомагають незрячій людині притримуватись напрямку руху.

Наведений список технологій є доволі великим, проте множина рішень, побудованих на їх основі, є ще більш різноманітною. Та оскільки список викликів і задач, що стоять перед людиною із вадами зору, є не менш обширним, то пошук нових конструктивних рішень залишається актуальною задачею у сфері вищезначеної проблематики.

Список використаних джерел

1. «Just Let the Cane Hit It»: How the Blind and Sighted See Navigation Differently / M. A. Williams, C. Galbraith, K. K. Shaun, A. Hurst. 16th International ACM SIGACCESS Conference on Computers & Accessibility : Proceedings. Rochester, 2014, p. 217–224. DOI: 10.1145/2661334.2661380.
2. A comprehensive review of navigation systems for visually impaired individuals / M. A. Haider, A. N. Siddiquee, H. Alkhalefah, V. Srivastava. Heliyon. 2024. v. 10, № 11. p. e31825. DOI: 10.1016/j.heliyon.2024.e31825.
3. Analysis of Navigation Assistants for Blind and Visually Impaired People: A Systematic Review / K. Sulaiman, N. Shah, K. U. Habib. IEEE Access. 2021. v. 9. DOI: 10.1109/ACCESS.2021.3052415.
4. A review of assistive spatial orientation and navigation technologies for the visually impaired / H. Fernandes [et al.]. Universal Access in the Information Society (01 March 2019). 2019. v. 18, № 1. p. 155-168.
5. Research advances of indoor navigation for blind people: A brief review of technological instrumentation / P. Darius. IEEE Instrumentation & Measurement Magazine . 2020. v. 23. DOI: 10.1109/MIM.2020.9126068.
6. Giudice N. A., Legge G. E. Blind Navigation and the Role of Technology. The Engineering Handbook of Smart Technology for Aging, Disability, and Independence. 2008. p. 479-500. DOI: 10.1002/9780470379424.ch25.
7. Kuriakose B., Shrestha R., Sandnes F. E. Tools and Technologies for Blind and Visually Impaired Navigation Support: A Review. IETE Technical Review. 2022. v. 39, №1. p. 3-18. DOI: 10.1080/02564602.2020.1819893.
8. Comprehensive Review: High-Performance Positioning Systems for Navigation and Wayfinding for Visually Impaired People / J. Feghali, C. Feng, A. Majumdar, W. Ochieng. Sensors. v. 24. 2024. p. 7020. DOI: 10.3390/s24217020.

ІНТЕЛЕКТУАЛЬНІ АЛГОРИТМИ ПІДВИЩЕННЯ ЯКОСТІ КОРИСТУВАЦЬКОГО ДОСВІДУ В АДАПТИВНИХ ВІДЕОСИСТЕМАХ СИТУАЦІЙНОГО ЦЕНТРУ

Кармазін К.В., kirillkarmazin2301@gmail.com

Інститут кібернетики НАН України імені В.М. Глушкова

У сучасних умовах зростання техногенних ризиків, урбанізації та військових загроз важливою стає здатність ситуаційних центрів (СЦ) забезпечувати стійкий відеомоніторинг подій у реальному часі. Основним інструментом таких систем є адаптивні відеопотоки, які надходять від різномірних джерел – безпілотників, мобільних роботів і сенсорних платформ технопарку робототехнічних систем (TPC) [1]. Водночас динамічність мережевого середовища, енергетичні обмеження та флуктуації пропускної здатності призводять до зниження якості користувацького досвіду Quality of Experience (QoE) операторів [2]. Для вирішення цієї проблеми запропоновано концепцію резильєнтності поточкових відеосистем, що поєднує прогнозування, навчання з підкріпленням Reinforcement Learning (RL), мультиагентну координацію [3] та генеративну реконструкцію відео. Запропоновано математичну модель Resilient Quality of Experience (RQE), яка оцінює сталість сприйняття відео користувачем, враховуючи середнє значення QoE, дисперсію якості та частку часу буферизації. На цій основі розроблено алгоритм RL-ABR-RQE, де функція винагороди базується на прирості RQE. Це дозволяє адаптивно обирати бітрейт з урахуванням прогнозованої пропускної здатності та довгострокової стабільності відео. Порівняння підходів до обробки і передачі відео наведено у таблиці 1. Особливу увагу приділено генеративному відновленню відео з використанням моделей типу GAN і Diffusion Models.

Підхід	Стабільність QoE	Буферизація	Складність	Примітки
Евристичний ABR	Середня	Середня/Висока	Низька	Чутливий до збурень
RL-ABR	Висока	Низька	Середня/Висока	Оптимізація під метрику RQE
RQE-ABR	Дуже висока	Низька	Висока	Компенсація втрат кадрів

Таблиця 1. Порівняння підходів до обробки/передачі відео

Це дає змогу компенсувати втрати кадрів і артефакти під час передачі, і підвищити суб'єктивну якість відео навіть при 10-15 % втрат даних. Запропоновано механізм компромісу між якістю та стабільністю, який дозволяє системі знижувати бітрейт у моменти перевантаження, зберігаючи візуальну цілісність через реконструкцію дрібних деталей.

1. Kateryna Melkumian, Julia Pisarenko, Alexander Koval. Organization of Regional Situational Centers of the Intelligent System “Control_TEA” using UAVs // Proceedings of the 2021 IEEE 6th International Conference on Actual Problems of Unmanned Aerial Vehicles Development (APUAVD-2021). – PP. 37-40. [DOI: 10.1109/APUAVD53804.2021.9615416].
2. Ajeypasaath K. B., Vetrivelan P. A Hybrid Machine Learning Approach for Improved QoE in Video Services over 5G Wireless Networks // Computers, Materials & Continua. – 2024. – Vol. 78. – PP. 3195–3213. [DOI: 10.32604/cmc.2023.046911].
3. Kateryna Mamedova, Julia Pisarenko, Alexander Koval, Rasya Mamedov. Chapter “Information System for Certifying UAV Operations within the Intelligent System Concept «Control TEA» // book «Advanced Information-Measuring Technologies and Systems II». – ISBN 978-3-032-02770-2. [link.springer.com/book/9783032027696].

ДОСЛІДЖЕННЯ ЕФЕКТИВНОСТІ НЕЙРОННИХ МЕРЕЖ У ЗАДАЧІ КЛАСИФІКАЦІЇ РЕАЛЬНИХ ТА ЗГЕНЕРОВАНИХ ЗОБРАЖЕНЬ ОБЛИЧ

Ківа М. В., maksym.kiva@gmail.com

Мацуга О. М., olga.matsuga@gmail.com

Дніпровський національний університет імені Олеся Гончара

Стрімкий розвиток генеративних моделей призвів до появи надзвичайно реалістичних зображень облич, які важко відрізнити від справжніх. Це створює загрозу поширення deepfake-контенту, що може використовуватися для маніпуляцій чи шахрайства. У даній роботі поставлено задачу порівняти архітектури нейронних мереж для класифікації справжніх та згенерованих зображень облич людей.

У межах дослідження було сформовано датасет, що включає три класи зображень: реальні обличчя, згенеровані моделями GAN (StyleGAN, StyleGAN2, FaceSwap), а також створені за допомогою моделей Stable Diffusion (версії 1.5, 2.1, SDXL 1.0). Було використано публічні набори на Kaggle: Stable Diffusion Face Dataset та Real vs AI Generated Faces Dataset.

Сформований датасет обсягом 27 000 зображень було розділено на 3 частини: 70% для навчання (18 900), 15% для валідації (4 050) та 15% для тестування (4 050). Попередня обробка включала приведення зображень до розміру 224×224 , аугментацію для навчального набору (RandomCrop, HorizontalFlip, ColorJitter) та нормалізацію за параметрами ImageNet.

У роботі порівнювалась ефективність шести архітектур нейронних мереж: MobileNetV2, ResNet50, EfficientNetB4, DenseNet121, ViT-B/16 і ConvNeXt-Tiny. Серед них перші чотири – класичні CNN, ConvNeXt-Tiny – сучасна архітектура, що поєднує CNN та Transformer, а ViT-B/16 відповідає архітектурі чистого Transformer.

Навчання моделей проводилося з використанням бібліотеки PyTorch у локальному середовищі, а також на платформі Google Colab. Усі нейронні

мережі попередньо натреновані на ImageNet, тому у роботі застосовано transfer learning з донавчанням лише останнього класифікаційного шару. Для CNN-архітектур використовувалися такі параметри: функція втрат – CrossEntropyLoss, оптимізатор Adam з $lr=3e-4$, планувальник швидкості навчання – ReduceLROnPlateau, кількість епох – 5, розмір батчу – 64. Для Transformer-архітектури ViT-B/16 застосовувалися: функція втрат CrossEntropyLoss, оптимізатор AdamW з $lr=1e-4$ та $weight_decay=0.05$, швидкість навчання регулювалася за допомогою CosineAnnealingLR, кількість епох – 20 епох, розмір батчу – 48.

У таблиці 1 наведено значення основних метрик якості навчених моделей, розрахованих на тестовій вибірці.

Таблиця 1 – Метрики якості моделей

Модель	Accuracy	Precision	Recall	F1-score
MobileNetV2	95.67	95.96	95.67	95.66
ResNet50	95.41	95.80	95.41	95.39
EfficientNetB4	96.17	96.38	96.17	96.17
DenseNet121	95.43	95.85	95.43	95.42
Vit-B / 16	95.45	95.75	95.45	95.44
ConvNeXt-Tiny	97.83	97.86	97.83	97.83

Одержані результати свідчать, що точність класифікації є дуже високою для усіх моделей, а різниця в архітектурах впливає мінімально на значення метрик. Навіть для моделі MobileNetV2, яка є найлегшою серед розглянутих, точність досягає 95.67%. Найкращі результати продемонструвала модель ConvNeXt-Tiny: accuracy – 97.83, precision – 97.86, recall – 97.83, F1-score – 97.83.

В подальшому планується розробка власної CNN архітектури для досягнення аналогічних результатів з меншою кількістю параметрів.

ДИНАМІЧНІ ЗАДАЧІ ОПТИМАЛЬНОГО РОЗБИТТЯ МНОЖИН: МАТЕМАТИЧНІ МОДЕЛІ, МЕТОДИ, АЛГОРИТМИ, ПРАКТИЧНЕ ЗАСТОСУВАННЯ

Кісельова О.М., Кузенков О.О., Закутній Д.

Дніпровський національний університет імені Олеся Гончара (www.dnu.dp.ua)

Оптимальне розбиття множин є однією з ключових проблем сучасної прикладної математики та теорії оптимізації, яка знаходить широке застосування в логістиці, інформатиці, біоінженерії, моделюванні складних систем та штучному інтелекті. Особливий інтерес становлять динамічні варіанти задач оптимального розбиття, коли умови задачі змінюються у часі, а розбиття має адаптуватися відповідно до динаміки системи. Оптимальне розбиття множин, в переважній більшості прикладних задач, безпосередньо пов'язано з мінімізацією цільового функціоналу, невід'ємною складовою якого є не тільки контури підмножин, але й інші визначальні параметри, що є ключовими для шуканих підмножин. В класичних постановках такими параметрами є, наприклад, центри підмножин. Прикладне застосування задач в такій постановці знаходять своє залучення в економіці, логістиці, медицині, архітектурі та інших галузях.

Ключова особливість таких задач полягає в тому, що навіть при наявності спрощених припущень щодо геометрії множини та характеру змін її внутрішніх властивостей, оптимальне розбиття здатне проявляти складну топологічну організацію. Це зумовлює потребу в розробці спеціалізованих методів аналізу, спрямованих на виявлення умов існування, єдиності та стійкості розв'язків відносно малих збурень у вхідних параметрах.

Динамічні задачі оптимального розбиття множин із варіативним розташуванням центрів постають у ситуаціях, коли об'єкти потребують не лише класифікації, а й просторової локалізації. Типовим прикладом таких задач є моделювання конкурентної взаємодії мобільних торговельних точок у межах урбанізованого середовища. У подібних сценаріях, окрім запланованих змін цінової політики, пересувні магазини здатні змінювати своє географічне положення з метою підвищення привабливості для споживачів. Такі переміщення безпосередньо впливають на геометрію зон обслуговування.

МЕТОД РОЗВ'ЯЗАННЯ НЕЧІТКОЇ ДВОЕТАПНОЇ ЗАДАЧІ ОПТИМАЛЬНОГО РОЗБИТТЯ МНОЖИН

Кісельова¹ О.М., Притоманова² О.М., Лебедєв¹ Д.,
mstr.danila@gmail.com

¹*Oles Honchar Dnipro National University (www.dnu.dp.ua)*

²*Kyiv National Economic University named after Vadym Hetman (www.kneu.edu.ua)*

Пропонується метод розв'язання двоетапної неперервно-дискретної транспортної задачі оптимального нечіткого розбиття-розподілення, що базується на поєднанні методів теорії оптимального розбиття множин, теорії нечітких множин та дискретної транспортної теорії. Теорія оптимального розбиття множин є перспективним напрямом розв'язання задач нескінченновимірної оптимізації, а дослідження задач з нечіткістю у розв'язку дозволяє більш реалістично моделювати прикладні задачі.

Поставлену задачу формалізовано шляхом введення функцій належності нечітких підмножин та коефіцієнту нечіткості. Доведено теорему, що визначає вид оптимального розв'язку задачі, який включає визначення оптимального нечіткого розбиття множини на підмножини та оптимальних об'ємів транспортування ресурсів між центрами першого та другого етапів. На основі теореми сформульовано чисельний алгоритм розв'язання задачі, у якому для розв'язання задачі безумовної оптимізації застосовується один із варіантів субградієнтного методу недиференційовної оптимізації r -алгоритм Шора.

На основі запропонованого алгоритму створено програму реалізацію мовою C++. Програма ефективно розв'язує задачу, надає проміжні та кінцеві числові результати, а також виконує інтерактивну візуалізацію оптимального розв'язку. Передбачено можливість регулювання точності обчислень шляхом зміни роздільної здатності дискретної сітки простору, параметрів чисельного інтегрування, параметрів градієнтного спуску. Проведено ряд експериментів з модельними задачами, на яких досліджено роботу алгоритму, проведено порівняльний аналіз оптимального нечіткого розбиття в залежності від заданого коефіцієнта нечіткості.

ВПЛИВ ПОВЕРХНЕВОЇ ПРУЖНОСТІ НА НАНОТРІЩИНУ В ІЗОТРОПНОМУ МАТЕРІАЛІ

Клецьков О.М., alex.kl87@i.ua,

Шевельова А.Є. shevelevaaea@dnu.dp.ua,

Лобода В.В. loboda@dnu.dp.ua

Дніпровський національний університет імені Олеся Гончара (www.dnu.dp.ua)

Досліджено вплив поверхневої пружності на пружно-деформований стан тріщини типу III, що виникає за умов антиплоских зсувних деформацій у лінійно-пружному середовищі. Механічні ефекти, які проявляються поблизу поверхонь, зокрема на берегах тріщин, враховано в рамках континуальної моделі поверхневої пружності Гуртіна–Мердока. Отримано уточнені граничні умови на верхньому та нижньому берегах тріщини, які далі досліджуються із застосуванням методів теорії функцій комплексної змінної.

Поставлена задача зводиться до сингулярного інтегро-диференціального рівняння першого порядку з ядром типу Коші. Для його розв'язання використано представлення невідомих функцій у вигляді розкладень за многочленами Чебишова першого роду та застосовано метод колокації у вузлах цих многочленів. Розв'язання отриманої системи лінійних алгебраїчних рівнянь дало змогу визначити коефіцієнти зазначених розкладень. На основі отриманих результатів виведено формулу для обчислення напружень на продовженні тріщини, яка виражається інтегралом із ядром типу Коші.

Виконано детальний аналіз запропонованого числового алгоритму, який охоплює варіювання кількості членів у розкладі невідомих функцій за многочленами Чебишова та зміни кількості вузлів у квадратурних формулах Гауса, що використовуються для числового інтегрування. Показано, що вплив поверхневої пружності стає істотним у випадку тріщин нанометрового масштабу – коли їх довжина не перевищує одного мікрметра. Подальше зменшення розмірів тріщини суттєво змінює характер розподілу напружень поблизу її вершини: зникає коренева особливість напружень, і значення напружень у вершинах набувають скінченних величин.

ІНФОРМАЦІЙНА ТЕХНОЛОГІЯ ОБРОБКИ ВИМІРЮВАНЬ ТА ЇХ АНАЛІЗУ З ВИКОРИСТАННЯМ ПАРАМЕТРИЧНИХ ТА НЕПАРАМЕТРИЧНИХ КРИТЕРІЇВ

Клименко О.Д., klymenko_o@365.dnu.edu.ua

Дніпровський національний університет імені Олеся Гончара

Стрімкий розвиток інформаційних технологій та інтелектуальних систем зумовлює потребу в підвищенні точності, достовірності та швидкодії обробки результатів вимірювань. Одним із важливих напрямів сучасної метрології та технічної діагностики є створення адаптивних методів статистичного аналізу даних, здатних ефективно працювати в умовах невизначеності, шумів і відхилень від класичних припущень щодо розподілу випадкових величин.

Традиційні параметричні критерії статистичного аналізу, такі як t-критерій Стюдента, F-критерій Фішера, критерій Пірсона, базуються на припущенні нормальності розподілу даних. У випадках, коли вибірки містять викиди або мають асиметричні властивості, застосування цих критеріїв може призводити до помилок у прийнятті рішень. Тому актуальною є розробка інформаційних технологій, які поєднують застосування параметричних та непараметричних методів для підвищення достовірності обробки вимірювань та їх подальшого аналізу.

Метою роботи є створення інформаційної технології обробки результатів вимірювань з використанням параметричних та непараметричних критеріїв для забезпечення підвищеної точності статистичних оцінок і виявлення аномалій у вибірках. У межах дослідження розроблено концепцію інтегрованої технології, яка включає такі етапи: первинну фільтрацію та нормалізацію даних; перевірку гіпотези про однорідність та нормальність розподілу (критерії Андерсона-Дарлінга, Колмогорова-Смирнова, Буша-Вінда, Шапіро-Уїлка); вибору методу

обробки даних залежно від результатів перевірки (параметричний чи непараметричний аналіз); формування узагальнених висновків щодо достовірності даних.

Розроблено інноваційну інформаційну технологію, що не має аналогів в Україні. Створена інформаційна технологія надає користувачам доступ до генерації, візуального аналізу та статистичного порівняння вибірок, зокрема – перевірки їх на нормальність розподілу за допомогою сучасних критеріїв статистики (рис.1).

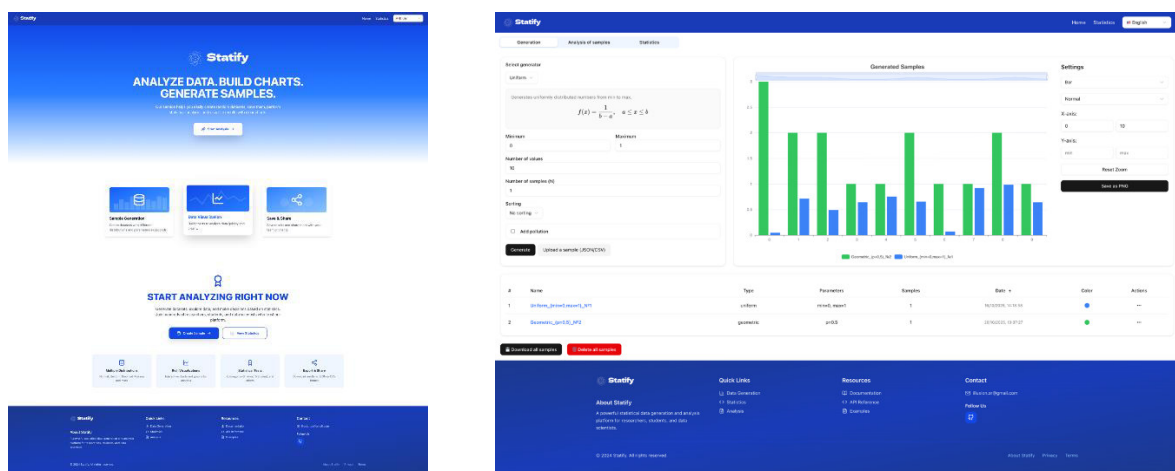


Рис. 1. Приклади інтерфейсу інформаційної технології обробки вимірювань та їх аналізу з використанням параметричних та непараметричних критеріїв.

Інформаційна технологія має двомовний інтерфейс (українська та англійська мови), що робить її зручною для використання як у міжнародному науковому середовищі, так і в освітньому процесі у вищих навчальних закладах України. Запропонована інформаційна технологія є сучасним інструментом цифрової статистики, що поєднує функціональність, наочність і доступність для дослідників та студентів.

ФРАГМЕНТАРНІ МОДЕЛІ ДЛЯ ЗАДАЧ ПОКРИТТЯ МНОЖИНИ ТОЧОК

Козін І.В., ainc00@gmail.com,

Сардак О.В., karnelcore@gmail.com

Запорізький національний університет

Задачі покриття множини точок мають численні приклади використання. Вони часто з'являються у металургії, текстильній та деревообробній промисловості, будівництві, землеустрої, логістиці, військовій справі. Ці задачі мають безліч різних постановок і здебільшого відносяться до важкорозв'язних. Класичний варіант задачі покриття множини точок на площині виглядає так: на площині задано множину точок своїми координатами і заданий набір плоских фігур, для яких допускаються певні ізометричні перетворення (зрушення, повороти). Задача полягає в пошуку мінімального за загальною площею набору фігур, які при відповідному розміщенні покривають всю множину точок. Умови таких задач також можуть бути різними. Наприклад, потрібно, щоб елементи покриття не перетиналися або перетиналися лише межами. До таких задач належать багато задач упаковки та розкрою.

Узагальнена задача покриття вершин на зваженому графі має наступну постановку: визначити набір підграфів певного виду, які містять усі вершини графа та мають мінімальну вагу. Прикладами таких задач є задачі покриття графа ребрами, зірками, регулярними графами та інші. Задача покриття графу зірками має значні прикладні аспекти. Зокрема, вона з'являється в задачі землекористування, задачах розміщення виробництва, оптимальній класифікації тощо.

Для пошуку субоптимальних рішень задач покриття пропонується спосіб заснований на побудові фрагментарної моделі задачі[1].

Кожне припустиме рішення задачі покриття представляється як об'єднання окремих частин-фрагментів рішення. Спосіб з'єднання та вибір послідовності фрагментів залежить від конкретної постановки задачі. Множина фрагментів скінчена і, таким чином, задача зводиться до вибору певної перестановки фрагментів відповідно до заданого способу їх приєднання.

Це уявлення породжує ін'єктивне відображення множини перестановок у множину допустимих розв'язків задачі. Таким чином, вдається звести задачу оптимізації до безумовної задачі пошуку оптимальної перестановки фрагментів.

Для пошуку оптимальної перестановки пропонується використовувати метаевристики різних типів.

Найбільш цікаві результати отримані для еволюційних метаевристик з геометричним оператором кросоверу [2]. Наявність метрики у просторі перестановок дозволяє визначити поняття відрізка і, відповідно, опуклої множини. Геометричний оператор кросоверу обирає результат кросоверу двох перестановок як точку на відрізку, що їх з'єднує. Таким чином, операція кросоверу призводить до зменшення простору пошуку. З іншого боку, мутація розглядається як оператор, який виводить за межі опуклої оболонки існуючої популяції рішень і таким чином призводить до розширення простору пошуку.

На базі фрагментарних моделей можна будувати гібридні методи, засновані на комбінації різних метаевристик. Наприклад, комбінація еволюційного алгоритму, алгоритму перемішаних стрибавих жаб, алгоритму імітації відпалу з алгоритмом локального пошуку дозволяє значно покращити рішення з точки зору їх наближення до оптимуму.

Література

1. I.V. Kozin, N. K. Maksyshko, V. A. Perepelitsa Fragmentary Structures in Discrete Optimization Problems, Cybernetics and Systems Analysis November 2017, Volume 53, Issue 6, P 931–936. <https://doi.org/10.1007/s10559-017-9995-6>
2. A.Moraglio R.Poli, Topological Crossover for the Permutation Representation// *Intelligenza Artificiale*, – 2011. volume 5, issue 1, –P. 49-70.

БАГАТОКРИТЕРІЙНА ОПТИМІЗАЦІЯ ІНВЕСТИЦІЙНОГО ПОРТФЕЛЯ

Колесніков Д.С., danylokolesnikov@gmail.com

Інститут кібернетики імені В.М. Глушкова НАН України

Відома оптимізація портфеля акцій за класичною моделлю Марковіца [1–3] базується на мінімізації дисперсії портфеля при заданій очікуваній дохідності. Такий підхід дозволяє отримати ефективну межу “ризик – дохідність”, проте не враховує багатокритерійний характер реальних фінансових задач.

Метою даної роботи є побудова багатокритерійної математичної моделі оптимізації інвестиційного портфеля, що дозволяє отримати множину Парето-оптимальних рішень і використовується як основа для систем підтримки прийняття рішень. Управління ризиком трактується як процес кількісної оцінки та мінімізації можливих втрат [4].

Такий підхід лежить в основі критерію ризику

$$f_1(x) = w^T \Sigma w = \sum_i \sum_j w_i w_j \sigma_{ij} \quad (1)$$

Математична модель визначається як задача багатокритерійної оптимізації:

$$\min F(x) = \{f_1(x), f_2(x), f_3(x), f_4(x), f_5(x)\}, \quad (2)$$

де

$$f_1(x) = w^T \Sigma w = \sum_i \sum_j w_i w_j \sigma_{ij} - \text{ризик портфеля (дисперсія)}, \quad (3)$$

$$f_2(x) = -\sum_i w_i r_i - \text{дохідність (максимізація)}, \quad (4)$$

$$f_3(x) = \sum_i \frac{w_i}{L_i}, - \text{ліквідність}, \quad (5)$$

$$f_4(x) = \sum_i c_i w_i, - \text{транзакційні витрати}, \quad (6)$$

$$f_5(x) = \sum_i w_i^2 - \text{диверсифікація.} \quad (7)$$

Для побудови Парето множини використано еволюційний алгоритм NSGA-II [5]. У результаті роботи алгоритму формується множина Парето, що містить усі ефективні комбінації ризику та дохідності. Порівняння результатів показало, що при фіксованих рівнях дохідності багатокритерійна оптимізація забезпечує зниження портфельного ризику на 14–17% у порівнянні з класичною моделлю Марковіца [3]. Це пояснюється урахуванням додаткових критеріїв ліквідності, транзакційних витрат і диверсифікації, що зменшує коваріаційний ефект між активами [6]. Практичне значення запропонованої моделі полягає у можливості формування індивідуальної траєкторії інвестора на множині Парето-оптимальних розв’язків. Модель може бути основою систем підтримки прийняття рішень для фінансової аналітики.

Список літератури:

1. Markowitz H.M. Portfolio Selection. The Journal of Finance. 1952. 7, N 1. P. 77-91.
2. Markowitz H.M. Portfolio Selection: Efficient Diversification of Investments New York: John Wiley & Sons, Inc.; London: Chapman & Hall, Ltd.; Cowles Foundation for Research in Economics at Yale University, 1959. 344 p.
3. Kolesnikov D.S., Semenova N.V. (2024). Optimization of the Investment Portfolio. XXXIX International Conference “*Problems of decision making under uncertainties*”. PDMU-2024. (September 9–10, 2024, Brno, Czech Republic). Видавництво “Людмила”, 152 с. P. 76–77.
4. Гуменюк В.Я., Міщук Г.Ю., Олійник О.О. (2009). Управління ризиками: навчальний посібник. Рівне: НУВГП, 156 с.
5. Pital P. et al. (2025). An integrated TOPSIS and ARAS multi-criteria decision-making approach for optimizing investment portfolios using goal programming and genetic algorithm model. Scientific Reports. 15, 34450. <https://doi.org/10.1038/s41598-025-17604-y>.
6. Guarino A. et al. (2024). EvoFolio: a portfolio optimization method based on multi-objective evolutionary algorithms. *Neural Computing & Applications*.

РЕКУРСИВНІ ЗАПИТИ В SQL: СЕМАНТИКА, УМОВИ ЗАВЕРШУВАНOSTІ ТА ОПТИМІЗАЦІЯ

Компан С.В., s.kompan@duikt.edu.ua

Державний університет інформаційно-комунікаційних технологій

Рекурсивні запити в SQL-подібних мовах є основним засобом опрацювання ієрархічних структур даних, однак питання їхньої формальної семантики, завершуваності та ефективності виконання залишаються відкритими. У дослідженні розвинено підходи, запропоновані у праці Д.Б. Буй та С.А. Полякова [1], і отримано результати, що уточнюють семантику Common Table Expressions (CTE) запитів та вдосконалюють методи їх оптимізації.

На основі аналізу композиційної семантики сформульовано узагальнену модель виконання рекурсії як послідовності таблиць, що стабілізуються при досягненні фіксованої точки. Визначено умови завершуваності запитів і показано, що їхню поведінку можна описати через монотонні відображення та замикання у просторі мультимножин. Окрему увагу приділено інтерпретації процесу рекурсії як поступового наближення до стаціонарного стану, у якому множина результатів перестає змінюватися. Такий підхід поєднує практичні механізми виконання запитів із теоретичними моделями обчислень і виявляє залежність між структурою запиту та його збіжністю. Наукова новизна роботи полягає у переосмисленні ролі ациклічності в рекурсивних SQL-виразах на основі класичної теорії фіксованих точок. Запропоноване трактування уточнює межі коректності рекурсії та має як теоретичне, так і практичне значення для розвитку ефективних методів оптимізації SQL-запитів.

1. Буй Д.Б., Поляков С.А. Рекурсивні запити в SQL-подібних мовах: приклади, змістова і формальна семантика. *Проблеми програмування*, № 2–3, 2010, с. 434–439.

ВИКОРИСТАННЯ R-АЛГОРИТМУ ДЛЯ ЗНАХОДЖЕННЯ ПАРАМЕТРІВ ДВОКЛАСОВОЇ SVM МОДЕЛІ

Корабльов М.М. m.m.korablov@gmail.com

Інститут кібернетики імені В.М. Глушкова НАН України

Позначимо через $D = \{(x_i, y_i): x_i \in \mathbb{R}^d, y_i \in \{-1, +1\} \forall i \in \overline{1, n}\}$ вибірку даних розміру n , розділену на тренувальну D_1 і тестову D_2 підвибірки розміром n_1 та n_2 , відповідно. Розв'язання задачі бінарної класифікації полягає у знаходженні такої функції $f: \mathbb{R}^d \rightarrow \{-1, +1\}$, яка гарно описуватиме тренувальні дані і при цьому мінімізує кількість помилок класифікації на нових, небачених до цього, даних [1].

На практиці для розв'язання цієї задачі навчання з учителем часто використовують метод опорних векторів (Support Vector Machines, SVM). Тренування даних моделей можна звести до розв'язання задач опуклого програмування або ж двоїстих до них [1]. Найвідомішими реалізаціями таких процедур є методи з бібліотек *libsvm* та *liblinear*, які включені до бібліотеки *scikit-learn* [2] мови програмування Python. Якщо ж підійти до обґрунтування SVM моделей з точки зору мінімізації емпіричного ризику та використати функцію втрат Хінджа [1], то їх тренування зводиться до мінімізації наступної опуклої негладкої функції (Hinge loss SVM):

$$\min_{w \in \mathbb{R}^d, b \in \mathbb{R}} \left\{ \frac{1}{2} \|w\|_2^2 + C \cdot \sum_{i=1}^{n_1} \max\{0, 1 - y_i \cdot (\langle w, x_i \rangle + b) \} \right\}. \quad (1)$$

Регуляризаційний параметр $C > 0$ дозволяє керувати тим, наскільки добре отримана модель $f(x) = \text{sign}(\langle w, x \rangle + b)$ буде описувати тренувальні дані, тим самим дозволяючи уникати перенавчання класифікатора [1].

Для розв'язання задачі (1) було розроблено та програмно реалізовано мовою Python алгоритм *ralgsvm*. В його основі лежить розроблена стохастична модифікація B-форми $r(\alpha)$ -алгоритму з адаптивним регулюванням кроку, в якому використовується незміщена статистична оцінка субградієнта, значення якої на кожній ітерації k обчислюється на

випадковій підвибірці $D_1^k \subseteq D_1$ розміром $m \in [1, n_1]$. При $m = n_1$ даний алгоритм співпадає з відомим детермінованим алгоритмом *ralgb5a* [3].

Порівняння роботи алгоритму *ralgsvm* та методів бібліотеки *scikit-learn* при $C = 1$ проводилося на задачі побудови бінарного класифікатора для набору даних *Breast Cancer Wisconsin (Diagnostic)* [2] розміром $n = 569$.

Табл. 1. Значення метрики Accuracy отриманих моделей

	<i>libsvm</i>	<i>liblinear</i>	<i>ralgsvm</i> ($m = 1$)
ACC_{train}	0.9912	0.9912	0.6352
ACC_{test}	0.9737	0.9737	0.6491
	<i>ralgsvm</i> ($m = 0.25n_1$)	<i>ralgsvm</i> ($m = 0.5n_1$)	<i>ralgsvm</i> ($m = n_1$)
ACC_{train}	0.9868	0.9890	0.9912
ACC_{test}	0.9825	0.9825	0.9737

Отримані результати показали, що алгоритм *ralgsvm*, завдяки використанню різних схем обчислення значення оцінки субградієнта, дозволив отримати модель класифікації, яка має кращу узагальнюючу спроможність та точніше класифікує тестові, небачені до цього, дані. А саме, при $m = 0.5n_1$ отримали модель з $ACC_{train} = 0.9890$ та $ACC_{test} = 0.9825$.

Подяки: Автор висловлює свою вдячність Стецюку П.І. за консультації під час проведення даних досліджень, які були підтримані проектом №2М-2024 (0124U002162) ДО «ВЦП КНУ імені Тараса Шевченка при НАН України» та НДР В.П.120.33 (0124U001061).

Список використаних джерел

1. Deisenroth M.P., Faisal A.A., Ong C.S. Mathematics for Machine Learning: textbook. Cambridge, 1st Edition. 2020.– 398 p.
2. Pedregosa et al. Scikit-learn: Machine Learning in Python, JMLR 12, pp. 2825-2830, 2011.
3. Стецюк П.І., Бєлих Т.В., Криворучко О.І. Теорія та програмні реалізації г-алгоритмів Шора, у Субградієнтні алгоритми та задачі на комбінаторних конфігураціях / Стецюк П.І., Донець Г.П., Ненахов Е.І. та ін.; за загал. ред. П.І. Стецюка. – Київ : Унів. вид-во ПУЛЬСАРИ, 2019. – 235 с.

ВИКОРИСТАННЯ МОДЕЛЕЙ ГЛИБОКОГО НАВЧАННЯ ДЛЯ КЛАСИФІКАЦІЇ РИТМУ СЕРЦЯ ЗА РЕЗУЛЬТАТАМИ ЕКГ

Костюченко А.Д., kostiuchenko_a@365.dnu.edu.ua

Герасимов В.В., herasymov_v@365.dnu.edu.ua

Дніпровський національний університет імені Олеся Гончара

У наш час моделі глибокого навчання дозволяють вирішувати значний перелік прикладних задач – від класифікації повідомлень електронної пошти до побудови складних інтелектуальних агентів. Окрема увага приділяється впровадженню моделей штучного інтелекту для покращення наявних методів діагностики в медицині. Одним із основних досліджень стану організму людини є ЕКГ (електрокардіограма), що дозволяє попередити прогресування різних серцево-судинних хвороб. Так як дослідження являє собою біомедичний часовий ряд, який відображає зміну електричних потенціалів серця у часі, доцільним є аналіз такого ряду та пошук аномальних станів за допомогою моделей глибокого навчання, а саме LSTM (англ. – Long Short-Term Memory) та GRU (Gated Recurrent Unit) [1].

Для навчання мереж було обрано відкритий набір даних з платформи Kaggle «ECG Heartbeat Categorization Dataset», який представлено 109446 навчальними прикладами, розподіленими на 5 класів. Мажоритарний клас, приклади якого відображають ритм здорового серця, складається з 72471 об'єктів для тренування, мінорний – 641, відповідно. Такий значний дисбаланс класів може призвести до ситуації, коли стандартне обчислення точності як основної міри якості прогнозування моделі не підходить для класифікації менш представлених класів [2].

Навчання моделей глибокого навчання проводилось із використанням фреймворку TensorFlow на мові програмування Python, у середовищі розробки Kaggle Notebook. Функцією втрат мереж було обрано багатокласову крос-ентропію. Результати навчання, які було отримання в ході проведення дослідження, подано у табл. 1 та табл. 2.

Таблиця 1 – Значення метрик якості класифікації моделі на основі GRU, %

GRU	Accuracy	Precision	Recall	F1-Score
train	98,66	95,92	88,14	91,5
test	97,99	93,13	84,94	88,39
validation	98,24	94,36	86,73	89,98

Відповідно до табл. 1, модель демонструє високу абсолютну точність, що наближається до 98 % на всіх вибірках. Такий результат зумовлений нерівномірним розподілом тренувальних прикладів між класами, внаслідок чого більшість прогнозів належать до мажоритарного класу, а поодинокі помилки класифікації не мають суттєвого впливу на загальне значення метрики Accuracy. Водночас показники Precision та Recall враховують співвідношення між правильністю передбачень та здатністю моделі виявляти усі релевантні приклади, знаходячись при цьому в діапазоні 85-95%.

Таблиця 2 – Значення метрик якості класифікації моделі на основі LSTM, %

LSTM	Accuracy	Precision	Recall	F ₁ -Score
train	95,47	83,38	71,26	75,89
test	95,19	81,99	70,7	74,94
validation	95,43	82,8	70,73	75,11

Модель LSTM також демонструє високі значення метрики Accuracy, однак Precision та Recall суттєво зменшились до рівня 70-83%, що вказує на гіршу узагальнюючу здатність моделі на основі LSTM у порівнянні з GRU архітектурою.

Список використаних джерел

- [1] A. Kjærø, E. S. Bugge, and C. B. Vennerød, "Long Short-term Memory RNN," arXiv preprint arXiv:2105.06756, 2024. [Online]. Available: <https://arxiv.org/pdf/2105.06756>
- [2] Shayan Fazeli, "ECG Heartbeat Categorization Dataset," Kaggle, 2020. [Online]. Available: <https://www.kaggle.com/datasets/shayanfazeli/heartbeat>

ЦИФРОВИЙ ЖИВОПИС ЯК ЕВОЛЮЦІЙНА ФОРМА ОБРАЗОТВОРЧОГО МИСТЕЦТВА В ЦИФРОВУ ЕПОХУ

Кочержинська А.Д.*, Сірик С.Ф.*, Лисиця Н.М.*, Шишканова Г.А.**

siryk600@gmail.com, lisitsa_natalya1971@ukr.net, shganna@ukr.net

**Дніпровський національний університет імені Олеся Гончара*

***Національний університет «Запорізька політехніка»*

У сучасному науковому дискурсі Digital Painting дедалі частіше розглядається не лише як форма образотворчого мистецтва, а й як приклад складної взаємодії між естетичними практиками, прикладною математикою та комп'ютерними технологіями. Його розвиток став можливим завдяки синтезу алгоритмічних методів обробки графіки, чисельного моделювання, а також інтерфейсних рішень, що реалізуються у спеціалізованих програмних середовищах. На відміну від загального поняття Digital Art, яке охоплює широкий спектр візуальних форм – від 3D-моделювання до інтерактивних інсталяцій – Digital Painting репрезентує напрям, що зберігає спадкоємність із класичними техніками станкового живопису, водночас інтегруючи математичні моделі для симуляції фізичних властивостей матеріалів, світлотіньових ефектів та динаміки кистей.

Digital Painting базується на використанні переважно двовимірних графічних редакторів, які реалізують складні алгоритми растрової та векторної графіки. Зокрема, моделювання мазків кисті здійснюється за допомогою математичних функцій, що описують тиск, нахил, текстуру та в'язкість фарби. Алгоритми згладжування, інтерполяції, фільтрації та симуляції освітлення забезпечують реалістичність зображення, а багат шарові структури дозволяють здійснювати нелінійне редагування композиції. З математичної точки зору, цифровий живопис є прикладом застосування чисельних методів (інтерполяція, апроксимація, згортки), теорії графів (структура шарів, зв'язність елементів), теорії кольору (моделі RGB, CMYK, HSV), а також алгоритмів машинного навчання.

На рис. 1 представлена серія з чотирьох зображень маків, де за літерою а) надано олійна картина (Кочержинська А.), б)-г) – цифрові зображення, створені за її мотивами. На рис.1 б) реалізовано техніку *chiaroscuro* – контрастне освітлення на темному тлі, що підкреслює об’ємність пелюсток. Використано пакети Corel Painter який дозволяє симулювати олійні мазки з високою текстурною деталізацією. Зображення в) створено у MyPaint, з акцентом на плавність ліній і природну динаміку кистей, а г) – за допомогою штучного інтелекту. Роботи демонструють послідовне ускладнення композиції: від статичної контрастної сцени до інтеграції квітів у світловий простір, що свідчить про застосування математичних моделей згортки, інтерполяції та симуляції текстур у цифровому середовищі.



а)



б)



в)



г)

Таким чином, можна сказати, що Digital Painting постає не лише як технічна інновація, але як культурне явище, що репрезентує синтез традиційної художньої спадщини з цифровими технологіями.

АНАЛІЗ ВАРІАБЕЛЬНОСТІ СЕРЦЕВОГО РИТМУ

Кошель В.В., viktoriakoshel05@gmail.com

Єгошкін Д.І., KnightDanila@i.ua

Сірик С.Ф., siryk600@gmail.com

Дніпровський національний університет імені Олеся Гончара

На сьогодні одним із найбільш поширених способів оцінювання стану адаптації організму є аналіз варіабельності серцевого ритму. Результати численних фізіологічних і клінічних досліджень підтвердили, що комплексні показники серцево-судинної системи можуть слугувати індикатором адаптивних реакцій організму, а також предиктором ризику розвитку захворювань. І донині вчені продовжують розвивати цю тему для більш глибокого розуміння причинно-наслідкових зв'язків між серцевим ритмом та майбутніми захворюваннями. Наприклад, в роботі [1] було досліджено вплив терапії місцевими анестетиками на варіабельність серцевого ритму.

Оскільки варіабельність серцевого ритму відображає змінність інтервалів між серцевими скороченнями, існує багато способів отримати серцебиття людини. Методи в часовій області базуються на прямому аналізі послідовності RR-інтервалів, як от SDNN та RMSSD. Вони прості у розрахунку і найчастіше застосовуються на практиці. Також існують методи нелінійні, як Sample Entropy, які можуть виміряти хаотичність і мета таких методів полягає в з'ясовуванні схожості фрагментів сигналів. Отримані результати можуть бути використані для розробки систем ранньої діагностики серцево-судинних порушень

Здоровій людині притаманно мати не ідеальний ритм, він повинен мати деяку хаотичність. Однак, якщо ритм стає надто регулярним, тобто набуває меншої хаотичності, то є підозра на певну патологію.

Для вимірювання варіабельності серцевого ритму, можна використовувати різні бази даних із зібраними сигналами ЕКГ. У роботі [2]

розглянуто, що для SDNN краще використовувати сигнали, що були виміряні впродовж 24 годин, адже цей метод потребує великих інтервалів для валідної оцінки. Для RMSSD достатньо використати сигнал, який було виміряно впродовж 5 хвилин. Отримані дані можуть бути використані для навчання моделей автоматичної оцінки стану серцево-судинної системи або для виявлення відхилень у роботі серця.

Таким чином, аналіз варіабельності серцевого ритму є актуальною задачею сучасної біомедицини, оскільки дозволяє оцінити стан регуляторних систем організму та виявити ранні ознаки патологій. Цей підхід пропонує неінвазивний і водночас інформативний спосіб моніторингу серцево-судинного здоров'я, що робить його перспективним для подальших клінічних і дослідницьких застосувань.

Бібліографічні посилання:

1. Effects of therapy with local anesthetics (TLA) on heart rate variability (HRV) over 24 hours [Електронний ресурс] // ResearchGate. – 2024. – Режим доступу:
https://www.researchgate.net/publication/395971491_Effects_of_therapy_with_local_anesthetics_TLA_on_heart_rate_variability_HRV_over_24_hours
2. Shaffer F., Ginsberg J. (2017). An overview of heart rate variability metrics and norms [Електронний ресурс]. – PMC. – Режим доступу:
<https://pmc.ncbi.nlm.nih.gov/articles/PMC5624990/>

ЗАСТОСУВАННЯ НЕЙРОННИХ МЕРЕЖ ДЛЯ ПРОГНОЗУВАННЯ ВАРТІСНИХ ПОКАЗНИКІВ ОБ'ЄКТІВ

Красюк М.В., aweyer42@gmail.com

Наконечна Т.В., naktanya@ukr.net

Дніпровський національний університет імені О. Гончара

Прогнозування вартісних показників – ціни акцій, товарів, електроенергії, криптовалют, тощо – є критично важливим завданням для фінансового сектора, енергетики, торгівлі, інвестицій та управління ризиками. Традиційні статистичні моделі (AR, MA, ARIMA, експоненційне згладжування і подібні) добре працюють за умов стабільності, лінійності, низької волатильності та за наявності великих історичних даних, але часто виявляються недостатніми: через нелінійність часових рядів, через високі коливання (волатильність), через шум, сплайси, різні зовнішні шоки, через мультифакторність і залежність від зовнішніх факторів (мінітрендів, індикаторів, макроекономіки тощо).

Нейронні мережі й методи глибокого навчання дають змогу моделювати складні нелінійні залежності, витягувати ознаки з великих масивів даних, обробляти часові залежності високого порядку, і, за деяких умов, підвищити точність прогнозів. Інтерес до застосування нейронних мереж на вирішення прикладних завдань прогнозування почав активно зростати у 80-90-ті роки XX століття і залишається актуальним нині.

Автор статті [1] порівняв результати застосування штучної нейронної мережі та моделі множинної лінійної регресії для моделювання структури капіталу в економічно розвинених країнах і дійшли висновку, що навчена штучна нейронна мережа виявилася ефективнішою завдяки більшій чутливості до зміни показників. Оцінка досліджень щодо ефективності нейронних мереж у вирішенні завдань прогнозування, описаних у статті [2] показала, що у 19 із 22 (86%) досліджених випадках нейронні мережі показали більш високі результати порівняно з іншими моделями.

Отже, мета нашого дослідження це аналіз ефективності архітектури нейронних мереж для прогнозування вартісних показників, у порівнянні з традиційними статистичними моделями, а також розробка й тестування роботи програмної моделі нейронної мережі для наглядного підтвердження кращої продуктивності моделей глибокого навчання порівняно з традиційними статистичними методами у прогнозуванні цін.

Для розв'язання поставленої задачі була розроблена рекурентна нейронна мережа. Для розуміння роботи нейронної мережі була побудована математична модель. Але, моделювання таких математичних обчислень в більш складних мережах глибокого навчання може бути дуже складною та трудомісткою роботою, а програмний код занадто громіздким. Через це були застосована програмна бібліотека, TensorFlow а також її надбудова Keras.

Була протестована робота побудованої рекурентної нейронної мережі. Мережа навчалась на випадково згенерованій вибірці даних, що представляє собою тренд вартісних показників за деякий час. Вибірка була поділена на тренувальну, затверджувальну та випробувальну. Функція втрат – середньоквадратична похибка. Після навчання мережі можемо побачити прогноз вартісних показників та порівняти його з реальними даними, що було відображено на відповідному графіку. Як можна побачити рекурентна нейронна мережа доволі добре впоралася з задачею прогнозування. Були розглянуті три найпопулярніші для рекурентних нейронних мереж функції активації. Програма показала, що найефективнішою функцією є reLU.

Таким чином, можна зробити висновок, що використання методів глибокого навчання дійсно підвищує точність прогнозування та може суттєво сприяти покращенню ринкової ефективності та фінансових показників компаній.

Бібліографічні посилання

1. Hsiao-Tien Pao, A comparison of neural network and multiple regression analysis in modeling capital structure. Hsiao-Tien Pao. Expert Systems with Applications. – 2008. – № 35. – P.720-727.
2. Adya, M. How Effective are Neural Networks at Forecasting and Prediction? A Review and Evaluation. M. Adya, F. Collopy. Journal of Forecasting. – 1998. – Vol. 17. – P.481-495.

ПОРІВНЯЛЬНИЙ АНАЛІЗ ЕВРИСТИЧНИХ АЛГОРИТМІВ РОЗВ'ЯЗАННЯ ЗАДАЧІ ОПТИМАЛЬНОГО ПРОЄКТУВАННЯ РЕЗЕРВУАРІВ ПІД ТИСКОМ

Кривцов О.С., kryvtsov.o21@365.dnu.edu.ua

Шевельова А.Є., shevelevaae@dnu.dp.ua

Дніпровський національний університет імені Олеся Гончара (www.dnu.dp.ua)

Зростання складності інженерних систем зумовлює потребу в ефективних методах пошуку оптимальних конструкцій. Класичні чисельні методи не завжди забезпечують знаходження глобального мінімуму, тому сучасні дослідження спрямовані на використання евристичних алгоритмів, здатних ефективно працювати в умовах нелінійності та багатьох обмежень [1].

Проблема оптимального проєктування резервуарів під тиском має значне промислове та економічне значення, оскільки впливає на ефективність використання матеріалів, собівартість виробництва й безпечність експлуатації обладнання [2]. Застосування метаввристичних алгоритмів дає змогу уникнути обмежень класичних методів, що робить їх перспективними для інженерних задач оптимізації.

Розроблено математичну модель задачі, у якій цільова функція відображає собівартість виготовлення резервуара залежно від товщини оболонок, радіуса та довжини циліндричної частини. Для знаходження оптимального розв'язку реалізовано метаввристичні методи: імітації відпалу (Simulated Annealing, SA), метод зозулі (Cuckoo Search, CS), оптимізація рою частинок (Particle Swarm Optimization, PSO), алгоритм світлячків (Firefly Algorithm, FA), алгоритм запилення квітів (Flower Pollination Algorithm, FPA), гравітаційний алгоритм пошуку (Gravitational Search Algorithm, GSA) та метод Хука–Дживса (Hooke–Jeeves, PS) [3].

Програмна реалізація виконана мовою C++ у середовищі WinUI 3, що забезпечує сучасний графічний інтерфейс і дозволяє користувачу

змінювати параметри задачі, проводити експерименти та порівнювати результати методів у єдиній системі.

Проведено серію обчислювальних експериментів для оцінки точності, швидкодії та стабільності алгоритмів. Встановлено, що найкраще співвідношення між швидкістю збіжності та точністю розв'язку показали алгоритми PSO та Firefly Algorithm, тоді як Simulated Annealing продемонстрував високу стабільність результатів при повторних запусках.

У роботі визначено доцільність застосування евристичних методів для інженерних задач оптимального проєктування. Аргументовано ефективність підходів, що поєднують глобальний та локальний пошук. Створене програмне забезпечення може бути використане як основа для подальших наукових досліджень і впровадження нових оптимізаційних методів у прикладній механіці та машинобудуванні.

Бібліографічні посилання

1. Al-Milli N. (2014). Hybrid Genetic Algorithm with great deluge to solve constrained Optimization Problems. *Journal of Theoretical and Applied Information Technology*, 59(2), P. 385–389. <https://www.jatit.org/volumes/Vol59No2/19Vol59No2.pdf>
2. Hassan S., Kumar K.K., Deva Raj C., Sridhar K. (2013). Design and Optimization of Pressure Vessel Using Metaheuristic Approach. *Applied Mechanics and Materials*, 465–466, P. 401–406. <https://doi.org/10.4028/www.scientific.net/AMM.465-466.401>
3. Harish Garg. (2014). Solving Structural Engineering Design Optimization Problems using an Artificial Bee Colony Algorithm. *Journal of Industrial and Management Optimization*, 2014, 10(3): 777-794. <https://doi.org/10.3934/jimo.2014.10.777>

МАТЕМАТИЧНЕ МОДЕЛЮВАННЯ ЗАДАЧ ОПТИМАЛЬНОГО РОЗБИТТЯ ТА ПЛАНУВАННЯ В УМОВАХ НЕВИЗНАЧЕНОСТІ

Кузенков О.О., Козакова Н.Л., Григоренко О., Балейко Н.В.

Дніпровський національний університет імені Олеся Гончара (www.dnu.dp.ua)

Дослідження присвячене розробці методів для різних класів математичних моделей, які переважно застосовуються до задач оптимального розбиття множин різної природи й наповнення. Пошук розв'язків задач ОРМ із урахуванням змінних обмежень дає змогу не лише знайти оптимальне розміщення підмножин, але й відповісти на суміжні практично важливі питання: яка мінімально необхідна кількість підмножин, які оптимальні координати центрів, а також як мінімізувати сумарні витрати у відповідній задачі оптимізації.

Значний внесок у розвиток теорії ОРМ у різних формулюваннях та з різними критеріями пошуку оптимального розв'язку зробили вчені наукової школи доктора фізико-математичних наук, професора Кисельової О.М. Основні положення та прикладні дослідження з цієї проблематики викладені в працях Кисельової О.М. у співавторстві з Шором Н.З., Притомановою О.М., Шевельовою А.Є., Гарт Л.Л., Ус С.А. та ін.

Метою роботи було розроблення нових підходів до дослідження та розв'язання різних класів динамічних задач оптимального розбиття множин, зокрема з обмеженнями та без обмежень, однопродуктових і багатопродуктових, з заданими та невідомими координатами центрів підмножин, їх адаптація до конкретних прикладних задач та узагальнення умов застосування цих підходів у практичних ситуаціях.

Розв'язання задачі в математичній постановці дало змогу отримати прикладні результати для низки проблем, що ґрунтуються на вже наявних задачах і можуть бути зведені до задач розміщення-розбиття. Мова йде, зокрема, про такі прикладні задачі, як розміщення центрів управління протиповітряною обороною, побудова оптимальної траєкторії 3-D друку, математичне моделювання поширення інфекційних захворювань з розміщенням центрів вакцинації та організацією допомоги постраждалим.

АНАЛІЗ ТА КЕРУВАННЯ БІФУРКАЦІЯМИ В МОДЕЛЯХ ПРИРОДНИХ ПРОЦЕСІВ РОЗМІЩЕННЯ-РОЗБИТТЯ

Кузенков О.О., Козакова Н.Л., Лупинський С., Балеєнко Н.В.

Дніпровський національний університет імені Олеся Гончара (www.dnu.dp.ua)

У загальній формі, теорія біфуркацій – це теорія знаходження точок рівноваги нелінійних рівнянь, де під точками рівноваги розуміють стаціонарні, періодичні та квазіперіодичні розв'язки. Використання функціонального аналізу в задачах теорії біфуркацій насамперед пов'язане з обґрунтуванням можливості спрощення задач високої та навіть нескінченної розмірності до одновимірних, двовимірних або тривимірних задач. Привабливість задач малої розмірності головним чином зумовлена можливістю проекції задач на простір власних функцій та простотою аналітичного розв'язання поліноміальних, диференціальних та інтегральних підзадач у контексті теорії біфуркацій. Під час дослідження якісних властивостей систем диференціальних рівнянь, однією з основних проблем є проблема нелінійності інтегральних кривих, коли параметри систем проходять через певні критичні значення. Дослідження та прикладні інтерпретації катастрофічних біфуркацій системи можуть сформулювати чисельні та аналітичні умови, що визначають якісні зміни властивостей систем і сприяють розумінню сутності цих змін.

Математична модель, що розглядається у даній роботі, формулюється наступним чином:

$$\frac{dx_j}{dt} = \sum_{i=1}^n A_{ji} \cdot a_i \cdot \left(1 - \frac{1}{K} \sum_{l=1}^n x_l\right) x_i, \quad j = \overline{1, n} \quad (1),$$

де x_j – кількість j -тих груп, $f_i(x)$ – функція, що описує загальні можливості i -ої групи до зростання, A_{ji} – доля зростання i -ої групи, що припадає на j -ту підгрупу, a_i – відображає спроможності групи з індексом i до зростання, K – ємність системи. Припускаємо, що $\sum_j A_{ij} = 1$ для будь-якого i . З формули (1) очевидно, що динаміка групи затухає у випадку, коли x_i прямує до нуля або $\sum_{l=1}^n x_l$ наближається до ємності системи K .

МЕТОДИ РОЗВ'ЯЗАННЯ ДИНАМІЧНИХ ЗАДАЧ ОПТИМАЛЬНОГО РОЗМІЩЕННЯ-РОЗБИТТЯ

Кузенков О.О., Козакова Н.Л., Шведов В., Сірик С.Ф.

Дніпровський національний університет імені Олеся Гончара (www.dnu.dp.ua)

Актуальність теми зумовлена необхідністю розвитку наявного математичного апарату теорії оптимального розбиття множин у задачах, що характеризуються динамічністю факторів, які в класичній теорії вважаються сталими. Цей підхід є актуальним у широкому переліку прикладних задач, пов'язаних із розміщенням мережі магазинів, стільникових станцій, систем громадського сповіщення, візуального та звукового контролю тощо. Актуальність цього напрямку підкреслювалася в роботах багатьох авторів, однак цілісний підхід не був запропонований.

Одними з найвагоміших робіт, у яких узагальнено вдосконалення теорії оптимального розбиття множин, є праці таких авторів, як О.М. Кисельова та Л.С. Коряшкіна [1–2], у яких, зокрема, приділено особливу увагу питанням математичної теорії оптимального розбиття множин n -вимірному евклідовому простору, які належать до некласичних задач нескінченновимірною математичного програмування. У зазначених роботах особливий акцент зроблено на задачах, які характеризуються нелінійністю критеріїв оптимальності, наявністю додаткових обмежень та варіативністю параметрів задачі. Також у роботах приділено увагу застосуванню теорії оптимального розбиття множин до розподілених систем і специфіці розв'язків, що виникають при застосуванні цієї теорії до систем параболічного типу. Окрім того, розглядається багато прикладів неперервних задач теорії оптимального розбиття множин та задач оптимального покриття, геометричного проектування тощо. У роботах сформульовано коло задач для подальших досліджень цього класу математичних проблем.

Бібліографічні посилання:

1. Kiseleva E.M., Koryashkina L.S. Models and methods of solving continuous problems of optimal partitioning of sets: linear, nonlinear, dynamic problems: monograph. K.: Naukova dumka, 2013. 606 p.
2. Kiseleva E.M., Koryashkina L.S. Continuous problems of optimal partitioning of sets and r-algorithms: monograph K.: Naukova Dumka, 2015. 400 p.

ДИНАМІЧНІ ЗАДАЧІ ОПТИМАЛЬНОГО РОЗБИТТЯ МНОЖИН ЗІ ЗМІННИМ ПОПИТОМ

Кузенков О.О., Лозовський А.В.

Oles Honchar Dnipro National University (www.dnu.dp.ua)

Оптимальне розбиття множини - класична задача математичного програмування та комбінаторної оптимізації, що знаходить широке застосування в різних галузях науки й техніки. Особливу актуальність становить розв'язування динамічних задач оптимального розбиття, коли вхідні дані, параметри чи умови задачі змінюються з часом, ускладнюючи традиційні методи розв'язання та вимагаючи застосування спеціалізованих моделей і алгоритмів. Важливою підзадачею в цій сфері є оптимальне розбиття множини з фіксованими центрами, за якої центри підмножин задані наперед і не змінюються, а завдання полягає в визначенні розподілу елементів між ними з мінімізацією певного цільового функціоналу.

Динамічний підхід до розв'язання задачі дає змогу враховувати змінні умови, такі як коливання витрат на транспортування, попиту або інші зовнішні фактори, що впливають на структуру оптимального розбиття. У цій роботі використано математичну модель, раніше розроблену авторами, яка формалізує динамічну задачу оптимального розбиття множини з фіксованими центрами та дозволяє адаптивно перерозподіляти елементи у відповідь на зміни вхідних параметрів.

Такий клас задач виникає у різних прикладних контекстах - від логістики до систем спостереження та управління, - де необхідно оптимально розподілити елементи множини між наперед заданими центрами з урахуванням змін, що відбуваються з часом. Запропонована математична модель дозволяє враховувати динаміку вхідних даних для побудови адаптивної структури розбиття. Застосовані чисельні методи забезпечують ефективність розв'язання задач великої розмірності у реальному часі.

Постановка задачі. Знайти розбиття $\varpi = \{\Omega_1, \dots, \Omega_N\} \in \Sigma_\Omega^N$ множини $\Omega \subset E_n$ і векторну функцію $\rho(x, t)$, визначену майже всюди (м.в.) для $x \in \Omega$ та всіх $t \in [0, T]$, які забезпечують

$$\inf_{\varpi \in \Sigma_\Omega^N; \rho(x, t) \in L_2^N(\Omega \times [0, T])} F(\varpi, \rho(x, t)),$$

де

$$F(\varpi, \rho(x, t)) = \int_0^T \sum_{i=1}^N \int_{\Omega_i} (c_i(x, \tau_i) + a_i) \rho(x, t) dx dt$$

за умов

$$\begin{aligned} \frac{\partial \rho(x, t)}{\partial t} &= f(\rho(x, t)), 0 \leq t \leq T \\ \rho(x, t_0) &= \rho_0(x), i = 1, \dots, N \end{aligned}$$

Тут $\rho(x, t)$ - шукана дійсна функція, визначена на $\Omega \times [0, T]$, яка неперервно диференційована за аргументом t на відрізку $[0, T]$ м.в. для $x = (x^{(1)}, \dots, x^{(n)}) \in \Omega$, обмежена та вимірна за аргументом x на Ω для всіх $t \in [0, T]$; $\rho_0(x)$ - задана невід'ємна функція, обмежена та вимірна за аргументом $x \in \Omega$; $f(\rho(x, t))$, - задана дійсна функція, неперервна та ліпшицева в межах свого визначення; $c_i(x, \tau_i)$ - визначені на $\Omega \times \Omega$ задані дійсні обмежені функції, вимірні за аргументом $x = (x^{(1)}, \dots, x^{(n)})$ при будь-якому фіксованому $\tau_i = (\tau_i^{(1)}, \dots, \tau_i^{(n)}) \in \Omega_i \subset \Omega$; $a_i, i = 1, \dots, N$, - задані, переважно, невід'ємні числа; $T > 0$ та $t_0 \in [0, T]$ задані числа.

ДО РОЗРОБКИ АЛГОРИТМІВ КЛАСИФІКАЦІЇ ЖЕСТОВОЇ МОВИ НА ОСНОВІ ПСЕВДООБЕРНЕННЯ І ГРУПУВАННЯ ОЗНАК

Кузнєцов В.О.^{1,2}, Крак Ю.В.^{1,2} (Iurii.krak@knu.ua), Куляс А.І.¹, Кудін Г.І.¹

¹ Інститут кібернетики імені В.М. Глушкова НАН України

² Київський національний університет імені Тараса Шевченка

Швидкий розвиток систем спілкування користувача з комп'ютером, призвів до необхідності обробки альтернативних засобів введення інформації, зокрема жестів, та її розпізнавання із прийнятною швидкістю.

Це вимагає певних алгоритмів аналізу даних, їх подання їх у зручній для обробки формі, та виділення характерних ознак. Існує декілька підходів до аналізу жестів – структурний, на основі аналізу особливостей і на основі вихідного графічного зображення. В даній роботі пропонується дослідити структурний підхід, оскільки він має можливості візуальної верифікації коректного просторового розташування структурних ознак.

З цією метою було проаналізовано набір типових конфігурацій руки у дактильній мові і отримано набір даних із 400 унікальних зображень руки (рис 1), з яким можна ознайомитися на GitHub за посиланням [1].



Рис 1. Зображення ручних жестів А, Г, Х із вибірки зразків

Для обробки даних було використано бібліотеку MediaPipe та OpenCV, що призначені для видобутку структурних ознак руки та обробки зображення. Отримані метадані зберігалися і надалі завантажувалися під час навчання методів класифікації за допомогою бібліотеки scikit-learn.

Для оцінки ефективності ознак було застосовано методи групування і кластеризації [1]. Для цього аналізувалася кореляція ознак і величина подібності ознак на основі спектру власних значень матриці кореляції.

На основі аналізу ознак, було відібрано набори із 50, 80, та 100 ознак та проаналізовано їх ефективність на декількох наборах даних, що містили 5, 10, 14 та 24 класи типових конфігурацій руки відповідно.

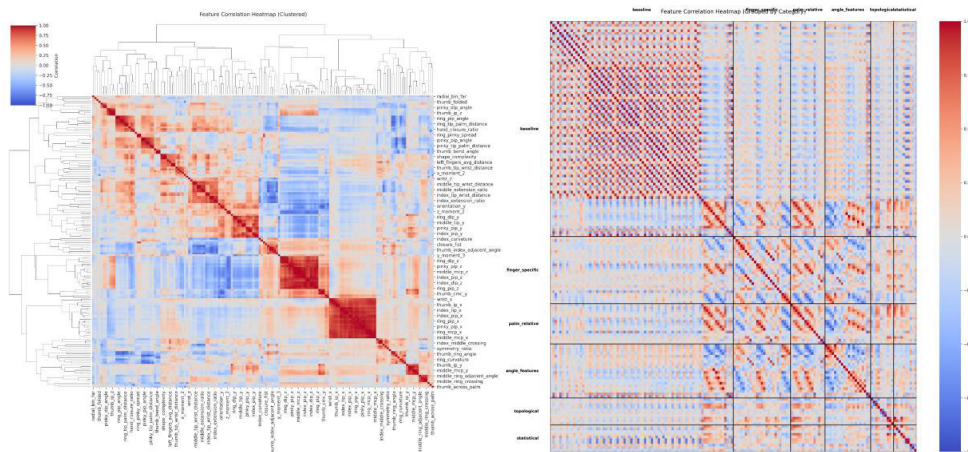


Рис 2. Візуалізація матриці кореляції і результат групування за подібністю

В результаті аналізу різних методів, було відібрано 2 методи, які показали найвищу точність розпізнавання в декількох експериментах – логістична регресія і дерева рішень із точністю 90,53% і 95% відповідно.

В подальших дослідженнях пропонується дослідити розширений набір даних і інші методи класифікації і перетворення розмірності ознак.

Бібліографічні посилання

1. Репозиторій github із набором даних: <https://github.com/Icyb/usl-gesture-dataset>
2. Krak, I., Kudin, H., Efremov, M., Samoylov, A., Kuznetsov, V., Amirgaliyev, Y., Kasianiuk, V. (2021). Hyperplane Clasterization of the Small Data Based on Pseudo-Inverse and Projective Matrices. CEUR Workshop Proceedings (с. 137–145). <https://ceur-ws.org/Vol-3101/Short10.pdf>
3. Kondratiuk, S., Krak, I., Kulas, A., Kuznetsov, V. (2021). Three-dimensional convolutions and temporal data for sign language recognition. CEUR Workshop Proceedings (с. 220–228). http://ceur-ws.org/Vol-3132/short_2.pdf.

МЕТОДИ УМОВНОЇ ОПТИМІЗАЦІЇ У МОДЕЛЯХ РОЗШАРУВАННЯ ГІРСЬКИХ ПОРІД

Кузьменко В.І., vasilkuzmenko50@gmail.com

Дніпровський національний університет імені Олеся Гончара

Розглядаються технологічні проблеми розшарування порід під лією рідини або газу високого тиску. Матеріал шарів вважається лнійно-пружним. На межі шарів діють сили тиску рідини або нестисливою рідини. Зміна тиску призводить до зміни зазору між шарами. Задачі такого класу виникають при аналізі процесів видобутку рідких або газоподібних копалин. Відзначимо сучасні технології розриву пластів рідиною висоого тиску.

Сформульована задача лінійної теорії пружності з умовами у вигляді нерівностей. Дослідження ґрунтується на варіаційному формулюванні. Виникає задача пошуку мінімуму квадратичного функціоналу на множини допустимих переміщень. Дискретизація варіаційної задачі здійснювалася з використанням методу скінченних елементів. Розв'язання задачі квадратичного програмування високої вимірності здійснювалося з використанням модифікованого алгоритму методу послідовної верхньої релаксації, що враховує особливості задачі. Досліджено зміну напружено-деформованого стану у процесах закачування або відкачування рідини.

Запронований підхід може бути використаний при розробці технологій видобутку корисних копалин методами розриву пластів рідиною високого тиску.

СЕМАНТИЧНИЙ АНАЛІЗ РЕКЛАМНИХ ВІДЕО ЗА ДОПОМОГОЮ ВЕЛИКИХ МОВНИХ МОДЕЛЕЙ

Куліш С. П. kulish@vntu.edu.ua, **Ткаченко О. М.** ant@vntu.edu.ua

Вінницький Національний Технічний Університет (www.vntu.edu.ua)

Стрімке зростання обсягів онлайн-відеореклами зумовлює актуальність автоматизованого семантичного аналізу її контенту. Ручна обробка рекламних відео є надзвичайно ресурсозатратною, що підкреслює потребу в інтелектуальних системах, здатних точно інтерпретувати зміст і маркетингове навантаження таких матеріалів.

Великі мовні моделі (LLM) відкривають нові можливості для аналізу відео, оскільки здатні працювати з мультимодальною інформацією — візуальним рядом, мовленням і текстовими елементами. На відміну від традиційного покадрового тегування або розпізнавання об'єктів, LLM забезпечують глибше розуміння сюжетної структури та контексту, дозволяючи виявляти цілісний сенс і приховані маркетингові меседжі.

У пропонованому підході використовується попередній аналіз відео засобами комп'ютерного зору для виділення ключових об'єктів, сцен і транскрипції мовлення. Ці структуровані дані подаються на вхід LLM, яка формує узагальнене семантичне представлення відео: визначає тематику, категорію продукту, наявність закликів до дії (СТА), промо-елементів тощо. Наприклад, модель може виявити фрази типу «купити зараз», наявність логотипів бренду, промоакцій або особливих пропозицій — незалежно від того, в якому вигляді вони подані (текст на екрані чи мовлення диктора).

На відміну від підходів, що розглядають відео як сукупність розрізнених елементів, LLM враховують послідовність подій у часі, емоційний фон і стиль подачі, що дає змогу точніше інтерпретувати мету

та характер реклами. Запропонована архітектура була реалізована у вигляді прототипу, який складається з модуля попереднього відеоаналізу та модуля мовної моделі. Система генерує детальні анотації для рекламних відео, що можуть бути використані для подальшої класифікації, фільтрації або оцінки ефективності креативу.

Таким чином, поєднання можливостей комп'ютерного зору з великими мовними моделями відкриває нові горизонти в автоматичному аналізі відеоконтенту, особливо в рекламній сфері. У перспективі можливий перехід до повністю відео-орієнтованих LLM, здатних безпосередньо аналізувати сюжет як єдину структуру з урахуванням просторово-часових зв'язків.

АЛГОРИТМИ РОЗВ'ЯЗАННЯ СЛАЙДІНГ-ПАЗЛІВ

Лавушник Я.П., lavushnyk.y22@365.dnu.edu.ua, Степанова Н.І.

Дніпровський національний університет імені Олеся Гончара

Задача розв'язання слайдінг-пазлів (Sliding Puzzle) є класичною комбінаційною проблемою, що полягає у знаходженні оптимальної послідовності дій для переведення у обмеженому ігровому полі набору пронумерованих плиток з початкового розташування у задане кінцеве. Такі задачі широко застосовуються для дослідження пошукових алгоритмів, оскільки поєднують у собі обмежений простір можливих станів, необхідність побудови ефективних стратегій пошуку та можливість порівняння різних методів за швидкістю та точністю.

Дана робота досліджує базові алгоритми розв'язку поставленої задачі, зокрема методи пошуку в ширину (BFS), пошуку в глибину (DFS) та евристичний алгоритм A^* . Для кожного із зазначених алгоритмів розглянуто принцип роботи, особливості застосування до слайдінг-пазлів і проаналізовано ефективність на прикладі конкретних конфігурацій.

Вибір саме цих методів зумовлений тим, що вони представляють різні підходи до розв'язання задач пошуку: вичерпний, рекурсивний та евристичний. Це дозволяє дослідити вплив стратегії пошуку на швидкість, обсяг пам'яті та якість знайденого рішення у задачі розв'язання слайдінг-пазлів.

- BFS — проходить усі можливі стани послідовно, рівень за рівнем, поки не знайде розв'язок.
 - Перевага: завжди знаходить найкоротший шлях.
 - Недолік: займає багато пам'яті та часу.
- DFS — рекурсивно переходить до наступних ходів, доки не дійде до цілі або до «тупикового положення».

- Перевага: потребує менше пам'яті.
- Недолік: може піти в неправильну гілку і не знайти оптимального шляху.
- A^* — оцінює кожен стан за формулою:
$$f(n) = g(n) + h(n),$$

де $g(n)$ — ціна шляху від початку, $h(n)$ — евристична оцінка «відстані» до мети.

- Перевага: працює набагато швидше, бо не перевіряє всі стани.
- Недолік: ефективність залежить від правильно обраної евристики.

У даній роботі розроблено програмний інструмент, який демонструє процес розв'язання слайдінг-пазлу за допомогою алгоритму BFS. Програма дозволяє наочно відстежувати послідовність ходів і аналізувати ефективність обраного підходу. У подальшому планується також виконати програмну реалізацію двох інших зазначених алгоритмів з метою дослідження особливостей застосуванні і оцінки ефективності.

Слайдінг-пазл – це не просто іграшка для розваги або розвитку логічного мислення, це насамперед модель, яка має широке застосування у різних прикладних задачах. На таких моделях відбувається тренування евристик у галузі штучного інтелекту, вони використовуються у нейропсихології для оцінки просторового мислення, пам'яті, здатності до планування. Це підтверджує актуальність і практичну цінність даної роботи.

Список використаних джерел

1. Algoua. «Пошук в ширину | Algoua» [Електронний ресурс]. – Режим доступу: <https://algoua.com/algorithms/graphs/bfs/>.
2. Yingpeng, Y. *A Comparative Analysis of Depth-First Search (DFS) and Breadth-First Search (BFS)* [Електронний ресурс]. – Режим доступу: <https://medium.com/@yingpeng0221/a-comparative-analysis-of-depth-first-search-dfs-and-breadth-first-search-bfs-a2d830794479>.
3. Горчинський, М.І. (або інші автори). «...» [Електронний ресурс]. – М.І. Горчинський та ін. – Мит2023-... [PDF]. – Режим доступу: <http://dspace.opu.ua/jspui/bitstream/123456789/14199/1/MIT2023-Горчинський.pdf>.

СТІЙКІСТЬ ТА РЕГУЛЯРИЗАЦІЯ ЗА ОБМЕЖЕННЯМИ ВЕКТОРНИХ ЗАДАЧ ЦІЛОЧИСЛОВОЇ ОПТИМІЗАЦІЇ

Лебедєва Т.Т., Семенова Н.В., Сергієнко Т.І.

lebedevatt@gmail.com, nvsemenova@meta.ua, taniaser62@gmail.com

Інститут кібернетики імені В.М. Глушкова НАН України

Досліджена векторна задача оптимізації вигляду:

$Q_P(F^u, X^u): \max \{F(x) | x \in X\}$, де $F(x) = (f_1(x), \dots, f_\ell(x))$, $\ell \geq 2$, $X \neq \emptyset$,

$X = G \cap Z^n$, $f_i: R^n \rightarrow R, i \in N_\ell$, $Z^n \subset R^n$ – множина цілочислових векторів,

$G = \{x \in R^n | \langle B_i x, x \rangle + \langle p_i, x \rangle + h_i \leq 0 \leq 0, i \in N_m\}$ – обмежена множина,

$p_i \in R^n$, $h_i \in R^1$, $B_i \in R^{n \times n}$, B_i – симетричні невід’ємно визначені матриці, $i \in N_m$. Множиною розв’язків задачі $Q_P(F, X)$ ($Q_{S\ell}(F, X)$) вважаємо

множину Парето-оптимальних розв’язків $P(F, X) = \{y \in X | \pi(x, F, X) = \emptyset\}$

$\pi(x, F, X) = \{y \in X | F(y) \geq F(x), F(y) \neq F(x)\}$ (множину Слейтера

$S\ell(F, X) = \{y \in X | \sigma(x, F, X) = \emptyset\}$, $\sigma(x, F, X) = \{y \in X | F(y) > F(x)\}$).

Під дослідженням стійкості векторної задачі $Q_P(F^u, X^u)$ будемо розуміти вивчення поведінки множини її Парето-оптимальних розв’язків при невеликих збуреннях у вхідних даних, що стосуються обмежень задачі. Позначимо $B = (B_1, \dots, B_m)$, $p = (p_1, \dots, p_m) \in R^{m \times n}$, $h = (h_1, \dots, h_m)$.

Вхідні дані обмежень $u_2 = (B, p, h) \in U_2 \subset R^{n \times n \times m} \times R^{m \times n} \times R^m$.

Визначимо для довільного $\delta > 0$ множину $O_\delta(u_2) = \{u_2(\delta) \in U_2 | \max_{i \in N_m}$

$\|B_i(\delta) - B_i\| < \delta, \|p(\delta) - p\| < \delta, \|h(\delta) - h\| < \delta\}$, й $\forall \delta > 0, u_2(\delta) \in O_\delta(u_2)$ розглянемо

збурену допустиму множину

$$X_{u_2(\delta)} = \{x \in Z^n | \langle B_i(\delta)x, x \rangle + \langle p_i(\delta), x \rangle + h_i(\delta) \leq 0, i \in N_m\}.$$

Отримані необхідні й достатні умови різних $T_1, \dots, T_5$ – типів стійкості за критеріями та за обмеженнями [1–3] свідчать про те, що якими б незначними не були збурення вхідних даних, множини $P(F_{u_1(\delta)}, X)$, $P(F_{u_1(\delta)}, X_{u_2(\delta)})$ і $P(F, X_{u_2(\delta)})$ у загальному випадку не можна вважати наближеннями множини $P(F, X)$ оптимальних розв’язків вхідної задачі. Це приводить до необхідності створення регуляризуючого оператора, що є певним видом збурень вхідних даних для того, щоб замінити можливо нестійку задачу напевно стійкою еквівалентною задачею. Опишемо процедуру регуляризації за обмеженнями, що ґрунтується на роботах [1–4] та передбачає заміну можливо нестійкої до збурень задачі $Q_P(F, X_{u_2(\delta)})$ $(Q_{Sl}(F, X_{u_2(\delta)}))$ збуреною стійкою задачею $Q_P(F, X^\tau)$ $(Q_{Sl}(F, X^\tau))$, де $X^\tau = G^\tau \cap Z^n$, $G^\tau = \{x \in R^n \mid g_i^\tau(x) \leq 0, i \in N_m\}$, $g_i^\tau(x) = g_i(x) \leq \tau e$, $i \in N_m$, $\tau \in R^1$, $\tau > 0$, $e = (1, \dots, 1) \in R^m$. В основі регуляризації лежать твердження.

Твердження 1. Існує $\tau_0 > 0$: $\forall \tau \in (0, \tau_0]$: $X = G^\tau \cap Z^n$, і $\forall \tau \in (0, \tau_0)$: $X = X^\tau$.

Теорема 1. Існують $\tau_0 > 0$ і $\delta > 0$ такі, що $\forall \tau \in (0, \tau_0)$, $\forall u_2(\delta) \in O_\delta(u_2)$:

$$P(C, X) = P(C, X^\tau) = P(C, X^\tau(\delta)), \quad Sl(C, X) = Sl(C, X^\tau) = Sl(C, X^\tau(\delta)).$$

[1] Таким чином, регуляризовані задачі $Q_P(F, X^\tau)$ і $Q_{Sl}(F, X^\tau)$ мають властивості $T_1, \dots, T_5$ – типів стійкості за обмеженнями.

Список літератури

1. Lebedeva T.T., Semenova N.V., Sergienko T.I. Stability and regularization of vector optimization problems under possible criteria disturbances. *Cybernetics and Systems Analysis*. 2022. **58**, N 5. P. 721–726. <https://doi.org/10.1007/s10559-022-00505-7>.
2. Семенова Н.В. Стійкість за обмеженнями векторних задач цілочислової оптимізації з опуклими квадратичними функціями обмежень. *Теорія оптимальних рішень*. Київ: Ін-т кібернетики імені В.М. Глушкова НАН України, 2007. №6. С. 132–139.
3. Kozerskaya L.N., Lebedeva T.T., Sergienko T.I. Regularization of integer vector optimization problems. *Cybernetics and Systems Analysis*. 1993. 29, N. 3. P. 455–458.
4. Lebedeva T.T., Semenova N.V., Sergienko T.I. Regularization of the vector problem with quadratic criteria of Pareto optimization. *Cybernetics and Systems Analysis*. **59**, N 5. P. 561–566. <https://doi.org/10.1007/s10559-023-00591-1>.

ТОЧНИЙ АЛГОРИТМ ДЛЯ ЗАДАЧІ CVRP НА ОСНОВІ МЕТОДУ BRANCH-PRICE-AND-CUT З GPU-ПРИСКОРЕННЯМ

Ленський М.М., Михальчук Г.Й., lmichael1003@gmail.com

Дніпровський національний університет імені Олеся Гончара (www.dnu.dp.ua)

Задача маршрутизації транспортних засобів з обмеженнями на вантажопідйомність (Capacitated Vehicle Routing Problem, CVRP) [1] є однією з ключових задач комбінаторної оптимізації, розв'язання якої дозволяє суттєво скоротити операційні витрати в логістиці. Класичні точні методи, що базуються на повному переборі, є непридатними для розв'язання великомасштабних задач CVRP через експоненційне зростання кількості можливих маршрутів.

Сучасним стандартом для точного розв'язання задач маршрутизації є метод Branch-Price-and-Cut (BPC) [2]. На відміну від повного перебору, BPC не генерує апріорі всі можливі маршрути. Він використовує механізм генерації стовпців (Column Generation) для динамічного додавання лише тих маршрутів, які потенційно можуть покращити поточний розв'язок. Цей підхід переформулює вихідну задачу у вигляді головної задачі (Master Problem) та підзадачі ціноутворення (Pricing Problem).

Головна задача, що формулюється як задача про розбиття (Set Partitioning Problem), на кожній ітерації обирає оптимальну комбінацію з уже наявних маршрутів, а її розв'язок надає двоїсті змінні. Ці змінні є економічними індикаторами, що передаються в підзадачу ціноутворення для пошуку нових, більш вигідних маршрутів.

Ефективність методу BPC обмежується значною обчислювальною складністю підзадачі ціноутворення. Вона формулюється як задача про елементарний найкоротший шлях з обмеженнями на ресурси (Elementary Shortest Path Problem with Resource Constraints, ESPPRC). Ця задача сама по собі є NP-важкою, що робить її розв'язання головним вузьким місцем всього алгоритму.

Для розв'язання ESPPRC у роботі [3] застосовано алгоритм маркування. Алгоритм ітеративно розширює часткові шляхи, представлені у вигляді «марок» (labels), що зберігають інформацію про накопичену вартість, поточне завантаження та відвіданих клієнтів.

Для пришвидшення алгоритму маркування пропонується використовувати масовий паралелізм GPU. Він забезпечує одночасне виконання ключових етапів алгоритму: розширення тисяч марок, паралельної перевірки на домінування для відсікання неоптимальних шляхів та ущільнення масиву активних марок для наступної ітерації.

Реалізація такого підходу пов'язана зі значними архітектурними викликами, специфічними для GPU: жорсткі вимоги до управління пам'яттю, затримки під час передачі даних, а також необхідність використання атомарних операцій для уникнення станів гонки (race condition). Успішне подолання цих труднощів дозволить значно скоротити час розв'язання підзадачі ціноутворення.

Очікується, що запропонований гібридний підхід дозволить знаходити оптимальні розв'язки для великомасштабних екземплярів CVRP за прийнятний час, що є недосяжним для класичних процесорних реалізацій. Подальші кроки включають програмну реалізацію алгоритму та проведення обчислювальних експериментів на стандартних наборах даних для валідації його ефективності та масштабованості.

Бібліографічні посилання

1. Toth P., Vigo D. The Vehicle Routing Problem. Siam, 2002. 367 p.
2. Pessoa A., Sadykov R., Uchoa E. et al. A Generic Exact Solver for Vehicle Routing and Related Problems. Mathematical Programming. 2020. Vol. 183. P. 483-523.
3. Feillet D., Dejax P., Gendreau M., Gueguen C. An exact algorithm for the elementary shortest path problem with resource constraints: Application to some vehicle routing problems. Networks. 2004. Vol. 44, No. 3. P. 216-229.

**ЗАСТОСУВАННЯ МЕТОДІВ ОПТИМІЗАЦІЇ
НУЛЬОВОГО ТА ПЕРШОГО ПОРЯДКІВ
ДЛЯ РОЗВ'ЯЗУВАННЯ ЕКОНОМІЧНИХ ЗАДАЧ**

Лісняк В.С., nikalisniak@gmail.com,

Волошко В.С., VVL56@i.ua

Дніпровський національний університет імені Олеся Гончара

Оптимізація є важливим інструментом у прикладній математиці, комп'ютерному моделюванні та економіці, оскільки дає змогу знаходити найкращі рішення з множини допустимих. Методи оптимізації розподіляються на кілька класів, зокрема — методи нульового порядку, які не вимагають знаходження похідних, та методи першого порядку, які базуються на використанні градієнтних значень цільової функції. Метою даної роботи є дослідження характеристик методів нульового та першого порядків, оцінка їхніх переваг і недоліків, а також використання для прикладних.

Були розглянуті такі методи [1].

1.Метод покоординатного спуску є прикладом методу нульового порядку. Його особливість полягає в поетапній оптимізації функції за кожною змінною окремо. Переваги – можливість роботи з недиференційованими функціями та простота реалізації.. Недолік – відносно повільна збіжність при зростанні розмірності задачі.

2.Методи першого порядку. Наприклад, градієнтний метод, що спирається на застосування похідних, і полягає у зміні оптимізаційної функції у напрямку її максимального збільшення(зменшення), гарантуючи високу швидкість збіжності та точність. Недоліками методу є потреба у обчисленні градієнтів і залежність результату від вибору початкових даних.

Розглянемо задачі із застосуванням цих методів та результати розв'язку.

1. Оптимізація функції витрат: $f(x, y) = 2x^2 + 3y^2 - 8x - 12y + 24$,

де x, y — обсяги виробництва продуктів.

$$f(x, y) \rightarrow \min, \quad x \geq 0, \quad y \geq 0.$$

2. Оптимізація функції прибутку:

$$P(x, y) = 4x + 6y - x^2 - y^2,$$

де x, y — обсяги виробництва, $4x, 6y$ — лінійні доходи, x^2, y^2 — штрафи.

$$P(x, y) \rightarrow \max, \quad x \geq 0, \quad y \geq 0, \quad x + y \geq 5$$

3. Оптимізація структури інвестиційного портфеля, мінімізація ризику:

$$R(x, y) = (x - 3)^2 + (y - 5)^2 + 2xy,$$

де x, y — два активи з інвестиційними обсягами, сума фіксована.

$$R(x, y) \rightarrow \min, \quad x \geq 0, \quad y \geq 0, \quad x + y = 8.$$

Задача 1: Мінімізація витрат

Координатний спуск: [2, 2] Значення: 4 Ітерації: 30

Гradientний спуск: [1.99947037, 1.99999149] Значення: 4.000000561233136 Ітерації: 2055

Задача 2: Максимізація прибутку

Координатний спуск: [2, 3] Прибуток: 13 Ітерації: 30

Гradientний спуск: [1.9998844, 2.99976881] Прибуток: 12.999999933187823 Ітерації: 902

Задача 3: Оптимізація портфеля

Координатний спуск: [0, 8] Ризик: 18 Ітерації: 34

Гradientний спуск: [0, 8.002] Ризик: 18.012004000000005 Ітерації: 4120

За цими задачами отримані такі чисельні результати

Метод покоординатного спуску показав стабільність і точність у простих моделях із незалежними змінними, тоді як gradientний метод продемонстрував вищу швидкість збіжності та кращу адаптацію до задач зі складною геометрією та взаємозалежними параметрами, хоча і виявив вразливість до початкових умов.

Бібліографічні посилання

1. Кісельова О.М., Шевельова А.Є. Чисельні методи оптимізації: навч. посіб. Д.: Вид-во ДНУ.2008. – 212 с.

МЕТРИКИ ОЦІНЮВАННЯ ЕФЕКТИВНОСТІ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ ДЛЯ ОБРОБКИ ПРИРОДНОЇ МОВИ

Логвин Д.А., lohvyn_d@365.dnu.edu.ua,

Божуха Л.М., bozhukha_l@365.dnu.edu.ua

Дніпровський національний університет імені Олеся Гончара

Зі зростанням обсягів даних і поширенням систем штучного інтелекту питання оцінювання якості та продуктивності програмного забезпечення для оброблення інформації набуло особливої актуальності. Сучасні рішення (напр. пошукові системи, мовні моделі) вимагають не лише високої точності результатів, але й стабільної швидкодії, економного використання ресурсів та енергоспоживання.

У межах дослідження запропоновано класифікувати метрики оцінювання систем на три основні групи.

1. **Метрики якості результату.** Визначають відповідність отриманого результату очікуваному або еталонному. Для задач класифікації, загалом, це точність (Accuracy), повнота/чутливість (Recall/Sensitivity), специфічність (Specificity), точність прогнозного значення (Precision) та F-міра (F1-score), крива ROC (Receiver Operating Characteristic).
2. **Метрики продуктивності.** Характеризують швидкодію, ефективність використання ресурсів і стабільність виконання програмних модулів. Основними є час відповіді (Response Time), затримка (Latency), пропускна здатність (Throughput), використання оперативної пам'яті (Memory Usage), апаратне навантаження та енергоспоживання.
3. **Метрики стабільності та надійності.** Відображають стійкість системи до змін умов виконання та навантаження: варіацію часу обробки, частоту помилок, стабільність при довготривалому використанні, температуру пристрою та швидкість деградації продуктивності.

Запропонована методологія оцінювання передбачає багаторівневий підхід до оцінки ефективності програмного забезпечення:

1. **На рівні алгоритму** оцінюється якість опрацювання даних – точність класифікації, релевантність генерації або перекладу, коректність синтаксичного розбору.
2. **На рівні системи** – продуктивність виконання, використання ресурсів, масштабованість та стійкість до навантажень.
3. **На рівні інтеграції** – взаємодія з іншими підсистемами, затримки під час виклику API, ефективність паралельного виконання запитів тощо.

Для зручності порівняння різних реалізацій запропоновано застосувати інтегральний показник ефективності, який об'єднує результати декількох метрик у єдину шкалу. Кожна метрика нормалізується до діапазону $[0;1]$, після чого зважується за важливістю для конкретного типу застосунку. Загальна формула показнику ефективності може бути подана як:

$$EI = \sum_{i=1}^n w_i \cdot M_i^{norm},$$

Де w_i – ваговий коефіцієнт i -тої метрики, M_i^{norm} – нормалізоване значення i -тої метрики, а сума всіх $w_i = 1$.

Сформована система метрик дозволяє цілісно оцінювати ефективність програмного забезпечення для оброблення інформації. Поєднання якісних та продуктивних характеристик у межах єдиної методології відкриває можливість стандартизувати процес оцінювання.

Бібліографічні посилання

1. Liu H., Liu X., Hu G. *Metrics and evaluations for computational and sustainable AI efficiency*. arXiv:2510.17885, 2025. DOI: 10.48550.
2. Schmidová P., Mahamood S., Balloccu S. et al. *Automatic Metrics in Natural Language Generation: A Survey of Current Evaluation Practices*. arXiv:2408.09169, 2024. DOI: 10.48550.

ЗАСТОСУВАННЯ ГЛИБИННОГО НАВЧАННЯ ДЛЯ СЕМАНТИЧНОГО АНАЛІЗУ ТЕКСТІВ ВАКАНСІЙ

Луг О. О., lug.oksana.l@gmail.com

Кузенков О. О., kuzenkov1986@gmail.com

Дніпровський національний університет імені Олеся Гончара

У сучасних умовах розвитку ІТ-ринку обсяг інформації у сфері працевлаштування постійно зростає. Сотні і тисячі нових вакансій публікуються щодня, і при цьому кожна з них має власну структуру, лексику та стиль, що робить неможливим простий пошук необхідних компетенцій на основі ключових слів.

Для ефективного аналізу таких текстів виникає потреба у створенні моделей, здатних семантично інтерпретувати описи вакансій та автоматично виділяти затребувані **працедавцем** навички.

Задачею дослідження є навчання моделі для виконання аналізу текстових описів вакансій та формування списку навичок на основі навчальних пар «опис вакансії – навички». Вхідними даними є тексти вакансій, зібрані на dou.ua, а вихідними - структурований список навичок, розділених комами.

Для вирішення такої задачі доречно задіяти моделі глибинного навчання та трансформери, що здатні враховувати контекст і семантику, а не лише окремі терміни. Для цього застосовано NLP на базі глибинного навчання: модель mT5 враховує контекст слів у текстах із вільною структурою. Для донавчання використано Hugging Face Transformers із PyTorch, а обробку даних — бібліотеку pandas.

На даний момент було навчено модель mT5 на датасеті train (182 приклади) і validation (80 прикладів) протягом 10 епох, із поступовим зниженням втрат від ~19.5 до ~3.7, що підтверджує її здатність правильно зіставляти описи вакансій із релевантними навичками. Модель показує результати, де успішно повторює зміст потрібної частини тексту, в якій знаходиться перелік навичок, виписує принаймні деякі з них і використовує ту саму мову, що застосована в описі вакансії.

Результати свідчать про ефективність підходу та його перспективність для масштабування і подальшого навчання на більшій кількості прикладів, що має значно покращити точність відповідей.

Бібліографічні посилання

1. Google's mT5. URL: <https://huggingface.co/google/mt5-small>

**ПРОГРАМНИЙ ІНСТРУМЕНТАРІЙ ДЛЯ АНАЛІЗУ І ВІЗУАЛІЗАЦІЇ
ХАРАКТЕРИСТИЧНИХ ОЗНАК В ЕЛЕКТРОКАРДІОГРАМАХ**

Ляшко А.В.¹ (andrey_liashko@knu.ua), **Стеля О.Б.**², **Куляс А.І.**²,
Єфремов М.С.¹., **Крак Ю.В.**^{1,2}

¹*Київський національний університет імені Тараса Шевченка*

²*Інститут кібернетики імені В.М. Глушкова НАН України*

В дослідженні наведено опис деяких функціональних можливостей програмного інструменту для обробки електрокардіограм (ЕКГ). Створені інструменти об'єднані графічним інтерфейсом та призначені для підтримки прийняття рішення при постановці діагнозу лікарем. Важливість правильної обробки артефактів в даних ЕКГ важко переоцінити. Навіть один пропущений при аналізі артефакт може вплинути на остаточний діагноз. Таким чином, надійні математичні алгоритми обробки даних ЕКГ є запорукою їх впевненого аналізу.

Створений інструментарій розвивається та наповнюється новими алгоритмами та відображенням результатів їх роботи на графіках та діаграмах. На вхід програмного комплексу подаються дані ЕКГ з багатьох апаратів, які містять шуми різної природи та інтенсивності. Для подолання високочастотних шумів здійснюється фільтрація або згладжування. Одним з базових елементів обробки даних ЕКГ є точне визначення R-зубців. Для цього розроблено варіант порогового методу [1], який показав високу точність на даних з бази аритмій МІТ-ВІН [2]. За знайденими значеннями абсцис R-зубців знаходяться R-R інтервали. Далі серед R-R інтервалів знаходяться артефакти у вигляді інтервалів з аномальними довжинами та здійснюється їх графічне відображення (рис. 1) для першочергового перегляду лікарем. Пошук та виділення артефактів може здійснюватися як в автоматизованому, так і в ручному (з візуальним контролем) режимах. Хоча розроблений метод визначення R-зубців не вимагає попереднього вирівнювання базової лінії (ізоелектричної лінії) існують випадки, коли цю лінію треба вирівнювати. В ідеалі ізоелектрична лінія є прямою. В

закладеному алгоритмі здійснюється пошук точок, які належать ділянкам ЕКГ близьким до прямої лінії. По знайденим точкам будується інтерполяційний або згладжуючий сплайни. Далі знайдена за допомогою сплайна лінія віднімається від вихідного сигналу ЕКГ.

Створено інтерфейс, який дає змогу графічно візуалізувати отриману інформацію [3]. За характером поведінки графіка можна зробити припущення щодо наявності аномалій у роботі серця (рис. 1).

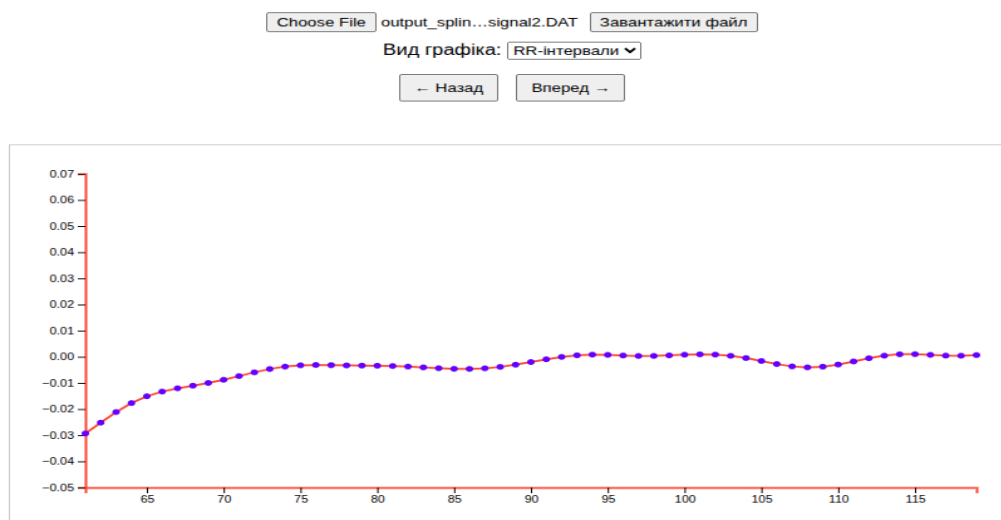


Рис. 1. Головне меню програми

Бібліографічні посилання

1. Крак Ю.В., Стеля О.Б., Єфремов М.С., Ляшко А.В. Інформаційна технологія обробки даних електрокардіограм для знаходження R-піків. *Доповіді Національної академії наук України*. 2024. No 5, С. 44-52. DOI: 10.15407/dopovidi2024.05.044.
2. Moody G., B., & Mark R., G., The impact of the MIT-BIH Arrhythmia Database. *EEE Engineering in Medicine and Biology Magazine*, 2001, vol. 20, issue 3, pp. 43-50, DOI: 10.1109/51.932724.
3. Ляшко А. В., Єфремов М. С. Алгоритмічно-графічні засоби попередньої комп'ютерної обробки даних Холтера. *Журнал обчислювальної та прикладної математики*. №1. С. 30-39. DOI: <https://doi.org/10.17721/2706-9699.2024.1.02>.

ОГЛЯД ЗАСТОСУВАННЯ МЕТОДІВ МАШИННОГО НАВЧАННЯ ДО ОБЕРНЕНОЇ ЗАДАЧІ ІДЕНТИФІКАЦІЇ ПАРАМЕТРІВ ТОНКОСТІННОЇ ОБОЛОНКИ

Marac O., omahas.edu@gmail.com

Гук Н., natalyguk29@gmail.com

Oles Honchar Dnipro National University (www.dnu.dp.ua)

Задача ідентифікації параметрів тонкостінної оболонки належить до класу нелінійних та некоректно поставлених задач. У цих задачах класичні методи чисельного аналізу є нестійкими. Щоб подолати цю проблему, можна застосувати перехід до умовно-коректної постановки задачі та використати методи машинного навчання.

До таких методів, у тому числі, належать методи прямого глибокого навчання, фізично-інформовані нейронні мережі (PINNs), сучасні варіанти нейроеволюційних алгоритмів (такі, як RNEAT), а також пропонується модифікація нейроеволюційного підходу на основі нейроеволюції топологій (NEAT+BP).

Методи прямого глибокого навчання орієнтовані на побудову апроксимації між даними та невідомими параметрами без прямого залучення рівнянь моделі [1]. Вони ефективні для задач із великими вибірками експериментальних даних, але втрачають стабільність та здатність до узагальнення на невеликих вибірках.

Фізично-інформовані нейронні мережі повною мірою реалізують принцип умовно-коректної постановки, коли додатковою інформацією для стабілізації розв'язку виступають власне фізичні співвідношення [2]. У функцію втрат таких мереж безпосередньо включаються рівняння рівноваги, конститутивні залежності та геометричні умови. Це забезпечує фізичну узгодженість розв'язку, але й значно ускладнює процес навчання через можливу погану обумовленість системи та високі вимоги до доступних обчислювальних ресурсів.

Сучасні варіанти нейроеволюції, такі як RNEAT, поєднують механізми класичної еволюції топологій із принципами картезіанського

генетичного програмування, що дозволяє автоматизовано формувати складні архітектури мереж [3].

У модифікації гібридного алгоритму NEAT архітектура нейромережі формується еволюційним шляхом із покроковим ускладненням структури з додатковим донавчанням класичними методами для уточнення вагових параметрів [4]. До переваг такого підходу належить можливість побудови мережі без попереднього знання фізичних законів. Це дозволяє розв'язувати задачі з неповною або нечітко визначеною математичною моделлю.

Отже, вибір методу визначається метою, наявними ресурсами та природою доступної інформації про систему. Методи прямого навчання переважають у випадках великих емпіричних баз даних. Фізично-інформовані нейронні мережі є ефективними якщо відомі рівняння рівноваги та доступні обчислення на графічних процесорах.

При цьому, NEAT+BP залишається актуальним, коли пріоритетом є: структурна адаптивність мережі, необхідність її автоматичного і контрольованого формування без попереднього проєктування архітектури, простота фінальної архітектури нейронної мережі, невисокі вимоги до обчислювальних ресурсів, невисока складність реалізації.

Бібліографічні посилання

1. Fuchs, F., et al. Comparative Study of Various Neural Network Types for Direct Inverse Material Parameter Identification in Numerical Simulations. *Applied Sciences*, 2022, 12(24), 12793. <https://doi.org/10.3390/app122412793>
2. Raissi, M., Perdikaris, P. & Karniadakis, G. E. (2019). Physics-Informed Neural Networks: A Deep Learning Framework for Solving Forward and Inverse Problems Involving Nonlinear Partial Differential Equations. *Journal of Computational Physics*, 378, 686–707. <https://doi.org/10.1016/j.jcp.2018.10.045>
3. Turner, A.J., Miller, J.F. Recurrent Cartesian Genetic Programming of Artificial Neural Networks. *Genet Program Evolvable Mach* 18, 185–212 (2017). <https://doi.org/10.1007/s10710-016-9276-6>
4. Mahas, O. S., Guk, N. A. Neural network approach to the selection of an observed system model. *Problems of Applied Mathematics and Mathematical Modeling*, 24, 119-126 (2024). <https://doi.org/10.15421/322412>

ВИКОРИСТАННЯ МЕТОДУ МОНТЕ-КАРЛО ДЛЯ ГЕНЕРАЦІЇ ПОЧАТКОВИХ УМОВ У ПРЯМІЙ ЗАДАЧІ ДЛЯ БАГАТОШАРОВОГО ПАКЕТУ НА ТВЕРДІЙ ОСНОВІ

Мажара К.О., mazhara.kirill@gmail.com,

Трофімов О.В., atrof2222@gmail.com,

Дніпровський національний університет імені Олеся Гончара

Метод Монте-Карло є одним із найпоширеніших статистичних підходів до моделювання складних фізичних процесів, у яких аналітичний опис є надто складним або неможливим. Його суть полягає у використанні великої кількості випадкових випробувань для оцінювання поведінки системи за певних параметрів. Основною ідеєю є апроксимація очікуваного значення функції через середнє значення результатів випадкових вибірок:

$$E[f(x)] \approx \frac{1}{N} \sum_{i=1}^N f(x_i),$$

де x_i — випадкові реалізації параметрів системи, N — кількість симуляцій.

Розглянемо використання методу Монте-Карло для генерації початкових умов у прямій задачі моделювання тришарового дорожнього пакету (асфальтобетон — щебенева основа — ґрунт) на **твердій основі**. Пряма задача полягає у визначенні напружено-деформованого стану конструкції при заданих властивостях шарів та зовнішньому навантаженні. Метод Монте-Карло дозволяє згенерувати широкий набір випадкових початкових параметрів, серед яких: геометричні параметри, механічні властивості, навантаження, умови основи.

Таблиця 1 – Приклади початкових параметрів

Параметр	Позначення	Тип розподілу	Діапазон (приклад)	Фізичний зміст
Товщина асфальту	h_1	Нормальний	0.08–0.12 м	Верхній шар

Модуль пружності	E_i	Нормальний	1–10 ГПа	Матеріальна жорсткість шарів
Коефіцієнт Пуассона	μ_i	Рівномірний	0.25–0.35	Співвідношення поперечної/поздовжньої деформації
Модуль зсуву	G_i	Обчислюється	залежить від E_i, μ_i	Опір зсуву шару
Інтенсивність Навантаження	q	Нормальний	0.5–1.0 МПа	Тиск від колеса
Геометрія межі	$S_i(x)$	Стохастична функція	± 2 –5 мм відхилення	Нерівність контактів
Густина	ρ_i	Рівномірний	1800–2500 кг/м ³	Маса матеріалу шару

Такий підхід до формування вхідних даних є універсальним і може бути використаний як у класичних чисельних методах розв’язування прямої задачі — таких як метод скінченних елементів, метод скінченних різниць, метод граничних елементів, — так і в сучасних підходах, заснованих на нейронних мережах. У разі застосування глибинних моделей метод Монте-Карло може виступати джерелом великої кількості синтетичних даних для навчання, що суттєво покращує узагальнювальну здатність моделі. Таким чином, використання цього методу для генерації параметрів шаруватого пакету дозволяє забезпечити статистичну різноманітність початкових умов, автоматизувати процес підготовки даних і створити основу для подальшого чисельного чи нейромережевого розв’язку прямої задачі.

Бібліографічні посилання

1. Metropolis, N., & Ulam, S. The Monte Carlo Method. *Journal of the American Statistical Association*, 1949.
2. Rubinstein, R. Y., & Kroese, D. P. *Simulation and the Monte Carlo Method*. Wiley, 3rd Edition, 2017.
3. Goodfellow, I., Bengio, Y., & Courville, A. *Deep Learning*. MIT Press, 2016.

ПРОЄКТ ІНФОРМАЦІЙНОЇ СИСТЕМИ НАДАННЯ РЕКОМЕНДАЦІЙ З ФОРМУВАННЯ ОСОБИСТОГО СТИЛЮ

Максимів І.А., ivankamaksymiv019@gmail.com

Басюк Т. М., taras.m.basyuk@lpnu.ua

Національний університет «Львівська політехніка»

Сфера краси та особистого стилю нині демонструє одне з найшвидших зростань серед сервісних галузей. Активний розвиток цифрових технологій сприяє появі нових рішень для підвищення якості послуг та розширення можливостей взаємодії з клієнтами. AI-рішення для рекомендацій у beauty-сфері ґрунтуються на поєднанні різних технологій: комп'ютерного зору, який аналізує зовнішність за фото, машинного навчання, що визначає характеристики обличчя й стильові особливості, інструментів NLP для розуміння потреб і побажань клієнтів [1].

Це надає можливість формувати індивідуальні поради з урахуванням великої кількості показників: від форми обличчя і кольоротипу до актуальних трендів та типу заходу, для якого створюється образ.

Оскільки робота візажиста і стиліста безпосередньо пов'язана з аналізом зображень та професійними правилами, її досить легко частково автоматизувати. Попит на такі інтелектуальні системи постійно зростає, адже мода швидко змінюється, а алгоритми здатні одразу адаптуватися до нових напрямів.

Для проєктування вебсайту візажиста-стиліста використовується об'єктно-орієнтований підхід, що надає засоби із проєктування з використанням UML [2]. Діаграма варіантів використання зображена на рисунку 1. На ній відображено акторів, можливі сценарії використання сайту та логічні зв'язки між ними, що дає змогу чітко структурувати вимоги до майбутньої системи.

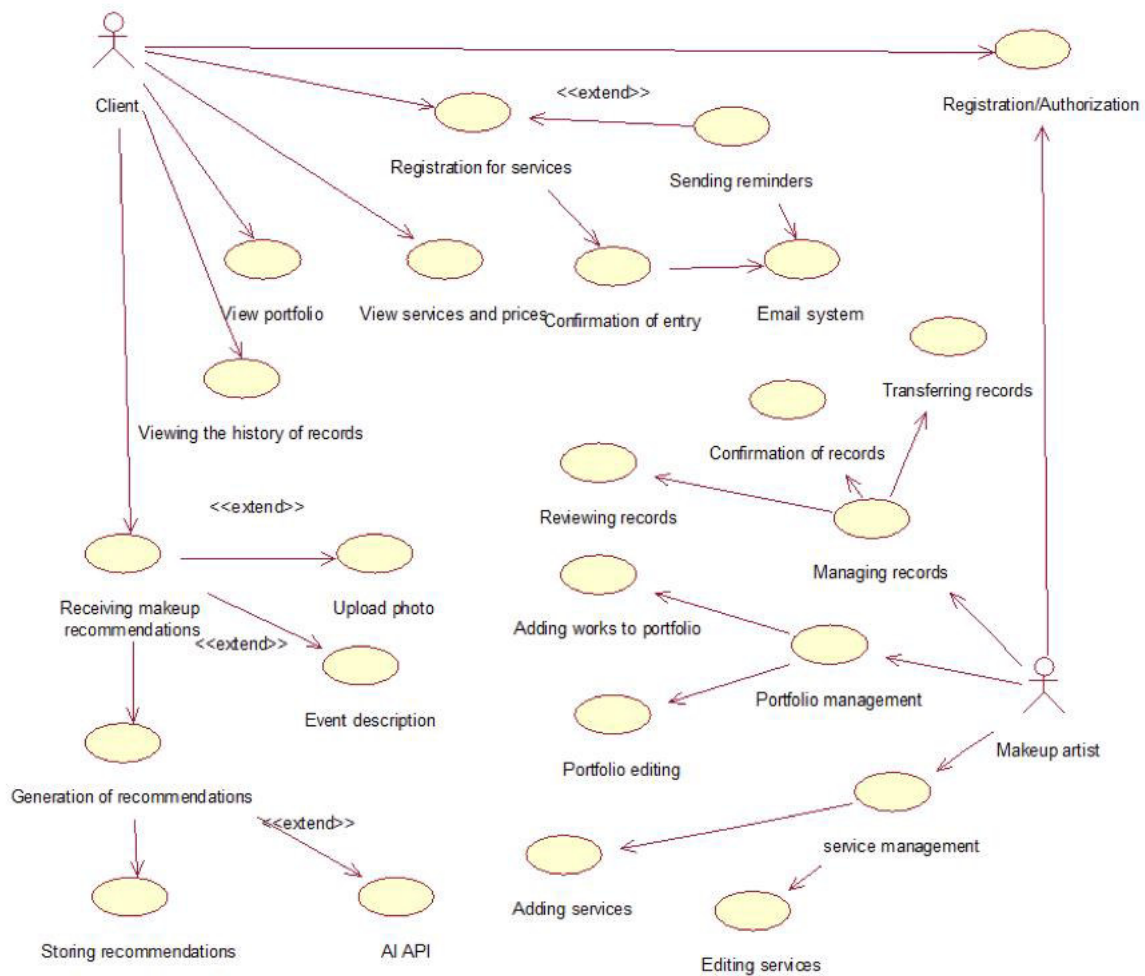


Рис. 1. Діаграма варіантів використання

Впровадження такої системи підвищуватиме якість обслуговування клієнтів, створюватиме конкурентні переваги та дозволить масштабувати послуги без пропорційного збільшення штату спеціалістів. Такий тип системи дозволяє надавати професійні консультації цілодобово, незалежно від географічного розташування клієнтів.

Список використаних джерел

1. Generative AI in Beauty Industry: Top Brands Examples
https://www.miquido.com/blog/generative-ai-in-beauty-industry/?utm_source=chatgpt.com
2. Use Case Diagram - Unified Modeling Language (UML)
<https://www.geeksforgeeks.org/system-design/use-case-diagram/>

ОПТИМІЗАЦІЯ КОМПОНУВАННЯ М'ЯКИХ ЕЛІПСОЇДІВ

Мелащенко О.П.^{1,2}, Романова Т. Є^{1,3}, Дюрягина З.А.⁴, Кравченко О.В.¹
o.p.melashenko@gmail.com

¹ Інститут енергетичних систем і машин ім. А.М. Підгорного НАН
України, Харків, Україна

² Харківський національний університет ім. В.Н. Каразіна, Харків, Україна

³ Університет Лідсу, Лідс, Велика Британія

⁴ Національний університет «Львівська політехніка», Львів, Україна

Задача оптимізації компонування еліпсоїдів є NP-складною, та має широкий спектр застосувань, зокрема, у матеріалознавстві, гірничодобувній промисловості, біоінженерії, нанотехнологіях, адитивному виробництві. Багато публікацій присвячено компонуванню 3D об'єктів фіксованої форми [1,2]. Сучасні застосування вимагають врахування нестандартних обмежень розміщення геометричних об'єктів, які можуть змінювати свою форму під тиском зовнішніх сил (м'яких об'єктів) [3,4]. У даному дослідженні розглядається задача компонування заданої множини м'яких еліпсоїдів $E_i, i=1, \dots, n$, у прямокутний контейнер мінімального об'єму. Еліпсоїди можуть змінювати форму в межах заданих обмежень деформації, зберігаючи свій об'єм під час деформації. Запропоновано засоби математичного моделювання умов розміщення м'яких еліпсоїдів із застосуванням методу ϕ -функції. Сформульована відповідна математична модель як задача нелінійного програмування. Розроблено метод розв'язання, що поєднує швидкий евристичний алгоритм для генерування допустимих стартових розміщень та стратегію декомпозиції [5] для пошуку локальних мінімумів. Наведено результати обчислювальних експериментів.

На рисунку 1 наведено фрагменти прикладів компонування м'яких еліпсоїдів у прямокутній контейнер кубоїд мінімального об'єму при різних значеннях параметру деформації $\mu_i, i=1, \dots, n$.

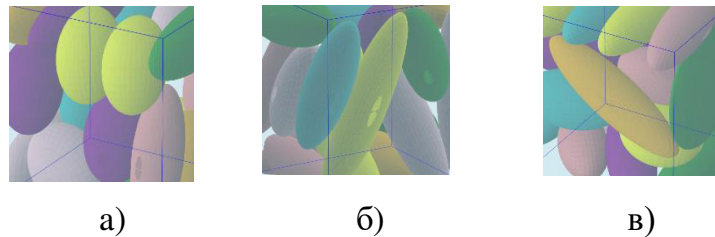


Рис. 1 – Компонування м'яких еліпсоїдів у кубоїд мінімального об'єму за умов: а) $1 \leq \mu_i \leq 2$ б) $1.5 \leq \mu_i \leq 2.5$ в) $0.5 \leq \mu_i \leq 2.5$

Обчислювальні експерименти підтвердили ефективність і надійність запропонованого методу, а також засвідчили його потенціал для практичного застосування — зокрема, під час моделювання гранульованих матеріалів, оптимізації мікроструктури порошкових сплавів і симулювання структурних перетворень м'яких частинок під дією зовнішніх сил. Подальші напрями досліджень передбачають розширення підходу на динамічні процеси ущільнення, урахування теплових та мультифізичних ефектів, а також розроблення паралельних стратегій розв'язання для великомасштабних промислових задач.

1. Kallrath, J. Packing ellipsoids into volume-minimizing rectangular boxes. *Journal of Global Optimization*, 67 (1-2), 151-185, 2017.
2. Romanova, T., Litvinchev, I., Pankratov, A. Packing ellipsoids in an optimized cylinder. *European Journal of Operational Research*, 285 (2), pp. 429-443, 2020.
3. Melashenko, O., Romanova, T., Litvinchev, I., Martínez Gomez, C. G., Yang, R., Sun, B.A Model-Based Heuristic for Packing Soft Rotated Rectangles in an Optimized Convex Container with Prohibited Zones. *Mathematics*. 13(3), 493, 2025.
4. Litvinchev, I., Infante, L., Romanova, T., Martinez-Noa, A., Gutierrez, L. (2024). Optimized Packing Soft Convex Polygons. In: Marmolejo-Saucedo, J.A., Rodríguez-Aguilar, R., Vasant, P., Litvinchev, I., Retana-Blanco, B.M. (eds) *Computer Science and Engineering in Health Services. EAI/Springer Innovations in Communication and Computing*. Springer, Cham.
5. Romanova, T., Stoyan, Y., Pankratov, A., Litvinchev, I., Marmolejo, J.A. Decomposition Algorithm for Irregular Placement Problems In: Vasant, P., Zelinka, I., Weber, GW. (eds) *Intelligent Computing and Optimization. Advances in Intelligent Systems and Computing*, 2020, 1072, 214-221. Springer, Cham.

The work is supported by the British Academy (Grant #100072) and the Ukrainian Government Project «Optimization of the Microstructure and Functional Properties of Critical-Application Components Made from Powder Alloys for Additive Manufacturing» (grant # 0125U001884)

ЗАСТОСУВАННЯ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ В РОЗПОДІЛЕНИХ БАЗАХ ДАНИХ

Миргородський А.В., Ліщинська Л.Б., mirgorodskijav@gmail.com

Вінницький національний технічний університет

Стрімкий розвиток розподілених баз даних супроводжується підвищенням вимог до їх швидкодії, узгодженості даних та відмовостійкості. Традиційні методи синхронізації і реплікації даних дедалі гірше справляються з динамічними умовами роботи сучасних розподілених систем, де навантаження змінюється нерівномірно, а обсяг інформації зростає експоненційно [1]. У зв'язку з цим **актуальним** є застосування інтелектуальних інформаційних систем та методів штучного інтелекту для адаптивного керування процесами узгодженості, сегментації та реплікації даних у розподілених середовищах.

Мета дослідження полягає у розробці підходів до інтеграції інтелектуальних методів у процеси керування розподіленими базами даних з метою підвищення їх адаптивності та продуктивності використання ресурсів.

Застосування штучного інтелекту у політиках узгодженості контекстно-орієнтованої моделі забезпечує автоматичну адаптацію системи без втручання адміністратора. Модель CAnDoR використовує час відповіді та час синхронізації як ключові метрики для динамічного розподілу сегментів даних між вузлами, поєднуючи високу продуктивність з доступністю [2]. Модель R-TBC/RTA організовує вузли у вигляді дерева, де центральні вузли забезпечують строгу узгодженість, а периферійні – кінцеву [3]. Такий підхід підвищує відмовостійкість системи та дозволяє динамічно змінювати топологію залежно від стану мережі.

У сфері сегментації даних розглянуто декілька інтелектуальних методів. Темпоральна сегментація класифікує дані як «гарячі», «теплі» або «холодні» залежно від частоти запитів і дозволяє ефективно розподіляти

ресурси. Концепція мікро-шардів у системі Akkio передбачає формування малих сегментів клієнтськими додатками для підвищення швидкодії та гнучкості реплікації [4]. Метод сегментації на основі мережі формує фрагменти даних на основі аналізу взаємодій користувачів, що зменшує кількість міжсегментних транзакцій. Сегментація на основі знань застосовує інтелектуальні алгоритми для поділу складних запитів на незалежні підзапити, що можуть виконуватися паралельно.

Отже, інтеграція інтелектуальних систем у розподілені бази даних підвищує їх адаптивність та швидкодію використання ресурсів. Результати дослідження підтверджують перспективність використання штучного інтелекту для побудови автономних розподілених СКБД нового покоління.

Список використаної літератури

1. TDSQL: tencent distributed database system / Y. Chen et al. Proceedings of the VLDB Endowment. 2024. Vol. 17, no. 12. P. 3869–3882. URL: <https://doi.org/10.14778/3685800.3685812>.

2. Mauffret E., Vernier F., Monnet S. CAnDoR: consistency aware dynamic data replication. 2019 IEEE 18th international symposium on network computing and applications (NCA), Cambridge, MA, USA, 26–28 September 2019. 2019. URL: <https://doi.org/10.1109/nca.2019.8935051> (date of access: 30.06.2025).

3. Advanced consistency management of highly-distributed transactional database in a hybrid cloud environment using novel R-TBC/RTA approach / J. Dizdarevic et al. Journal of cloud computing. 2021. Vol. 10, no. 1. URL: <https://doi.org/10.1186/s13677-021-00230-0> (date of access: 30.06.2025).

4. Sharding the shards: managing datastore locality at scale with Akkio / M. Annamalai et al. Proceedings of the 13th USENIX conference on operating systems design and implementation. USA, 2018. P. 445–460.

ДВОЕТАПНИЙ МЕТОД ОПТИМІЗАЦІЇ МАРШРУТІВ ДЛЯ ГЕТЕРОГЕННИХ ТРАНСПОРТНИХ СИСТЕМ

Мірзаєв Т. Р., Михальчук Г. Й., timur.mirzaiev@gmail.com

Дніпровський національний університет імені Олеся Гончара

Сучасні логістичні системи вимагають гнучких рішень. Для одних бізнес-моделей (експрес-доставка їжі) критичною є мінімізація часу доставки кожного окремого замовлення. Для інших (планова доставка товарів) пріоритетом є максимальне завантаження транспорту та мінімізація сукупних витрат. Статичні алгоритми не можуть задовольнити обидві потреби. Актуальною є розробка параметричних методів, які можна адаптувати під конкретні бізнес-цілі без зміни коду.

Для розв'язання задачі запропоновано двоетапний метод.

Етап 1: Інтелектуальне призначення. Вибір оптимального типу транспорту для кожного замовлення базується на мінімізації функції витрат TotalCost:

$$\text{TotalCost} = \alpha * \text{TimeCost} + (1 - \alpha) * \text{InefficiencyPenalty}$$

Ключовим елементом моделі є коефіцієнт пріоритету часу α . Цей параметр дозволяє гнучко налаштовувати стратегію призначення замовлень. Якщо $\alpha \rightarrow 1$, функція витрат практично повністю залежить від TimeCost. Система буде надавати перевагу найшвидшій доставці, обираючи транспорт, який виконає замовлення швидше, навіть якщо це економічно менш ефективно (наприклад, відправка автомобіля за малим пакетом). Якщо $\alpha \rightarrow 0$, домінує InefficiencyPenalty, тобто система буде прагнути до максимальної ефективності та групування, призначаючи замовлення транспорту з відповідною вантажопідйомністю, щоб мінімізувати "перевезення повітря", навіть якщо це призведе до незначного збільшення часу доставки окремого замовлення.

Етап 2: Побудова маршруту. Після розподілу замовлень на групи для кожної з них будуються маршрути за допомогою евристики найближчого сусіда.

Розроблено програмне забезпечення, що реалізує запропонований метод. Архітектура системи базується на принципах ООП та патерні "Стратегія". Коефіцієнт α є конфігурованим параметром класу `AssignmentModel`, що дозволяє легко змінювати стратегію.

Проведено серію обчислювальних експериментів, що підтверджують гнучкість методу. Зокрема підтверджено коректну реакцію моделі на домінантні фактори (відстань, ефективність, жорсткі обмеження), а також досліджено вплив коефіцієнта α . Проведено ключовий експеримент для замовлення, де автомобіль є швидшим, але скутер значно ефективнішим за завантаженням. За високого значення α (пріоритет швидкості) модель обирала автомобіль. У разі низького значенні α (пріоритет ефективності) для того ж самого замовлення модель обирала скутер. Це доводить, що метод є керованим і дозволяє свідомо обирати операційну стратегію.

Головною перевагою розробленого методу є його гнучкість. Завдяки параметризації функції витрат за допомогою коефіцієнта α , система може бути легко адаптована під різні бізнес-цілі, від максимальної швидкості до максимальної економічної ефективності, без необхідності зміни архітектури чи алгоритмів. Це перетворює її на практичний інструмент для широкого кола логістичних завдань. Подальші дослідження можуть бути спрямовані на розробку алгоритмів динамічної зміни α залежно від поточної завантаженості системи.

МАТЕМАТИЧНЕ МОДЕЛЮВАННЯ ДИНАМІКИ ЗВЕРНЕНЬ ДО МЕДИЧНОГО ЦЕНТРУ З ЕЛЕМЕНТАМИ АДАПТИВНОГО ПРОГНОЗУВАННЯ ТА ВИЯВЛЕННЯ АНОМАЛІЙ

Моїсєєнков Д.В., danil.moiseenkoff@gmail.com

Зайцев В.Г., vadymzaytsev65@gmail.com

Дніпровський національний університет імені Олеся Гончара

Зростання навантаження на медичні заклади потребує впровадження сучасних інтелектуальних систем підтримки прийняття рішень [6], [7]. Особливої актуальності набуває проблема моделювання та прогнозування потоків пацієнтських звернень з урахуванням сезонних коливань, пікових навантажень та аномальних сплесків [8], [9], [10]. Традиційні підходи не враховують нелінійні ефекти та змінність часових структур [5]. Тому метою даної роботи є розробка адаптивної математичної моделі з механізмами виявлення аномалій у реальному часі для підвищення ефективності управління ресурсами медичного центру.

У дослідженні реалізовано комбінований підхід з використанням наступних методів та моделей:

- статистичне згладжування часових рядів (ковзне середнє, STL-декомпозиція) [1];
- адаптивне прогнозування за допомогою моделей ARIMA та Holt–Winters [2], [3];
- оцінка стаціонарності через ADF-тест [4];
- автоматичне визначення оптимальних параметрів (p, d, q) на основі інформаційного критерію Акаїке (AIC) [1], [2];
- виявлення аномалій через статистичні порогові методи, контрольні графіки та адаптивні інтервали довіри [6], [7].

Особливістю моделі є здатність перебудовуватись при зміні динаміки потоку звернень, що дає змогу мінімізувати похибки прогнозу навіть в умовах різких флуктуацій [1], [3].

Результати дослідження. Апробацію методів проведено на реальних даних медичного центру за кілька місяців спостережень. Система показала:

- середню абсолютну похибку прогнозу $MAE < 0.1$;
- своєчасне виявлення аномальних сплесків навантаження;
- покращення розподілу робочих змін персоналу.

Отримані результати демонструють ефективність інтеграції математичного моделювання з адаптивними алгоритмами прогнозування та моніторингу [6], [7].

Висновки. Розроблена система дозволяє:

- точно прогнозувати динаміку звернень у медичні установи;
- автоматично фіксувати аномальні відхилення в потоках даних;
- підтримувати оперативне управління ресурсами та покращувати якість медичних послуг.

Запропонований підхід може бути масштабований для різних типів медичних закладів та інтегрований у системи цифрового менеджменту охорони здоров'я [9], [10].

Бібліографічні посилання:

1. Hyndman R. J., Athanasopoulos G. *Forecasting: Principles and Practice*. – 3rd ed. – Melbourne: OTexts, 2021. – 426 p.
2. Box G. E. P., Jenkins G. M., Reinsel G. C., Ljung G. M. *Time Series Analysis: Forecasting and Control*. – 6th ed. – Hoboken: Wiley, 2016. – 712 p.
3. Montgomery D. C., Jennings C. L., Kulahci M. *Introduction to Time Series Analysis and Forecasting*. – 3rd ed. – Hoboken: Wiley, 2021. – 650 p.
4. Shumway R. H., Stoffer D. S. *Time Series Analysis and Its Applications: With R Examples*. – 4th ed. – Cham: Springer, 2017. – 562 p.
5. Lutkepohl H. *Applied Time Series Econometrics*. – 2nd ed. – Cambridge: Cambridge University Press, 2019. – 356 p.
6. Мельник О. І., Кузьмін О. Є. *Аналітика великих даних у медицині: методи та технології*. – Львів: Видавництво ЛНУ, 2020. – 248 с.
7. Гриценко В. С., Бондаренко А. М. *Математичне моделювання в системах охорони здоров'я*. – Київ: КНЕУ, 2018. – 232 с.
8. European Centre for Disease Prevention and Control (ECDC). *Annual Epidemiological Report for 2019*. – Stockholm: ECDC, 2020. – 156 p.
9. World Health Organization (WHO). *World Health Statistics 2022: Monitoring health for the SDGs*. – Geneva: WHO Press, 2022. – 144 p.
10. European Commission, Eurostat. *Health statistics – database 2023*. – Luxembourg: Publications Office of the European Union, 2023. – 120 p.

ВИКОРИСТАННЯ ВЕЛИКИХ МОВНИХ МОДЕЛЕЙ (LLM) ДЛЯ ОПТИМІЗАЦІЇ БІЗНЕС-ПРОЦЕСІВ ПІДПРИЄМСТВА

Мойсеєнко В.М., victor.moyseyenko@gmail.com, **Антоненко С.В.**

Дніпровський національний університет імені Олеся Гончара

Нами було досліджено застосування великих мовних моделей (LLM) в автоматизації та оптимізації бізнес-процесів підприємства. LLM, що характеризуються своєю здатністю до розширеної генерації тексту та контекстного розуміння, надають значну можливість для підвищення операційної ефективності та прийняття рішень. В дослідженні пропонується комплексний аналіз архітектур LLM, включаючи авторегресивні (наприклад, GPT [1]), автокодувальні (наприклад, BERT [2]) та sequence-to-sequence (наприклад, T5) моделі, а також їхні відповідні сильні сторони в різних бізнес-застосунках. Ключові характеристики LLM, такі як кількість параметрів, контекстне вікно та гіперпараметри, розглядаються для з'ясування їхнього впливу на продуктивність моделі.

Крім того, у дослідженні оцінюються різні стратегії розгортання LLM для підприємств, що охоплюють LLM як послугу (LLM as a Service), впровадження з відкритим кодом та розгортання на основі ресурсів, детально описуючи їхні відповідні переваги та недоліки щодо вартості, контролю та безпеки даних.

Критичний компонент цього дослідження зосереджений на передових методах підвищення продуктивності LLM у конкретних бізнес-завданнях. До них належать генерація з доповненим пошуком (RAG), яка інтегрує зовнішні джерела знань для підвищення точності; виклик функцій, що забезпечує взаємодію із зовнішніми API для виконання складних завдань; та тонке налаштування, яке узгоджує попередньо навчені моделі з даними, специфічними для предметної області.

В дослідженні розглянуті можливість та ефективність застосування різних LLM в процесах банківської діяльності: в ІТ (генерації коду та документації, виявленні вразливостей та автоматизації тестування програмного забезпечення), автоматизації підтримки користувачів в контакт-центрах (впровадження інтелектуальних чат-ботів та віртуальних асистентів) та ризик-менеджменті (запобігання фроду, моделювання кредитоспроможності, аналіз ринкової інформації тощо).

Результати дослідження підкреслюють трансформаційний потенціал LLM у оптимізації бізнес-операцій. Дослідження також висвітлює необхідність ретельного врахування проблем, притаманних широкому впровадженню LLM в процесах організацій, включаючи витрати на розробку та експлуатацію, алгоритмічні упередження, етичні наслідки та вразливості безпеки. Стратегічне осмислення впровадження, використання таких методів, як RAG, виклик функцій та точне налаштування, має вирішальне значення для максимізації переваг LLM та одночасного зменшення пов'язаних з ними ризиків. Це дослідження сприяє глибшому розумінню практичного застосування LLM у корпоративних умовах та забезпечує основу для прийняття обґрунтованих рішень щодо їх впровадження та інтеграції.

ПОСИЛАННЯ

1. OpenAI. ChatGPT [Електронний ресурс]. – 2023. – Режим доступу: <https://openai.com/blog/chatgpt>
2. Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, 1, 4171–4186. – DOI: <https://doi.org/10.18653/v1/N19-1423>

ЗБАГАЧЕННЯ БАЗИСУ В МЕТОДІ ФУНДАМЕНТАЛЬНИХ РОЗВ'ЯЗКІВ ПОТЕНЦІАЛАМИ ПОДВІЙНОГО ШАРУ З В-СПЛАЙНОВОЮ ЩІЛЬНІСТЮ

Молчанов А.О., evandr73@gmail.com

Дніпровський національний університет ім. О. Гончара

Розглядається спосіб відновлення потенціальних фізичних полів, який називають метод фундаментальних розв'язків (MFS). Основна ідея методу полягає в тому, що гармонічна функція, яка описує потенціальне поле, апроксимується комбінацією лінійно незалежних розв'язків рівняння Лапласа:

$$\tilde{u}(x) = \alpha_0 + \sum_{i=1}^m \alpha_i \Phi_i(x), \quad (1)$$

де $\Phi_i(x) = \Phi(x - y_i)$, з точками сингулярності y_i , обраними навколо тої області, в якій потрібно відновити поле. Для 2D випадку $\Phi(x) = -\frac{1}{2\pi} \ln|x|$. Для визначення невідомих коефіцієнтів $\alpha = (\alpha_0, \dots, \alpha_m)$ у рівняння (1) підставляємо n відомих значень поля, що визначені (виміряні) в граничних вузлах $\{x_1, x_2, \dots, x_n\}$. Отриману систему $n \times (m+1)$ лінійних рівнянь розв'язуємо в сенсі мінімізації суми квадратів відхилень. Така система є некоректною за Адамаром й для отримання адекватного розв'язку вимагає певного методу регуляризації.

У випадку, коли маємо задачу з розривними умовами на границі:

$$\begin{cases} \Delta u = 0, x \in \Omega \\ u = g_0(x), x \in \Gamma_0, \\ u = g_1(x), x \in \Gamma_1 \end{cases} \quad (2)$$

де $\Gamma = \Gamma_0 \cup \Gamma_1$ - границя області Ω , й на кінцях цих сегментів маємо розриви, вищезгаданий метод не є підходящим, оскільки в (1) всі функції $\Phi_i(x)$ є аналітичними (в заданій області) й погано пристосовані для задач з

неаналітичними граничними умовами. Для подолання вказаних труднощів пропонується [1] лінійну комбінацію (1) доповнити функціями спеціального виду:

$$\tilde{y}(x) = \alpha_0 + \sum_{i=1}^m \alpha_i \Phi_i(x) + \sum_{r=0}^p \beta_r \Psi_r(x), \quad (3)$$

де остання сума – лінійна комбінація функцій, кожна з яких є потенціалом подвійного шару (double layer potential):

$$\Psi_r(x) = \int_{\Gamma_1} \partial_{\nu_y} \Phi(x-y) y^r ds_y, \quad (4)$$

ν_y - одиничний вектор нормалі в точках $y \in \Gamma$, що вказує назовні області Ω . Такі функції є гармонічними й аналітичними всередині Ω й їх сліди (trace) на ділянці границі Γ_1 є поліноміальними функціями x^r , $r = \overline{0, p}$. В [2] розглянуто застосування цього способу із степенями до $p = 18$.

Тепер ми пропонуємо в (4) замість поліноміальних функцій використати b -сплайни 3 степеня, носії яких належать ділянці Γ_1 :

$$\Psi_r(x) = \int_{\Gamma_1} \partial_{\nu_y} \Phi(x-y) b_{i,3}(y) ds_y, \quad (5)$$

Розглянуто сплайни з рівномірним розташуванням вузлів та із експоненційним згущенням вузлів до кінців ділянки Γ_1 . Для рівномірного розташування вузлів розроблено як чисельний розрахунок дискретизацією інтегралу так і аналітичне представлення функцій виду (5).

Бібліографічні посилання

1. Carlos J.S. Alves, Svilen S. Valtchev, Applied Mathematics and Computation. 320(3): 61-74 (2018).
2. Molchanov A.O. Problems of applied mathematics and mathematical modeling. Volume 24: 134-140 (2024).

МОДИФІКОВАНИЙ АЛГОРИТМ WINDOW-EMD ІЗ ННТ ДЛЯ РОЗПІЗНАВАННЯ КОРИСНИХ ТА ШУМОВИХ КОМПОНЕНТ У МОВНИХ АУДИОСИГНАЛАХ

Мороз В.В., Кулик Д.В., daniil.kulyk@stud.onu.edu.ua

Одеський національний університет імені І.І.Мечнікова

У роботі досліджується підхід до виділення корисних голосових компонент в реальних шумових умовах (офісні розмови, механічні шуми, клавіатурні клацання). Складність полягає у змінності та нестационарності фону, що маскує мовний сигнал і ускладнює подальший аналіз.

Запропоновано віконний варіант емпіричного модового розкладання (Window-EMD)[1] з подальшим аналізом у просторі Гільберта-Хуанга (ННТ)[2]. Така комбінація дає адаптивне розділення сигналу на внутрішні модові функції (IMF) та миттєві частоти, що дозволяє відокремлювати мовні й шумові складники та надалі застосовувати статистичні критерії для автоматизованого пригнічення шуму.

У роботі розглянуто низку класичних методів для знешумлення аудіосигналів. Тестування проведено на чотирьох реальних аудіозаписах із комбінованими шумами (фонові голоси, гвинт гелікоптера, реактивний двигун, клавіатура). Оцінювання за допомогою SNR та експертним прослуховуванням для виявлення суб'єктивних втрат якості.

Класичні методи виявилися нестійкими до нестационарного та змішаного шуму, що мотивує використання ННТ-підходів. Для подолання поставлених обмежень застосовано перетворення Гільберта-Хуанга (ННТ), яке поєднує EMD та миттєвий частотний аналіз. Отримані внутрішні модові функції (IMFs) дали змогу виокремити низько- та високочастотні компоненти сигналу, що дало змогу проводити адаптивну фільтрацію шумів. Для автоматичного визначення шумових мод використано F-тест[3], який визначає статистично незначущі коливання. Просте застосування віконного EMD та F-test показали незначні зміни SNR для

складних записів та мали великий час обчислення. Відмінним результатом стало введення модифікацій: віконний EMD з автоматичним розпізнавання ознак IMFs (енергетичні, частотні характеристики) та адаптивною фільтрацією. Такий підхід підвищив стабільність видалення фонових шумових компонент при збереженні характеристик голосу переднього плану.

Порівняння з моделлю на основі нейронних мереж (TensorFlow CNN)[4] показало, що глибокі мережі можуть досягати кращих значень по SNR, але вимагають суттєвих обчислювальних ресурсів на початку, часу на навчання та великої кількості анотованих датасетів.

У результаті дослідження запропонована комбінація модифікованого Window-EMD з ННТ з ознаками IMF (енергетика, частота) та адаптивною фільтрацією підвищує стабільність видалення фону без помітної втрати розбірливості голосу. При цьому потребує значно менше ресурсів, ніж навчання глибокої нейронної мережі. Виявлені недоліки: часткова втрата енергії сигналу під час Spectral Gating та чутливість параметрів алгоритму. Це вказує на необхідність подальшої оптимізації та адаптації порогових критеріїв.

Надалі планується поєднати ННТ-обробку з попереднім відбором/стисненням ознак (зокрема, через вейвлет-перетворення), а також із частотними та енергетичними фільтрами й легкими моделями машинного навчання.

Список літератури: 1. Calvo M.S. et al. Enhanced CEEMDAN with superior noise assistance for signal analysis. Engineering Reports (Wiley), 2024. 2. Ahmed A. et al. EMD-based feature extraction for environmental sound classification. Sensors, 2022. 3. Bian, Xihui & others (2022). Spectral denoising based on Hilbert–Huang transform combined with F-test. Frontiers in Chemistry. 4. Nagajyothi Dimmita & P. Siddaiah. (2018). Speech Recognition Using Convolutional Neural Networks.

НЕЙРОМЕРЕЖЕВІ МЕТОДИ ПІДВИЩЕННЯ ТОЧНОСТІ ТА ШВИДКОДІЇ МСЕ У ЗАДАЧАХ КОНТАКТНОЇ МЕХАНІКИ

Морозов Ю.С., yury.morozov@gmail.com, Зайцева Т.А., ztan2004@ukr.net

Дніпровський національний університет імені Олеся Гончара (www.dnu.dp.ua)

В останні роки спостерігається активний розвиток підходів, що поєднують класичні чисельні методи з алгоритмами глибокого навчання. Зокрема, Physics-Informed Neural Networks (PINNs) інтегрують фундаментальні фізичні закони безпосередньо у процес навчання нейронної мережі. Завдяки цьому PINNs здатні з достатньою точністю відтворювати розподіли напружень і деформацій навіть за обмежених або частково відсутніх даних, забезпечуючи прийнятну узгодженість отриманих результатів із фізичними принципами. Нейронні оператори забезпечують встановлення залежності між характеристиками матеріалів, граничними умовами та реакцією системи без необхідності традиційної дискретизації, що значно підвищує швидкодію обчислень. Згорткові нейронні мережі (CNN), навчання яких здійснено на даних, отриманих за допомогою МСЕ (метода скінчених елементів), здатні оперативно прогнозувати розподіли напружень і переміщень за нових умов. Такі моделі істотно скорочують обчислювальні витрати, забезпечуючи при цьому задовільну точність [1].

Гібридний підхід, який поєднує традиційний МСЕ з нейронною мережею прямого поширення (Feedforward Neural Network) також здатний підвищити ефективність розрахунків. Основна ідея полягає у використанні результатів МСЕ, отриманих на спрощеній сітці, та апостеріорної оцінки похибки, обчисленої методом Zienkiewicz–Zhu (ZZ-метод), для подальшого уточнення розподілу напружень за допомогою нейронної мережі [2]. Процес реалізації можна поділити на такі етапи:

1. Підготовка даних, де проводяться MSE-розрахунки на спрощеній і уточненій сітках та обчислення апостеріорних похибок. В результаті формується велика кількість навчальних пар, що включають напруження та похибки з одного боку та уточнені напруження з іншого.

2. Навчання мережі, де відбувається тренування моделі на великій кількості прикладів (понад 5000) для узагальнення закономірностей між грубою та уточненою сітками. На вхід подаються розв'язки на грубій сітці та похибки, вихідні дані порівнюють з розв'язками на уточненій сітці.

3. Застосування, коли відбувається прогнозування уточнених напружень навченою нейромережею у цільових точках для нових задач без необхідності ремешингу або повторного навчання.

В результаті модель глибокого навчання відтворює напруження, наближені до результатів розрахунків на високодеталізованій сітці, при скороченні обчислювальних ресурсів. Такий підхід дозволяє зменшити час моделювання, зберігаючи фізичну достовірність і придатність для інженерних застосувань. Таким чином, інтеграція MSE з алгоритмами глибокого навчання відкриває нові можливості для моделювання контактних задач, оптимізації конструкцій та розв'язання багатомасштабних проблем у механіці деформівного твердого тіла.

Бібліографічні посилання

1. Deshpande, S., Sosa, R.I., Bordas, S.P., & Lengiewicz, J. Convolution, aggregation and attention based deep neural networks for accelerating simulations in mechanics. *Frontiers in Materials*, 2022.
<https://doi.org/10.3389/fmats.2023.1128954>
2. Oishi A., Yagawa G. Finite elements using neural networks and a posteriori error. *Arch. Comput. Methods Eng.* 28, 3433-3456, 2021. <https://doi.org/10.1007/s11831-020-09507-0>

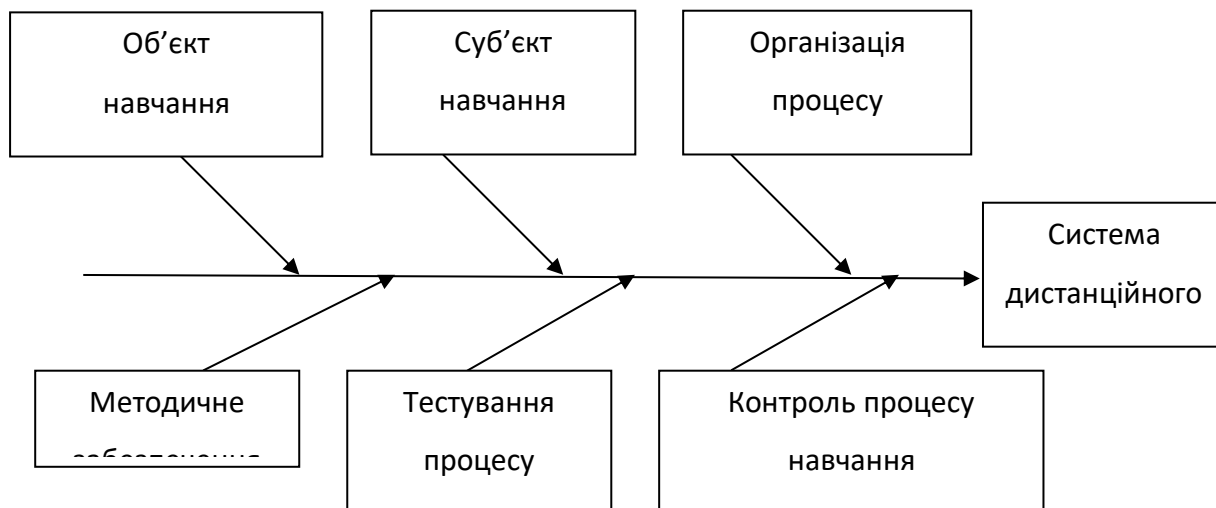
ПОБУДОВА МОДЕЛІ СТРУКТУРИ ТА ОПИСУ СИСТЕМИ ДИСТАНЦІЙНОГО НАВЧАННЯ

Наконечна Т.В., naktanya1961@gmail.com

Дніпровський національний університет імені О. Гончара

У зв'язку із зовнішніми обставинами, в яких знаходиться Україна, в організації навчання на всіх рівнях освіти відбулися значні зміни. До цих змін відноситься частий терміновий перехід дистанційного навчання з допоміжної форми навчання до основної. Тому потрібно неперервно вдосконалювати процес дистанційного навчання на основі використання системного підходу і моделювання. Дистанційна форма навчання базується на використанні традиційних та інноваційних методів і засобів, методичного забезпечення, заснованих на інформаційно-телекомунікаційних технологіях, що реалізують інтерактивну взаємодію учасників навчального процесу з метою досягнення здобувачами відповідного освітньо-кваліфікаційного рівня.

Для обґрунтування інформаційної технології підтримки процесів прийняття рішень щодо планування навчання фахівців розглянемо дослідження складових системи дистанційного навчання (СДН). До структурних складових математичних моделей ДН можна віднести моделі: об'єкта навчання, викладача, тестування, процесу навчання, СДН. Для зручності роботи можна почати з побудови каузальної моделі процесу навчання у виді діаграми Ісікави:



Ще одним важливим підходом до моделювання інформаційної діяльності об'єкта навчання в СДН є контроль вірності співвідношення

$$I = I_0 + I_1 + I_2 \leq \bar{I},$$

де I – сумарний інформаційний потік, I_0 – складова потоку вхідних даних, який прямує до об'єкту навчання без попередньої обробки, I_1 – оперативна інформація, яка є тією частиною вхідного інформаційного потоку, що прямує крізь обчислювальну систему, I_2 – нормативно-довідкова інформація, $\bar{I}$ – граничний допустимий потік інформації для об'єкта.

В процесі дослідження були отримано: *модель викладача* (побудовану на основі теорії нечітких множин); *модель тестування* (яка була представлена як сукупність трьох моделей: модель системи тестів, модель системи тестування, модель об'єкта навчання); *модель процесу дистанційного навчання* (що може бути записана таким чином $МПДН = \{B, H, СДН\}$, де B – множинне число викладачів, H – множинне число тих, хто навчається, СДН – система дистанційного навчання).

Таким чином, в результаті дослідження запропонованої моделі процесів дистанційного навчання встановлено, що застосування ДН при підготовці фахівців дозволяє підвищити ефективність процесу підготовки за рахунок забезпечення індивідуального підходу до кожного з об'єктів навчання та застосуванню системних модернізованих методів та засобів професійного навчання.

Список використаної літератури

1. Антофій, Н.Н. (2003). *Моделі і методи інформаційної підтримки у комп'ютеризованих системах навчання*. Херсон: Вид-во ХНТУ.
2. Исикава, К. (1988). *Японские методы управления качеством*. М.: Экономика,

ДОСЛІДЖЕННЯ МЕТОДІВ ПЕРЕДАЧІ ДАНИХ ЧЕРЕЗ ІНТЕРНЕТ У РОЗПОДІЛЕНИХ СИСТЕМАХ З МОБІЛЬНИМ КЛІЄНТОМ

Новіков С.О., supnovaaa@gmail.com, Антоненко С.В., Ізмайлова М.К.

Дніпровський національний університет імені Олеся Гончара

У сучасних умовах розвитку інформаційних технологій розподілені системи є базовим компонентом інфраструктури більшості веб- та мобільних застосунків. Одним із ключових аспектів їхньої ефективності є методи передачі даних між сервісами та клієнтами через Інтернет. Від вибору технології залежить швидкодія, стабільність і масштабованість усієї системи. Метою роботи є дослідження, реалізація та порівняння сучасних методів передачі даних між застосунками в умовах багатокомпонентної розподіленої системи з головним Android-клієнтом. Для реалізації було спроектовано архітектуру з шести застосунків, що взаємодіють між собою через різні протоколи та технології: REST, SOAP, WebSocket, FTP, GraphQL, gRPC і ETL.

REST застосовується для стандартних CRUD-операцій та інтеграції з базою даних. SOAP забезпечує формалізований XML-обмін із зовнішніми системами. WebSocket відповідає за передачу даних у реальному часі, а FTP — за завантаження великих файлів. gRPC використовується для швидкої передачі аналітичних даних у двійковому форматі. GraphQL слугує уніфікованим шлюзом для REST і SOAP, а ETL — автоматизує міжсерверний обмін даними.

Для перевірки ефективності кожного методу було проведено перевірку роботи системи в Docker-середовищі, де Android-застосунок виступав клієнтом, що отримував дані через різні протоколи. Результати тестування (10 МБ даних, середні 100 запитів) наведено у таблиці 1.

Результати показали, що gRPC є найефективнішим методом з точки зору швидкодії та стабільності, тоді як GraphQL і REST забезпечують

найкращий баланс між гнучкістю та простотою інтеграції. WebSocket демонструє найменшу затримку при постійних з'єднаннях, що робить його незамінним у застосунках із живими оновленнями. У свою чергу, ETL-процес підтвердив ефективність для автоматичної агрегації та міжсерверного обміну, що дозволяє реалізувати повний цикл обробки даних — від збору до збереження результатів на FTP-сервері.

Таблиця 1 — Результат порівняння методів передачі даних

Метод передачі	Середня затримка (мс)	Пропускна здатність (МБ/с)	Споживання пам'яті (МБ)	Надійність (%)
REST (JSON)	230	4.2	85	98
SOAP (XML)	410	2.9	105	97
WebSocket	95	6.8	90	99
FTP	160	8.1	75	98
GraphQL	210	4.8	92	99
gRPC	55	10.4	70	99
ETL (міжсервер)	180	7.2	78	100

Таким чином, реалізована архітектура довела можливість сумісного використання REST, SOAP, GraphQL, WebSocket, gRPC, FTP і ETL у межах єдиної системи. Проведене порівняння показало, що gRPC має найвищу швидкодію, тоді як REST і GraphQL залишаються найзручнішими для інтеграції мобільних клієнтів. Використання ETL дозволяє централізувати аналітичні процеси та підвищити узгодженість даних.

Отримані результати можуть бути використані при розробці корпоративних, аналітичних та мобільних систем, що вимагають стабільного обміну даними у розподіленому середовищі.

МАТЕМАТИЧНЕ МОДЕЛЮВАННЯ ТА АВТОМАТИЗОВАНА СИСТЕМА РОЗРАХУНКУ НАВЧАЛЬНОГО НАВАНТАЖЕННЯ ЗАКЛАДУ ВИЩОЇ ОСВІТИ

Овсов М.В., ovsov.m19@365.dnu.edu.ua, **Верба О.В.**,
verba@365.dnu.edu.ua, **Сафронова І. А.**, safronova_i@365.dnu.edu.ua,
Зайцева Т.А., zaitseva_t@365.dnu.edu.ua

Дніпровський національний університет імені Олеся Гончара

Актуальність роботи зумовлена необхідністю комплексної автоматизації процесів планування та управління освітнім процесом в умовах переходу закладів вищої освіти (ЗВО) на електронний документообіг. Ефективне формування навчального навантаження є критичним елементом організації освітньої діяльності, що безпосередньо впливає на раціональне використання кадрових та матеріальних ресурсів ЗВО. Автоматизація цього процесу дозволяє оптимізувати розподіл ресурсів, мінімізувати вплив людського фактора та істотно підвищити якість управління освітнім процесом [1, 2].

Метою роботи є розробка математичної моделі та створення комп'ютерної системи автоматизованого розрахунку навчального навантаження для закладу вищої освіти на прикладі Дніпровського національного університету імені Олеся Гончара.

Розроблена математична модель враховує ключові параметри навчальних планів, актуальний контингент студентів, форми здобуття освіти, типи аудиторних занять, а також чинні нормативи навантаження на викладачів [3].

Програмну реалізацію виконано мовою Python (у середовищі Microsoft Visual Studio), що забезпечує її гнучкість та кросплатформність. Для створення графічного інтерфейсу користувача (GUI) використано бібліотеку PySimpleGUI, а для ефективної роботи з файлами навчальних планів – Openpyxl, для управління локальною базою даних – SQLite3.

Розроблена система забезпечує комплексний функціонал, включаючи: автоматизований облік та обробку великих обсягів даних з навчальних планів; формування деталізованої звітності про навчальне навантаження за кафедрами, курсами та формами навчання; аналіз ефективності використання викладацьких ставок; гнучке моделювання різних сценаріїв розподілу навантаження шляхом зміни нормативних параметрів.

Практичну апробацію системи проведено на великому масиві даних (148 навчальних планів) бакалаврського та магістерського рівнів ДНУ, що підтвердило коректність математичної моделі, високу функціональність програмного забезпечення та ергономічність користувацького інтерфейсу.

Запропоноване рішення забезпечує комплексну автоматизацію процесів планування, обліку та аналізу навчального навантаження, сприяючи підвищенню ефективності управління освітнім процесом та раціональному використанню кадрових ресурсів. Модульна архітектура та використання відкритих стандартів створюють передумови для легкої адаптації системи в інших закладах вищої освіти з урахуванням їхніх специфічних вимог та нормативів.

Бібліографічні посилання

1. Баженов В. А., Лізунов П. П., Білощицький А. О., Білощицька С. В., Палій С. В. Математична модель планування та моніторингу обсягів навчальної роботи викладачів і студентів вищих навчальних закладів, *Управління розвитком складних систем*, 2014, Вип. 18, С. 148–161.
2. Григорович В. Інформаційні системи розрахунку навчального навантаження та розподілу штатів ВНЗ, *Вісник Національного університету "Львівська політехніка". Комп'ютерні науки та інформаційні технології*, 2016, № 843, С. 94–103.
3. Овсов М. В., Верба О. В. Моделювання системи розрахунку навчального навантаження закладу вищої освіти, *Математичне та програмне забезпечення інтелектуальних систем: матеріали конф.* (Дніпро, 20–22 листоп. 2024 р.), Дніпро, 2024, С. 216.

ПРОЄКТУВАННЯ ТА РЕАЛІЗАЦІЯ СЕРВЕРНОЇ ЧАСТИНИ ВЕБПЛАТФОРМИ ДЛЯ ПІДТРИМКИ НАВЧАЛЬНОГО ПРОЦЕСУ

Овчаренко Д.О., Михальчук Г.Й., ovcharenko.denys3110@gmail.com

Дніпровський національний університет імені Олеся Гончара (www.dnu.dp.ua)

Робота присвячена розробленню серверної частини вебплатформи для підтримки навчального процесу. Система забезпечує ефективну взаємодію між тьюторами та студентами, підтримує керування курсами, завданнями, навчальними матеріалами, а також обмін повідомленнями та сповіщеннями.

Основною технологічною платформою обрано Spring Boot, що забезпечує швидке створення, конфігурацію та розгортання серверних застосунків. Завдяки автоматичній конфігурації та тісній інтеграції з іншими компонентами екосистеми Spring, цей фреймворк дозволяє уникнути надлишкових налаштувань і підвищує продуктивність розроблення.

Для реалізації RESTful вебсервісів застосовано Spring Web MVC [1], що забезпечує взаємодію між клієнтською частиною (веб- або мобільним інтерфейсом) і сервером. Для спрощення розроблення рівня доступу до даних застосовано Spring Data JPA [2] – компонент сімейства Spring Data, що реалізує роботу з Java Persistence API (JPA). Механізми автентифікації та авторизації реалізовано за допомогою Spring Security [3], який підтримує різні стратегії безпеки та легко інтегрується з іншими компонентами Spring.

У структурі застосунку реалізовано такі основні модулі:

- автентифікація та авторизація користувачів;
- керування навчальними матеріалами та завданнями;
- збереження й обробка даних у базі даних;
- підтримка ролей користувачів – студент, тьютор.

Базу даних розроблено на основі MySQL з дотриманням принципів нормалізації, що забезпечує ефективне зберігання та швидкий доступ до інформації про курси, завдання, оцінки та користувачів.

Взаємодію між клієнтською частиною та сервером реалізовано через REST API, який відповідає сучасним стандартам безпеки та узгодженості даних. Авторизація користувачів здійснюється за допомогою JWT-токенів, що гарантує безпечну передачу інформації між клієнтом і сервером.

Проведене тестування підтвердило коректність роботи сервісів, правильність обробки запитів і стабільність системи у різних сценаріях використання, включно з випадками некоректних даних та повторних запитів.

Отримані результати демонструють можливість створення функціональної, безпечної та гнучкої серверної частини навчальної вебплатформи на базі сучасних технологій Java та Spring Framework. Розроблене рішення може бути інтегроване з веб- або мобільними клієнтами, забезпечуючи масштабовану підтримку навчального процесу.

Бібліографічні посилання

1. Web MVC Framework. URL: <https://docs.spring.io/spring-framework/docs/3.2.x/spring-framework-reference/html/mvc.html>
2. Spring Data JPA. URL: <https://spring.io/projects/spring-data-jpa>
3. Spring Security. URL: <https://spring.io/projects/spring-security>

ПІДХІД ДО ОЦІНЮВАННЯ ПОВЕДІНКОВИХ РИЗИКІВ КОРИСТУВАЧІВ З ПТСР ЗА ЦИФРОВИМ ПРОФІЛЕМ ЗАСОБАМИ NLP

Овчарук О.М., off4aruk@gmail.com,
Мазурець О.В., exe.chong@gmail.com
Хмельницький національний університет

Актуальність роботи зумовлена необхідністю раннього виявлення посттравматичних станів і супутніх ризиків у середовищах цифрової комунікації [1]. Класичні психодіагностичні інструменти не забезпечують потрібної оперативності та масштабованості, тоді як автоматизований аналіз користувацьких текстів дозволяє виявляти латентні патерни, асоційовані з ПТСР та потенційно небезпечною поведінкою [2]. Метою дослідження є побудова інтерпретовного нейромережевого підходу до інтегральної оцінки поведінкових ризиків на основі семантико-психолінгвістичного профілю користувача, сформованого з його цифрових текстових повідомлень.

Запропонований підхід реалізує багатокрокову обробку повідомлень з поєднанням трьох компонентів: детекція PTSD-контенту трансформерною моделлю; виявлення коморбідних ментальних станів у кожному повідомленні; аналіз емоційного контексту з фокусом на негативних реакціях. Далі ці висновки агрегуються у часовому вимірі в інтегральний індекс ризику з ковзним вікном, що відображає динаміку змін стану користувача. На рівні моделювання використовуються сучасні BERT-подібні архітектури: для PTSD-виявлення «microsoft/deberta-v3-small» (точність понад 90%), для коморбідних розладів п'ять спеціалізованих DistilBERT-моделей (з точністю не нижче 85%), для емоцій «roberta-base-go_emotions». Емоційний аналіз орієнтовано на виявлення негативної тональності як підсилювального фактора ризику.

Компонент коморбідностей репрезентує п'ятивимірний вектор показників «anger», «anxiety», «depression», «narcissistic», «panic», що оцінюється окремими донавченими трансформерними класифікаторами. Це забезпечує більш повну картину психоемоційного фону і зменшує ризик помилок, спричинених змішуванням симптомів у монолейблових моделях.

Експериментальна частина виконувалась у середовищі Google Colaboratory з використанням бібліотеки Transformers і фреймворку PyTorch; код організовано у вигляді Jupyter-ноутбука з фіксацією версій пакетів та конфігурації, що підвищує відтворюваність. Такий уніфікований стек спростив ітеративний аналіз часової динаміки ризику з урахуванням PTSD-контенту, коморбідностей та емоцій.

Методика поєднує точкові (рівень повідомлень) та агреговані (рівень користувачів) оцінки в єдиний часовий профіль, що дозволяє розрізняти флуктуації та стійкі тренди й визначати доменно обґрунтовані пороги реагування. Для узгодження апостеріорних імовірностей можливе байєсівське чи ізотонічне калібрування, а чутливісний аналіз порогів активації буст-факторів забезпечує контроль хибнопозитивних рішень. Схема також дає змогу розширювати індикатори за рахунок темпоральних і графових ознак, підвищуючи стійкість до доменних зсувів.

Обмеження стосуються чутливості до якості даних і платформа-специфічних артефактів. Наявність класового дисбалансу і доменних зсувів потребує систематичного застосування методів балансування, контролю витоків і регулярної повторної калібровки. Крім того, етичні аспекти обробки користувацьких даних вимагають процедур анонімізації, мінімізації даних і ведення аудит-логів доступу до пояснень моделі. Попри це, отримані метрики та пояснюваність свідчать про достатню робастність підходу в умовах обмежених обчислювальних ресурсів і варіативності даних.

Перелік посилань:

1. Юрченко Д. Ю., Овчарук О. М., Мазурець О. В., Шевчук П. О. Метод використання нейромережі гібридної архітектури для визначення емоційної тональності текстових повідомлень. Міжнародний науково-технічний журнал «Вимірювальна та обчислювальна техніка в технологічних процесах». 2025. № 2. С. 136–141. URL: <https://doi.org/10.31891/2219-9365-2025-82-18> (дата звернення: 30.10.2025).
2. Овчарук О., Мазурець О. Нейромережеве діагностування проявів птср у текстовому контенті з використанням помилко-орієнтованого навчального набору даних. Herald of Khmelnytskyi National University. Technical sciences. 2024. Т. 343, № 6(1). С. 195–200. URL: <https://doi.org/10.31891/2307-5732-2024-343-6-30> (дата звернення: 30.10.2025).

АВТОМАТИЗОВАНЕ РОЗТАШУВАННЯ МОДЕЛЕЙ НА ЦИЛІНДРИЧНІЙ ПЛАТФОРМІ 3D-ПРИНТЕРА ТИПУ FUGO

Орлов С.К., cetcat18@gmail.com

Наконечна Т.В., naktanya@ukr.net

Дніпровський національний університет імені О. Гончара

У системах циліндричного 3D-друку, таких як принтер Fugo[1], ефективність процесу значною мірою залежить від способу розташування моделей на внутрішній поверхні барабана. На відміну від традиційних планарних принтерів, де платформа має площину XY, у циліндричній системі робоча зона — це криволінійна поверхня, а напрямок нарощування матеріалу збігається з радіусом барабана. Тому оптимальне розташування моделей має враховувати радіальну геометрію, рівномірність навантаження та траєкторію лазера.

У ручному режимі компонування моделей виникають проблеми: нераціональне використання площі, нерівномірне балансування мас і перекриття траєкторій при скануванні. Автоматизована система розміщення дозволяє усунути ці недоліки, забезпечуючи максимальну продуктивність, стабільність обертання та якість друку (рис. 1).

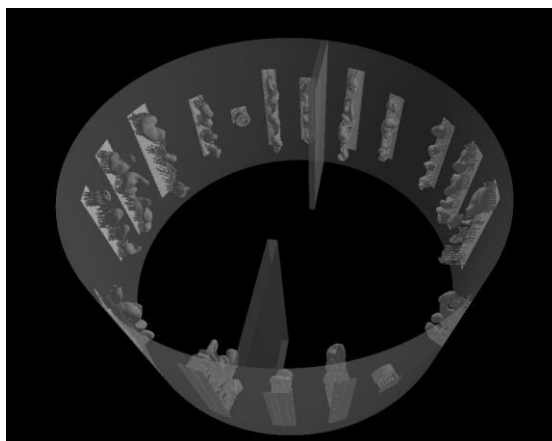


Рис.1 Приклад оптимального розташування моделей в циліндричному барабані

Основні задачі автоматизованого розташування:

1. **Максимізація використання площі друку.** Алгоритм має рівномірно заповнювати внутрішню поверхню циліндра моделями, мінімізуючи “мертві зони” та відстані між об’єктами.

2. **Балансування мас системи.** Оскільки барабан обертається на високих швидкостях, асиметричне розташування моделей може спричинити вібрації та дисбаланс. Автоматичний аналіз моменту інерції забезпечує симетричне завантаження відносно осі обертання. Оптимально збалансована система може вважатися такою, де центр мас лежить на уявній осі циліндра.

Розглядаються два потенційних метода при побудові алгоритму оптимального розташування моделей:

1. «Розгортання» циліндричної платформи в прямокутну форму і використання вже відомих алгоритмів прямокутного пакування [2][3], де кожна з моделей представляється прямокутником, який є проекцією описаного мінімального паралелепіпеда навколо моделі на площину розгорнутої циліндричної платформи. Перевагами цього алгоритму є простота реалізації та лінійна його складність. Недоліком можна вважати те, що така розстановка не завжди є оптимальною для ефективного балансування моделей всередині циліндра який обертається.

2. Алгоритм, який базується на розрахунку оптимального балансування системи – тобто знаходження такого положення моделей, при якому центр мас лежить на осі циліндра (або максимально наближений до нього). Тут можна розглянути різні алгоритми оптимізації цільової функції, такі як методи градієнту або методи дискретної оптимізації.

Бібліографічні посилання

1. Орлов, С. & Наконечна, Т. Сучасні методи адитивного виробництва тривимірних об’єктів на основі принтера типу FUGO: матеріали Міжнародної науково-практичної конференції, на тему «Математичне та програмне забезпечення інтелектуальних систем» (с. 318-319). 22-24 листопада, 2023, Дніпро, Україна.
2. Dowsland, K. A., & Dowsland, W. B. (1992). Packing problems. *European Journal of Operational Research*, 56 (1), 2–14
3. Huang, E., & Korf, R. E. (2010). Optimal rectangle packing on non-square benchmarks. In *AAAI’10: Proceedings of the 24th National Conference on Artificial intelligence*, pp. 317–324. AAAI Press

ПРОЄКТ ІНФОРМАЦІЙНОЇ СИСТЕМИ ДЛЯ АВТОМАТИЗОВАНОГО УПРАВЛІННЯ ЗАМОВЛЕННЯМИ ТА ДОСТАВКИ ФЕРМЕРСЬКОЇ ПРОДУКЦІЇ

Павлів М. Я., maxvill2016@gmail.com

Басюк Т. М., taras.m.basyuk@lpnu.ua

Національний університет «Львівська політехніка»

Аграрний сектор України переживає етап глибоких змін і потребує нових підходів до організації постачання продукції [1]. Складні логістичні умови, зростання витрат і дефіцит ресурсів роблять процес управління замовленнями та доставкою особливо вразливим. Запровадження автоматизованої інформаційної системи може стати дієвим інструментом підвищення ефективності фермерських господарств і забезпечення стабільності продовольчого ланцюга [2].

Метою дослідження є створення інформаційної системи, що забезпечує єдину платформу для прийому, обробки та доставки замовлень фермерської продукції. Така система сприяє прозорості процесів, зниженню ризиків помилок і скороченню часу між замовленням та його виконанням. На діаграмі варіантів використання (рис. 1) відображено основних учасників процесу — клієнта, менеджера, логіста та водія — і логіку їх взаємодії всередині системи.

Автоматизація управління замовленнями дозволяє зменшити витрати на транспортування, підвищити рівень сервісу для споживачів і створити умови для розширення ринку збуту фермерських господарств. Система забезпечує миттєву перевірку наявності продукції, формує оптимальні маршрути доставки з урахуванням географічних і часових факторів, а також дозволяє клієнтам відстежувати статус замовлення в реальному часі.

Очікуваним результатом упровадження є підвищення конкурентоспроможності малих і середніх фермерських підприємств, зниження операційних витрат та покращення доступу споживачів до якісної української продукції. Реалізація таких цифрових рішень є кроком

до створення сучасної екосистеми аграрної логістики, заснованої на інноваціях, довірі та сталому розвитку.

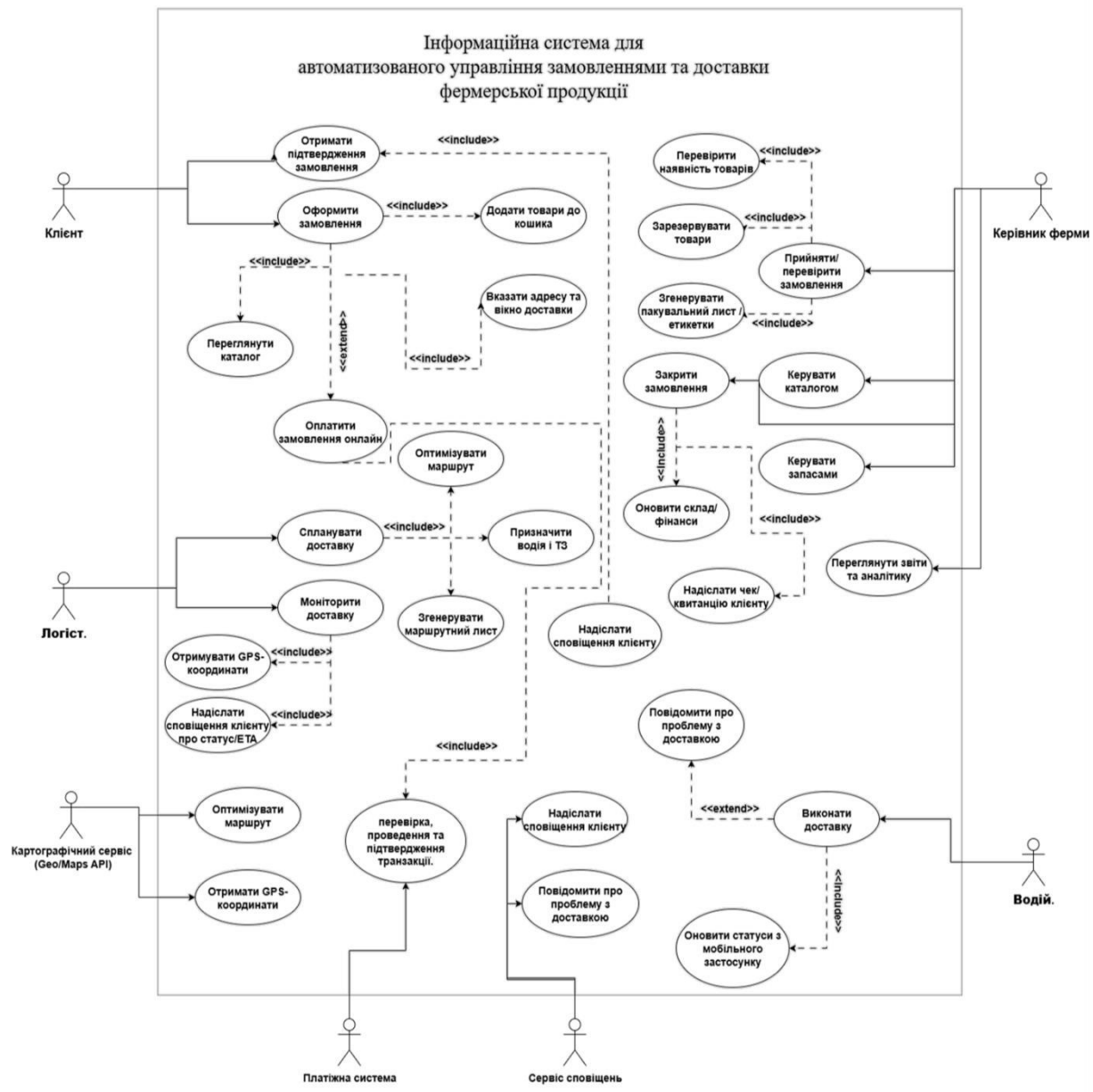


Рис. 1. Діаграма варіантів використання

Список використаних джерел

1. Міністерство аграрної політики та продовольства України. Офіційна статистика аграрного сектору [Електронний ресурс]. – Режим доступу: <https://minagro.gov.ua> (дата звернення: 08.07.2025).
2. Deloitte Ukraine. Agribusiness in Ukraine: Trends and Forecast 2024. – Київ : Deloitte, 2024. – 28 с. – Режим доступу: <https://www2.deloitte.com/ua> (дата звернення: 08.07.2025)

МУЛЬТИАГЕНТНА СИСТЕМА З RAG ДЛЯ МОДЕРАЦІЇ КОРИСТУВАЦЬКИХ ПОВІДОМЛЕНЬ

Павлюк Д.І., coster730@gmail.com

Байбуз О.Г., baibuz_o@365.dnu.edu.ua

Дніпровський національний університет імені Олеся Гончара

Вступ та аналіз досліджень. Зі зростанням обсягів користувацьких повідомлень у соцмережах їх модерація ускладнюється через різноманітність даних, множинність порушень та недостатню точність традиційних автоматизованих рішень. У дослідженнях, проведених у контексті аналізу настрою на основі великих мовних моделей, встановлено, що LLM більшого розміру зазвичай забезпечують вищу точність результатів порівняно з меншими моделями та спеціалізованими нейронними мережами, хоча й потребують значно більших обчислювальних ресурсів [1]. Водночас компактні моделі, тонко налаштовані для конкретного предметного домену, можуть демонструвати конкурентні або кращі результати [2]. Ефективність таких моделей може бути додатково підвищена методами, зокрема мажоритарного голосування [3].

Архітектура рішення. Модерація користувацьких повідомлень вимагає одночасно невисоких обчислювальних витрат (через велику кількість вхідних даних) та високої точності аналізу настрою для мінімізації участі модератора у процесі. Значна варіативність контенту та критеріїв модерації ускладнює застосування класичного *fine-tuning* мовних моделей. Альтернативним рішенням є *Retrieval-Augmented Generation (RAG)*, що дозволяє моделі використовувати попередньо підготовлену базу знань домену для відповіді на запити. Для оптимізації модерації застосовується мультиагентний підхід, який працює як послідовний ланцюг: Автоматичне отримання повідомлень з платформи через API (без залучення агентів) → Превалідація малою моделлю за базовими правилами (токсичність, погрози, реклама), що відсікає очевидні порушення →

*Виділення сутностей та базовий аналіз настрою, у тому числі щодо конкретних об'єктів → Агент із великим контекстним вікном та базою знань, який порівнює отриманий настрій з очікуваним, зберігає невідомі сутності та при необхідності передає запит далі → Агент з доступом до мережі (без бази знань), що уточнює інформацію про сутності у разі нестачі контексту. **Основні очікувані результати роботи системи:** повідомлення проходить перевірку; фільтрується з високою впевненістю; вимагає уваги модератора через низьку впевненість; потребує розширення контексту та втручання модератора. **Переваги:** зменшення ручної перевірки та фокус модератора на базі знань і складних випадках. **Недоліки:** складність підтримки мультиагентної інфраструктури, необхідність підготовки та оновлення бази знань.*

Висновок: мультиагентна система з RAG може суттєво автоматизувати модерацію та аналіз настрою, особливо для великих медіаплатформ з високою активністю користувачів, хоча потребує значних зусиль у розробці й підтримці інфраструктури.

Список використаних джерел:

1. Andrew Li. A case study of sentiment analysis on survey data using LLMs versus dedicated neural networks. Northeast High School Journal of Science. 2025. URL: <https://nhsjs.com/2025/a-case-study-of-sentiment-analysis-on-survey-data-using-llms-versus-dedicated-neural-networks/> (дата звернення: 29.10.2025).
2. Zhang, W., Deng, Y., Liu, B., Pan, S., Bing, L. Sentiment Analysis in the Era of Large Language Models: A Reality Check. Findings of the Association for Computational Linguistics: NAACL 2024. Mexico City, Mexico : Association for Computational Linguistics, 2024, P. 3881–3906. DOI: 10.18653/v1/2024.findings-naacl.246. URL: <https://aclanthology.org/2024.findings-naacl.246/> (дата звернення: 29.10.2025).
3. Junichiro Niimi. Dynamic Sentiment Analysis with Local Large Language Models using Majority Voting: A Study on Factors Affecting Restaurant Evaluation. 2024. URL: <https://arxiv.org/html/2407.13069v1> (дата звернення: 29.10.2025).

ІНТЕЛЕКТУАЛЬНА НОРМАЛІЗАЦІЯ ПОВІДОМЛЕНЬ У МУЛЬТИФОРМАТНИХ ПОДІЄВИХ СИСТЕМАХ

Панасюк Б.Ю., boris.panasyuk@gmail.com

Ліщинська Л.Б., llb@vntu.edu.ua

Вінницький Національний Технічний Університет

Подієві системи у мікросервісних архітектурах паралельно використовують Avro, JSON Schema та Protocol Buffers, що породжує несумісності на рівні типів, еволюції схем і семантики полів. Канонічна модель даних зменшує кількість парних інтеграцій, але не гарантує семантичну еквівалентність між форматами [1]. З'явилися структурні конвертори на кшталт Avrotize, які транслюють схеми через єдину інтеграційну модель, однак працюють переважно з синтаксисом, залишаючи невирішеними неоднозначності значень та одиниць виміру [2]. Інженерні кейси демонструють, що змішане використання JSON, Avro і Protobuf масштабно ускладнює підтримку та веде до ручних конверторів [3]. Для цифрових двійників критичним є семантично узгоджений потік подій, інакше погіршуються аналітика та симуляції.

Метою є підвищення коректності автоматичних конверсій між Avro, JSON Schema та Protobuf без ручних правок шляхом поєднання канонічної проміжної моделі даних з доменною онтологією, чого раніше не було реалізовано як цілісний процес із формальними метриками у подієвих системах.

Запропоновано метод канонічної нормалізації складається з двох опор: канонічна проміжна модель як універсальне внутрішнє представлення та онтологічний шар як семантичний посередник. Крок 1: проєкція схеми і екземпляра повідомлення у канонічну модель з фіксацією типів, варіантів, обов'язковості, обмежень і логічних типів, що узгоджується з принципами Canonical Data Model та практиками структурної конвертації [1; 2]. Крок 2: застосування онтології домену для

вирівнювання назв, структур та одиниць виміру, зняття неоднозначностей і доповнення залежностей, як в підходах семантичної інтеграції інформаційних потоків для промислових сценаріїв і цифрових двійників. Після семантичного узгодження канонічне представлення проєктується у цільовий формат із дотриманням правил еволюції, якість контролюється семантичними перевірками та двосторонньою round-trip валідацією. Очікується приріст Success@0-edit відносно структурних конверторів і стабілізація втрат при двосторонньому перетворенні на рівні, близькому до нульового, що зменшує ручні втручання у виробничих конвеєрах [2; 3].

Отже, поєднання канонічної проміжної моделі з доменною онтологією забезпечує семантичну інтероперабельність між Avro, JSON Schema та Protobuf у подієвих системах, підвищує частку коректних автоматичних конверсій і спрощує інтеграцію потоків даних у цифрових двійниках, перевершуючи суто структурні трансляції [1; 2; 3].

Список використаної літератури:

[1] Hohpe G., Woolf B. Enterprise Integration Patterns: Designing, Building, and Deploying Messaging Solutions. Boston: Addison-Wesley, 2004. 736 p.

[2] Vasters C. Avrotize: convert data structure definitions between schema formats using Apache Avro Schema as integration model [Електронний ресурс]. GitHub, 2023. Режим доступу:

<https://github.com/clemensv/avrotize>

(дата звернення: 30.10.2025).

[3] Clear Street. Switching to Protobuf (from Avro) on Kafka [Електронний ресурс]. 2023. Режим доступу: <https://www.clearstreet.io/news/blog/switching-to-protobuf-from-avro-on-kafka> (дата звернення: 30.10.2025).

**ПРОЄКТУВАННЯ ЦИФРОВИХ МОДУЛЬНИХ СИСТЕМ
ДЛЯ АВТОМАТИЗОВАНОГО КОНТРОЛЮ
ТЕХНОЛОГІЧНИХ ПРОЦЕСІВ**

Панкратов Є.В., evgennomad@gmail.com

НТУУ КІП ім. Ігоря Сікорського, ПБФ

Сучасне промислове виробництво та технологічні процеси дедалі більше потребують високого рівня точності, надійності та оперативності управління, що зумовлює необхідність застосування цифрових систем автоматизації. Традиційні централізовані системи часто не забезпечують достатньої гнучкості та масштабованості, особливо у виробництвах із різноманітними технологічними операціями та швидко змінними умовами. Модульний підхід до побудови систем автоматизації дозволяє інтегрувати різні компоненти – датчики, виконавчі пристрої, контролери, програмні модулі та інтерфейси взаємодії – у єдину архітектурну структуру, при цьому кожен модуль може функціонувати автономно та бути легко модернізованим або замінюваним без порушення роботи всієї системи [1].

Запропонована архітектура модульної цифрової системи управління технологічними процесами передбачає чіткий розподіл функцій між сенсорними, виконавчими та аналітичними блоками. Алгоритми управління реалізують адаптивне регулювання параметрів технологічних процесів у реальному часі, забезпечують послідовне та паралельне виконання операцій і автоматичне коригування дій при зміні зовнішніх або внутрішніх умов виробництва. Таке рішення дозволяє підвищити точність контролю технологічних параметрів, скоротити час реакції системи на відхилення та зменшити ризик помилок.

Інтеграція модульної архітектури з цифровими алгоритмами управління дозволяє підвищити ефективність управління, забезпечує зручність масштабування та інтеграції нових технологій, а також спрощує модернізацію наявного обладнання. Крім того, модульна структура сприяє

підвищенню надійності та стійкості системи до відмов окремих компонентів, що робить її придатною для використання у промислових середовищах з високими вимогами до безперервності процесу.

Особливу увагу приділено можливості інтеграції штучного інтелекту та методів машинного навчання для прогнозування параметрів технологічних процесів та оптимізації використання ресурсів. Використання алгоритмів адаптивного управління дозволяє передбачати потенційні відхилення і автоматично коригувати параметри процесу, що сприяє підвищенню продуктивності, зменшенню витрат енергії та матеріалів, а також забезпечує більш ефективну експлуатацію обладнання. Модульна цифрова система управління може бути адаптована для різних галузей промисловості, включаючи хімічну, харчову, металургійну, машинобудівну, енергетичну та інші, де потрібне інтегроване, гнучке та високоточне управління технологічними процесами.

Таким чином, поєднання модульної архітектури, цифрових алгоритмів управління та можливостей інтеграції інтелектуальних технологій створює основу для побудови сучасних вискоєфективних систем автоматизації. Використання таких рішень забезпечує не лише підвищення точності та надійності управління, але й відкриває нові перспективи для розвитку промислових підприємств за рахунок швидкої адаптації до змін технологічних процесів, впровадження інноваційних методів контролю та оптимізації. Застосування модульного підходу дозволяє забезпечити баланс між гнучкістю, масштабованістю та ефективністю, що робить запропоновану систему перспективним рішенням для сучасного промислового виробництва.

Література

1. Панкратов Є. В, N.D.Pankratova, H.S.Tymchyk “*Strategy of the cyber-physical system for a small business enterprise guaranteed functioning with the digital twin support.*”, *System Research And Information Technologies* No. 2 (2024) page 7 – 20. *Scopus*. DOI:10.20535/SRIT.2308-8893.2024.2.01

ДИНАМІКА РОЗВИТКУ ІНФОРМАЦІЙНИХ ТЕХНОЛОГІЙ І СТВОРЕННЯ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ

Пасічник А.М., Пасічник В.А., panukr977@gmail.com

Дніпровський державний технічний університет

Визначальною особливістю математизації та інформатизації сучасної науки стало створення нових можливостей для розвитку інформаційно-комунікаційних технологій (ІКТ) і систем автоматизації цілого ряду аспектів інтелектуальної діяльності. Функціонування ІКТ стало одним із основних компонентів науково-методичної та технологічної бази інформаційної індустрії яка поєднує математичні методи, інформаційні ресурси і програмно-технічні засоби для виробництва, передачі, обробки та використання інформації. Необхідно також зазначити, що розвиток ІКТ став потужним каталізатором подальшого розвитку математичних методів і алгоритмів.

Процес розвитку інформаційних технологій розпочинається з кінця 60-х років минулого сторіччя і умовно можна виділити такі основні етапи їх практичного застосування на шляху інформатизації суспільства:

1-й етап (1970-80 рр.). Характерною особливістю даного періоду є обмеженість потужностей обчислювальної техніки недостатня для обробки значних обсягів даних.

2-й етап (1980-90 р.). Характеризується появою великих електронно-обчислювальних машин (ЕОМ) типу IBM 360 / 370 та широким впровадженням комплексних інформаційних систем управління різними галузями економічної діяльності. Особливістю даного етапу є певна невідповідність можливостей апаратно-технічних засобів ЕОМ рівню розвитку програмного забезпечення.

3-й етап (1990-2000 рр.) – комп'ютерна техніка стає більш доступною для непрофесійного користувача, а інформаційні системи – засобом підтримки та прийняття управлінських рішень. Основні проблемні

питання ефективності застосування пов'язані з недостатнім рівнем оптимізації отриманих рішень, та недосконалістю користувацького інтерфейсу з комп'ютерним середовищем.

4-й етап (2000-2010 рр.) – запровадження сучасних інформаційно-комунікаційних технологій та систем розподіленої обробки даних. Основні питання стосуються розробки стандартів та протоколів для комп'ютерного зв'язку; забезпечення авторизованого доступу та захисту інформації.

5-й етап (2010-2020 рр.) – характеризується поширенням застосуванням алгоритмів і програм обробки мовних сигналів та графічних зображень, появою систем машинного навчання та робото-технічних систем, впроваджуються інформаційні інтелектуальні системи з використанням достатньо ефективних алгоритмів оптимізації генерованих рішень.

6-й сучасний етап із 2020 р. характеризується інтенсифікацією розробки та практичного використання досконалих алгоритмів оптимізації і відповідного програмного забезпечення інтелектуальних програмно-технічних систем. На даний час такі системи забезпечують виконання завдань для цілого ряду видів діяльності на високо-професійному рівні і класифікуються як системи штучного інтелекту (ШТІН).

Ефективність застосування ІКТ створило у ряді експертів ілюзію щодо настання «технологічної сингулярності» вже у 2025-29 рр., коли рівень ШТІН може досягнути інтелекту людини. Зазначимо, що такі прогнози не враховують той факт, що на даний час не існує ефективних математичних моделей та алгоритмів інтелектуальної діяльності людини. Тому, в сучасних умовах системи ШТІН, забезпечуючи неймовірну швидкість оптимізаційної обробки інформації за заданим алгоритмом, ще на достатньо тривалий час будуть залишатись ефективним інструментом проведення наукових досліджень і вирішення прикладних проблем.

НЕЙРОМЕРЕЖЕВЕ МОДЕЛЮВАННЯ ПРОЦЕСУ КОМП'ЮТЕРНОЇ ГРИ

Пашченко А.Ю., andrew.pashchenko2002@gmail.com, Книш Л.І.

Дніпровський національний університет імені Олеся Гончара

Прогрес в області штучного інтелекту та машинного навчання відкриває нові можливості для автоматизації та оптимізації складних процесів, включаючи моделювання інтелектуальної поведінки в імітованих середовищах, таких як комп'ютерні ігри. Нейронні мережі, будучи ключовим інструментом глибокого навчання, створюють складні нелінійні взаємозв'язки і приймають рішення на основі великого обсягу даних або доказів. Мета роботи полягає в створеній моделі гравця, який за допомогою нейронної мережі навчається грі в Тетріс, демонструючи стратегічну поведінку, спрямовану на максимальну продуктивність гри. Розробка реалізована за використанням мови програмування Python та її бібліотек для машинного навчання.

Тетріс - це гра з повною інформацією, що вимагає від гравця швидкого прийняття рішень у динамічному середовищі. На кожному кроці гравець повинен вибрати оптимальне положення, орієнтацію та точку приземлення поточної фігури.

Для моделювання ігрового процесу буде використано архітектуру нейронної мережі, яка виконуватиме роль оцінної функції (Q-функції) у контексті глибокого навчання з підкріпленням, зокрема алгоритму Deep Q-Network.

Основні елементи алгоритму це:

- вхідні дані (Стан): нейронна мережа прийматиме як вхідні дані поточний стан ігрового поля (наприклад, у вигляді двовимірного масиву), а також характеристики поточної та наступної фігури;

- вихідні дані (Дії): мережа виведе оцінки (Q-значення) для всіх можливих дій, які можна виконати з поточною формою (наприклад, переміщення ліворуч/праворуч, поворот, скидання);
- функція винагороди: нагорода розроблена так, щоб мотивувати програму очистити максимальну кількість рядків та підтримувати поле максимально рівним. Наприклад, вища нагорода за очищення чотирьох рядів одночасно («Тетріс»).

В результаті навчання нейронна мережа розробляє ефективну ігрову стратегію, яка значно перевершує випадковий набір дій. Критеріями успіху були обрані рахунок та кількість очищених ліній за гру. Результати дослідження демонструють ефективність глибокого навчання із підкріпленням як потужного підходу до моделювання інтелектуальних ігрових процесів. Це підтверджує здатність нейронних мереж освоювати складні стратегії прийняття рішень у динамічних дискретних середовищах. Ця робота може бути основою для подальших досліджень у галузі оптимізації стратегій.

РЕАЛІЗАЦІЯ ОБМІНУ ДАНИМИ МІЖ МОБІЛЬНИМ ДОДАТКОМ ТА ЕЛЕКТРОННИМ ТОНОМЕТРОМ

Перемітько М.В., mikhailperemitko@gmail.com,
Надригайло Т.Ж., ntatiana62@gmail.com
Дніпровський державний технічний університет

У сучасному програмному забезпеченні важливе місце посідають технології бездротової взаємодії між мобільними застосунками та широким спектром електронних пристроїв, таких як носимі гаджети, ваги, побутова техніка та навіть поштомати. Технологія передачі даних Bluetooth Low Energy (BLE) забезпечує стабільну комунікацію між пристроями за мінімального енергоспоживання та високої швидкості обміну даними. Було розроблено програмне забезпечення для взаємодії мобільного застосунку з електронним тонометром з метою зчитування показників артеріального тиску та частоти серцевих скорочень. Розроблений мобільний застосунок створено на платформі React Native, яка дає змогу створювати кросплатформні мобільні додатки із використанням JavaScript та React.

Робота системи починається з ініціалізації Bluetooth-модуля та отримання необхідних дозволів користувача на бездротову передачу даних. Після цього додаток запускає процес сканування навколишніх BLE-пристроїв і визначає серед них потрібний тонометр за ідентифікатором або назвою моделі. Після виявлення пристрою здійснюється підключення через унікальний ідентифікатор, а також пошук доступних сервісів і характеристик, які визначають канали обміну даними між пристроями. У разі успішного з'єднання мобільний застосунок підписується на характеристику, через яку тонометр надсилає дані вимірювань у вигляді коротких пакетів, що кодуються у форматі Base64.

Для обробки отриманих даних було реалізовано функцію, що отримує повідомлення з характеристики пристрою, декодує їх із Base64 у послідовність байтів і зчитує дані відповідно до визначеної структури

пакета. Кожне повідомлення містить службовий заголовок, тип пакета, значення систолічного та діастолічного тиску, частоту пульсу й контрольну суму. На початковому етапі перевіряється правильність заголовка пакету, що позначає початок нового повідомлення. Після цього виконується інтерпретація решти байтів — залежно від типу пакета (проміжні дані, завершення вимірювання тощо).

Кожен числовий параметр зчитується з двох байтів і перетворюється у ціле число без знаку. Для цього застосовуються стандартні методи роботи з буферами, що враховують порядок байтів у протоколі (наприклад, `readUInt16BE`). Після декодування формується структурований об'єкт із полями, що містять систолічний та діастолічний тиск, частоту серцевих скорочень і час вимірювання. Якщо під час перевірки контрольної суми виявлено помилку або пакет має неповну довжину, він відкидається.

Функція отримання даних від тонометра працює в асинхронному режимі, постійно реагуючи на надходження нових повідомлень. Вона забезпечує буферизацію у випадках, коли дані надходять фрагментами, об'єднуючи їх у повноцінне повідомлення перед передачею на обробку. У разі розриву з'єднання реалізовано механізм автоматичного повторного підключення, який дозволяє швидко відновити обмін даними без втручання користувача.

Архітектура додатку є модульною, тобто логіка з'єднання, обробки та збереження даних розділена між незалежними компонентами. Модуль керування BLE-підключенням функціонує окремо від користувацького інтерфейсу, що полегшує масштабування рішення та потенційне вбудовування в інші додатки. Отже, розроблене програмне забезпечення може бути використане як основа для більш розширеної системи домашнього контролю артеріального тиску та превенції серцево-судинних захворювань.

МЕТОД СПРЯЖЕНОГО СТАНУ В АНАЛІЗІ ОБЕРНЕНИХ ЗАДАЧ ДЛЯ БАГАТОШАРОВИХ ОСНОВ

Петров А.Д., @gmail.com,

Трофімов О.В., atrof2222@gmail.com.

Дніпровський національний університет імені Олеся Гончара

У роботі запропонована методика отримання компонент градієнта функції нев'язки в обернених задачах визначення невідомих фізичних та геометричних характеристик пружного багатошарового пакету. Суть метода полягає у певній параметризації відображень, що використовуються для формування регулярної сітки для кожного шару пакету. Для обчислення градієнту функції нев'язки застосовано метод спряженого стану. В якості «змінної стану» («фазової змінної») розглянемо вектор переміщень точок пружного багатошарового пакету (рис. 1).

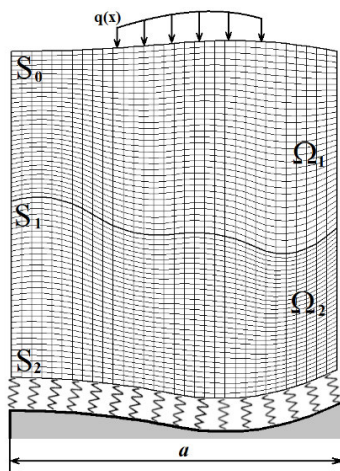


Рис. 1. Пружний багатошаровий пакет із криволінійними межами між шарами з побудованою регулярною сіткою.

У роботі запропоновано блочно-параметричний підхід до розв'язання обернених завдань у рамках «інформаційно-імовірнісної» концепції, який був реалізований у задачах визначення невідомих фізичних та геометричних параметрів пружного шаруватого пакета за заданими переміщеннями деяких точок його верхньої (навантаженої) поверхні.

Виявлено, що задачею, спряженою до сформульованої прямої задачі, буде задача про дію зосередженої одиничної сили, прикладеної в точці поверхні, де задане спостережуване значення променевого переміщення, і діючої вздовж вектору променю сканування, на початкове слоїсте тіло (розташоване на початковій же пружній основі). Тобто, по суті, розв'язок є функцією Гріна (фундаментальним розв'язком) для прямої задачі.

Розроблена методологія дозволить розробляти програмне забезпечення для нового покоління для високошвидкісних пристроїв, що використовуються для неруйнівного контролю та моніторингу технічного стану транспортних конструкцій, таких як дорожнє полотно, злітно-посадкові смуги, залізничні баластні призми тощо (так звані рухомі вимірювальні платформи).

Бібліографічні посилання

1. Trofimov A.V. Block-parametric approach for determining the geometric characteristics of elastic layered packages. *Comp. Appl. Math.* 44, 273 (2025).
<https://doi.org/10.1007/s40314-025-03223-w>
2. Trofimov A.V. Adjoint state method in the block-parametric approach to the inverse problem analysis for elastic layered packages. *Mech. Res. Comm.* 147, 104426 (2025).
<https://doi.org/10.1016/j.mechrescom.2025.104426>
3. Trofimov A.V. Block-parametric approach to non-destructive control data analysis of railroad layered foundations. *MATEC Web of Conferences* 390, 04003 (2024).
<https://doi.org/10.1051/matecconf/202439004003>

РОЛЬ МАТЕМАТИЧНОГО ЗАБЕЗПЕЧЕННЯ ПРИ ПРОЕКТУВАННІ МЕМС

Пилипенко О.О., Шишканова Г.А., Коротунова О.В.

kiko2005ty@gmail.com, shganna@ukr.net, okorotunova@gmail.com

Національний університет «Запорізька політехніка»

Розглянемо приклад проектування МЕМС-резонатора (мікроелектромеханічної системи), де математика та інтелектуальні системи відіграють ключову роль. МЕМС-резонатори – це мікроскопічні структури, які коливаються на певній резонансній частоті і використовуються в годинниках, фільтрах радіочастот та датчиках. Проектування МЕМС-резонатора часто зводиться до аналізу простої моделі "маса-пружина-демпфер" (MSD), що описує механічні коливання.

Диференціальне рівняння коливань для такої моделі описується лінійним диференціальним рівнянням другого порядку:

$$m \cdot \frac{d^2 x}{dt^2} + b \cdot \frac{dx}{dt} + kx = F(t),$$

де x – зміщення маси від положення рівноваги; m – ефективна маса коливної структури; b – коефіцієнт демпфування (втрати енергії, спричинені тертям, зокрема повітря); k – ефективна жорсткість пружини (підвісу); $F(t)$ – зовнішня сила, що діє на масу (наприклад, електростатична сила збудження).

Для заданої чутливості та стабільності необхідно, щоб резонатор мав точну резонансну частоту f_{target} . Для цього потрібно точно розрахувати параметри жорсткості k (залежить від розмірів, геометрії та матеріалу пружини) та маси m (залежить від розмірів та густини матеріалу).

Для збудження коливань часто використовується електростатична сила між двома електродами (рухомим і нерухомим). Сила F_e між паралельними пластинами (ідеалізована модель) з площею перекриття A на відстані g при прикладеній напрузі V та ємності C дорівнює:

$$F_e = -\frac{1}{2} \frac{\partial C}{\partial x} \cdot V^2.$$

Сила, що діє на рухомий електрод є нелінійною і ускладнює диференціальне рівняння коливань, вносячи додаткову (електростатичну) жорсткість.

Після встановлення математичної моделі, наступний крок – оптимізація геометричних параметрів (довжина, ширина, товщина резонатора, зазор між електродами) для досягнення потрібної частоти f_{target} та високого коефіцієнта добротності Q .

Застосуємо генетичний алгоритм для пошуку найкращої комбінації геометричних параметрів, які мінімізують помилку між бажаною та розрахованою частотами, і максимізують Q -фактор.

Процес генетичного алгоритму ітеративно генерує популяцію конструкцій (наборів параметрів), оцінює їхню пристосованість за допомогою математичних моделей (розв'язуючи диференціальні рівняння), застосовує оператори схрещування та мутації, відбираючи найкращі "особини" для наступного покоління.

Результат сходиться до набору геометричних параметрів, які забезпечують бажані f та Q з високою точністю, що значно прискорює процес проектування порівняно з ручним підбором.

Тобто проектування МЕМС – це приклад мультидисциплінарної інженерії, де диференціальні рівняння забезпечують точний аналіз механічної та електростатичної поведінки, а інтелектуальні системи (як-от генетичні алгоритми) ефективно керують складним, багатопараметричним процесом оптимізації для досягнення необхідних характеристик пристрою.

В подальшому будуть цікавими проведення дослідження впливу нелінійності електростатичної сили при великих амплітудах на стабільність та роботу резонатора, використовуючи нелінійні диференціальні рівняння, що застосовують математичний апарат нелінійного аналізу та хаосу.

ПРОГНОЗУВАННЯ ПРИБУТКУ ВІД КРЕДИТНИХ ПРОДУКТІВ НА ОСНОВІ МЕТОДІВ МАШИННОГО НАВЧАННЯ

Подвинцев В.В., podvin.2004@gmail.com

Козакова Н.Л., kozakova.natali@gmail.com

Дніпровський національний університет імені Олеся Гончара

У роботі розглянуто задачу прогнозування прибутковості кредитних продуктів банківських установ на піврічний період із використанням методів машинного навчання. Запропоновано алгоритмічний підхід, що дозволяє підвищити точність прогнозування фінансових результатів і сприяти ухваленню обґрунтованих управлінських рішень у банківській сфері. Реалізовано програмний інструмент на мові Python, який може бути інтегрований у системи підтримки прийняття рішень для фінансового менеджменту.

У сучасних умовах фінансова діяльність банківських установ безпосередньо залежить від ефективності управління кредитними портфелями. Висока конкуренція на ринку кредитування потребує впровадження аналітичних інструментів, які здатні точно прогнозувати прибутковість продуктів, враховуючи велику кількість змінних, невизначеність та нелінійні взаємозв'язки у даних. Методи машинного навчання дозволяють автоматизувати процес аналізу клієнтської поведінки та зменшити ризики, пов'язані з неправильними фінансовими рішеннями. Об'єктом дослідження є поведінка користувачів кредитних продуктів як складова фінансової діяльності банків, а також її вплив на формування прибутку в короткостроковому періоді. Метою роботи є розроблення ефективного алгоритму прогнозування прибутку від кредитування на піврічний період та створення програмного інструменту, який можна застосовувати у системах підтримки прийняття управлінських рішень.

У процесі роботи застосовано методи машинного навчання, статистичного аналізу та прогнозування. Програму було реалізовано мовою

Python у середовищі Google Colab. Модель прогнозування фінансового показника p180 побудовано з використанням бібліотек pandas, numpy, matplotlib, seaborn, scikit-learn та joblib.

Було проведено естування кількох алгоритмів: лінійна регресія – показала низьку ефективність через неоднорідність даних і значну кількість нерелевантних ознак, що ускладнювало навчання моделі. Алгоритм Random Forest продемонстрував кращу точність завдяки структурі з бінарних дерев, що дозволяє ефективно фільтрувати шум і враховувати нелінійні залежності. Для покращення результатів застосовано групування даних за ключовими параметрами, що частково нормалізувало вибірку та зменшило дисперсію похибки.

Отримані результати підтверджують, що випадковий ліс є ефективним інструментом для прогнозування прибутковості кредитних продуктів за умов обмеженого обсягу даних. Модель продемонструвала стабільність і здатність адаптуватися до структурної складності вхідної інформації. Попри це, дослідження виявило обмеження – зокрема, потребу у збільшенні обсягу даних для покращення точності прогнозів. У майбутніх роботах планується впровадження нейронних мереж, що дозволить виявляти складні взаємозв'язки між фінансовими показниками та поведінковими характеристиками користувачів.

ВПЛИВ ЗМАГАЛЬНИХ АТАК НА СТІЙКІСТЬ НЕЙРОННИХ МЕРЕЖ

Політов А.І., politov3601@gmail.com

Вовк С.М., vovk_s_m@ukr.net

Дніпровський національний університет імені Олеся Гончара

У дослідженні проаналізовано вплив змагальних атак (adversarial attacks) на стійкість згорткових нейронних мереж у задачі класифікації зображень на прикладі датасету CIFAR-10[1], що містить 10000 кольорових зображень, розподілених між 10 класами об'єктів (літаки, автомобілі, тварини, кораблі тощо). Атаки моделювалися методом Fast Gradient Sign Method (FGSM)[2], який додає до значень пікселів малоамплітудні модифікації, розраховані за знаком градієнта функції втрат; ці зміни майже непомітні людині, але здатні викликати помилкову класифікацію. Стійкість моделі було підвищено шляхом навчання з атакованими прикладами (adversarial training), коли частина даних під час навчання замінювалася атакованими зображеннями (приблизно половина навчальних прикладів). Такий підхід дає змогу моделі адаптуватися до змагальних впливів і зменшити кількість помилкових класифікацій для різних значень ϵ [3] (амплітуди змагальної атаки, яка визначає її силу).

Без застосування захисту модель демонструвала точність 68.40% на чистих зображеннях. Наведені в таблиці 1 результати показують, наскільки різко змінюється точність під дією FGSM-атак різної сили (параметра ϵ).

Таблиця 1 - Залежність точності моделі від сили атаки

ϵ (FGSM)	Точність без захисту
0.01	48.12%
0.03	21.58%
0.05	8.79%
0.10	1.13%

На рис.1 подано приклад змагальної атаки методом FGSM з параметром $\epsilon=0.01$. На рис.1а подано оригінальне зображення, яке модель вірно класифікує як літак; на рис.1б подано атаковане зображення з майже непомітними для людини змінами, яке модель помилково класифікує як корабель. На рис.1в подано різницю між оригінальним і атакованим зображенням. Отже, навіть за візуально непомітного втручання прогноз мережі змінився, що свідчить про її чутливість до змагальних атак.



Рисунок 1 – Приклад змагальної атаки

Отримані результати підтверджують, що незначні за амплітудою зміни пікселів можуть спричиняти суттєві помилки класифікації. Навчання з атакованими прикладами (adversarial training) частково зменшило вплив таких атак, підвищивши точність моделі з 21.58% до 40.95% при $\epsilon=0.03$, що свідчить про ефективність цього підходу як базового методу захисту.

Список використаних джерел

1. CIFAR-10 and CIFAR-100 datasets. <https://www.cs.toronto.edu/~kriz/cifar.html>
2. Goodfellow, I. J., Shlens, J., & Szegedy, C. (2014). Explaining and harnessing adversarial examples (arXiv:1412.6572). <https://arxiv.org/abs/1412.6572>
3. PyTorch Tutorials. (n.d.). Adversarial example generation (FGSM tutorial). https://docs.pytorch.org/tutorials/beginner/fgsm_tutorial.html

ВАРІАЦІЙНІ МЕХАНІЗМИ АДАПТИВНОГО КЕРУВАННЯ ПРОДУКТИВНІСТЮ В ГІБРИДНИХ СИСТЕМАХ

Попов О.В., alex50ropov@gmail.com, Павлюк А.В., kvadrupl@gmail.com

Інститут кібернетики ім. В.М. Глушкова НАН України

Актуальною задачею сучасних гібридних інтелектуальних систем, що функціонують у змішаному CPU–GPU середовищі, є динамічне балансування продуктивності за умов змінних обчислювальних та інформаційних потоків. Досвід експлуатації великих обчислювальних комплексів показує, що традиційні методи планування — зокрема класичні схеми OpenMP-scheduling або CUDA-stream-balancing — не враховують нелінійну взаємодію підсистем із різними часовими масштабами. Це призводить до локальної нестійкості або неефективного використання ресурсів. Застосування варіаційного підходу дозволяє трактувати зміну продуктивності не як емпіричне налаштування, а як задачу мінімізації узагальненого функціонала енергетичних і часових витрат у просторі станів системи. У наших чисельних експериментах ($N \approx 10^4$ вузлів) такий підхід забезпечував збереження стійкості навіть за різких змін зовнішнього навантаження.

Нехай $u_i(t)$ описує миттєвий рівень завантаження i -го обчислювального ядра, а $\tau_i(t)$ — ефективний часовий масштаб його реакції. Тоді функціонал узагальнених витрат можна записати як:

$$J[u, T] = \int_0^T [\alpha E(u, \dot{u}) + \beta T(u, \tau) + \gamma \Phi(u, \tau, \nabla u)] dt, \quad (1)$$

де $E(u, \dot{u})$ — локальна енергетична щільність обчислень; $T(u, \tau)$ — часовий тензор навантаження системи; $\Phi(u, \tau, \nabla u)$ — оператор взаємодії між підсистемами CPU і GPU.

Варіаційна умова стаціонарності

$$\delta J = 0 \Rightarrow \frac{d}{dt} \left(\frac{\partial \mathcal{L}}{\partial \dot{u}_i} \right) - \frac{\partial \mathcal{L}}{\partial u_i} + \lambda_i \frac{\partial \Phi}{\partial u_i} = 0, \quad (2)$$

визначає динаміку рівноважного розподілу продуктивності між компонентами системи, а параметр λ_i виступає множником Лагранжа.

Для задач реального часу можна сформулювати адаптаційний закон градієнтного типу:

$$\dot{u}_i = -\eta_u \frac{\partial J}{\partial u_i}, \dot{\tau}_i = -\eta_\tau \frac{\partial J}{\partial \tau_i}, \quad (3)$$

де $\eta_u, \eta_\tau > 0$ – коефіцієнти, які задають темп динамічної узгодженості системи.

Зазначений закон забезпечує збіжність до варіаційного екстремуму навіть за змін зовнішніх навантажень. На практиці це проявляється як стабілізація параметрів після короткочасних коливань енергетичних потоків. Чисельну реалізацію подано у вигляді операторної схеми $u_i^{(k+1)} = u_i^{(k)} - \eta_u \nabla_{u_i} J^{(k)}$, що зберігає стійкість при великих розмірах системи ($n \gg 10^5$). У тестах на GPU-кластері (CUDA 12.2) спостерігалася збіжність до екстремуму при збуреннях до 20 % від номінального навантаження.

Рівняння (1)–(3) описують адаптацію як еволюцію у просторі керувальних параметрів, спрямовану на мінімізацію інтегральних втрат продуктивності. Уже після 5–7 ітерацій відзначалася стабілізація середнього енергетичного потоку. Варіаційна форма постановки узгоджує локальну ефективність вузлів із глобальною стійкістю системи та відкриває можливість впровадження адаптивних схем керування продуктивністю в реальному часі, зокрема для балансування CPU–GPU.

Список використаних джерел

1. Khimich, O. M., Popov, O. V., Chistyakov, O. V., et al. (2023). Adaptive Algorithms for Solving Eigenvalue Problems in the Variable Computer Environment of Supercomputers. *Cybernetics and Systems Analysis*, 59, 480–492. <https://doi.org/10.1007/s10559-023-00583-1>

МОДЕЛЮВАННЯ ДИНАМІКИ ПАНДЕМІЙ ЗА ДОПОМОГОЮ АТТРАКТОРІВ І МЕТОДУ SINDy

Пригода О. О., slaughterhouse.leonardod@gmail.com

Золотько К. Є., zolt66@gmail.com

Дніпровський національний університет імені Олеся Гончара

Швидке зростання можливостей комп'ютерного моделювання та аналітичних алгоритмів відкрило нові підходи до дослідження складних біологічних систем. Однією з найважливіших сфер їх застосування є вивчення епідемічних процесів, що набули особливого значення після пандемії COVID 19. Класичні моделі типу SIR допомагають описати загальні тенденції поширення інфекцій, проте не відображають усієї складності нелінійної поведінки системи та її чутливості до початкових умов [2]. Тому перспективним напрямом є використання апарату нелінійної динаміки, зокрема побудови аттракторів та застосування методу SINDy (Sparse Identification of Nonlinear Dynamics), який дозволяє визначати приховані закономірності розвитку епідемій [1].

У роботі використано реальні часові ряди захворюваності на COVID 19 та SARS із відкритих баз даних, зокрема Johns Hopkins University. На етапі попередньої обробки дані нормалізували та згладжували для усунення шуму, що спотворює структуру фазового простору. На основі очищених рядів побудовано фазові портрети, які відтворюють нелінійну структуру процесу та дозволяють виділити стійкі області й режими хаотичної поведінки [3]. Для реконструкції аттракторів застосовано метод затриманих координат, що забезпечує відновлення траєкторій системи у просторі станів навіть за наявності обмежених спостережень.

Подальший етап полягав у використанні методу SINDy для автоматичного виведення системи диференціальних рівнянь, яка описує зміну кількості інфікованих у часі. Алгоритм відбирає лише найсуттєвіші члени моделі, зберігаючи її точність та інтерпретованість [1]. Точність перевіряли шляхом тестування на різних часових відрізках, а середньоквадратична похибка (RMSE) не перевищила 0.08 у нормалізованих одиницях, що свідчить про адекватність моделі.

Порівняння отриманих атракторів для COVID 19 і SARS засвідчило подібність топологічних структур динамічних режимів. Це вказує на наявність універсальних закономірностей поширення вірусів у замкнених популяціях [4]. У фазовому просторі обох систем спостерігається перехід від квазіперіодичних до хаотичних коливань залежно від параметрів, що підтверджує доцільність застосування підходів теорії хаосу для аналізу епідемій.

Поєднання атракторного аналізу з методом SINDy створює компактну, проте інформативну модель, здатну описувати складну поведінку системи з мінімальною кількістю параметрів. Такий підхід дозволяє не лише відтворити динаміку, а й здійснювати короткострокові прогнози поширення захворюваності на основі поточних тенденцій. Виявлено, що для COVID 19 модель демонструє поступове згасання коливань після досягнення певного рівня колективного імунітету, тоді як для SARS динаміка швидше стабілізується та має меншу амплітуду фазових змін [5].

Отримані результати можуть бути використані для вдосконалення прогнозних моделей, оцінки ефективності протиепідемічних заходів і створення систем раннього попередження. Розроблений підхід також дає змогу визначати критичні параметри переходів між стабільними та хаотичними фазами розвитку інфекційного процесу, що є корисним для побудови адаптивних стратегій реагування. Подальші дослідження планується спрямувати на уніфікацію введених параметрів, розширення набору первинних даних і впровадження нових методів прогнозування в межах мультисистемних моделей.

Бібліографічні посилання:

1. Brunton S. L., Proctor J. L., Kutz J. N. Discovering governing equations from data by sparse identification of nonlinear dynamical systems. *PNAS*, 113(15), 2016.
https://www.pnas.org/doi/10.1073/pnas.1517384113?utm_source=chatgpt.com
2. Heesterbeek H., Anderson R. M., et al. Modeling infectious disease dynamics. *Science*, 347(6227), 2015.
3. Datseris G., Wagemakers A. Nonlinear time series analysis and chaos reconstruction. *Chaos Journal*, 2022.
4. Bauch C. T., Earn D. J. D. Transients and attractors in epidemics. *Proceedings of the Royal Society B: Biological Sciences*, 270(1524), 2003.
5. Kutz J. Data-Driven Modeling & Scientific Computation. *Oxford University Press*, 2013.

ПІДВИЩЕННЯ МАСШТАБОВАНOSTI КОНТРОЛЮ ЦІЛІСНОСТІ ФАЙЛІВ У ХМАРНИХ СИСТЕМАХ НА ОСНОВІ BLAKE2B ТА ДЕРЕВА МЕРКЛА

Ріпка Є.В., ripka_e@365.dnu.edu.ua,

Сафронова І.А., safronova_i@365.dnu.edu.ua

Дніпровський національний університет імені Олеся Гончара

У роботі реалізовано та перевірено ефективність комбінованого підходу до контролю цілісності файлів, що поєднує алгоритм хешування BLAKE2b [1] та структуру дерева Меркла [2]. Метою дослідження було визначення доцільності використання цього поєднання для великих обсягів даних у сучасних системах обміну та хмарних сервісах.

Алгоритм BLAKE2b обрано завдяки його високій швидкодії та стійкості до криптографічних атак. Він формує хеш фіксованої довжини 256 біт і підтримує параметризацію, що робить його придатним для роботи з великими файлами [1]. Дерево Меркла, у свою чергу, забезпечує можливість перевірки цілісності окремих блоків даних без повторного обчислення хешу всього файлу, що є критично важливим при роботі з гігабайтними обсягами інформації [2].

Програмна реалізація виконана на мові Java із використанням стандартних бібліотек. Файли ділилися на блоки по 4 МБ, для кожного з яких обчислювався хеш за допомогою алгоритму BLAKE2b, після чого формувалося дерево Меркла з отриманих значень. Для оцінки ефективності було проведено порівняльний аналіз часу виконання двох підходів — повного хешування та блочного з побудовою Merkle Root [3].

Таблиця 1. - Порівняльна характеристика методів.

Розмір файлу	Час Blake2B (мс)	Час Merkle Tree + Blake2B (мс)
1 107 991 байт	33.924	29.117
162 477 байт	13.489	7.309
343 017 147 байт	2390	1347
6.4 ГБ	—	25254

Отримані результати свідчать, що при обробці невеликих файлів обидва методи демонструють подібну швидкодію, однак зі збільшенням розміру даних Merkle-підхід показує значну перевагу. Це пояснюється поетапним обчисленням над блоками, що усуває потребу завантаження всього файлу в оперативну пам'ять [3].

Таким чином, поєднання алгоритму BLAKE2b із деревом Меркла забезпечує ефективний і масштабований механізм контролю цілісності інформації. Отримані результати підтвердили доцільність використання цього підходу у системах з великими обсягами даних, зокрема у P2P-мережах, хмарних сховищах та платформах резервного копіювання.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Aumasson J.-P., Neves S. BLAKE2: simpler, smaller, fast as MD5. *Journal of Cryptographic Engineering*, 2013, 4(2), pp. 167–188.
2. Merkle R. Protocols for public key cryptosystems. *IEEE Symposium on Security and Privacy*, 1980.
3. Ільченко О.В., Лисенко П.П. Продуктивність алгоритму BLAKE2b у системах контролю цілісності великих файлів. *Інформаційна безпека*, 2022, №2, с. 15–22.

РОЗРОБКА СИСТЕМИ МОНІТОРИНГУ DOCKER-КОНТЕЙНЕРІВ

Рябоволенко В.А., Байбуз О.Г.

v.a.ryabovolenko@gmail.com, baibuz_o@365.dnu.edu.ua

Дніпровський національний університет імені Олеся Гончара

У сучасних умовах швидкого розвитку DevOps-підходів та хмарних технологій контейнери стали одним із ключових елементів побудови масштабованих, надійних і незалежних середовищ виконання. Проте із зростанням кількості одночасно запущених контейнерів у системах виникає потреба у керуванні та моніторингу за використанням її ресурсів [1].

Контейнеризація – це технологія віртуалізації на рівні операційної системи, яка дозволяє ізолювати застосунки разом з усіма необхідними залежностями в автономні контейнери. Основна ідея полягає у створенні відокремленого середовища, в якому програмне забезпечення працює однаково незалежно від системи, де воно розгорнуто [2].

У процесі розгортання та експлуатації контейнеризованих систем постає проблема забезпечення стабільності їхньої роботи при змінних навантаженнях. Зі збільшенням кількості контейнерів та складністю взаємозв'язків між ними звичайне спостереження за системними метриками стає неефективним. Необхідність моніторингу полягає не лише у фіксації поточного стану контейнерів, а й у виявленні тенденцій, прогнозуванні перевантажень та прийнятті рішень щодо оптимального розподілу ресурсів. Тому створення системи моніторингу відкриває можливість для реалізації методів автоматизованого перерозподілу навантаження між контейнерами.

Для реалізації системи потрібно інтегрувати Docker API для збору детальної інформації про запущені контейнери. На основі цього створюється Python-сервіс, який реалізує логіку обробки метрик у режимі реального часу. Зокрема, за допомогою бібліотек docker. Дана бібліотека

дозволяє працювати з API Docker, надаючи методи для отримання списку контейнерів (`client.containers.list()`), перегляду їхнього стану (`container.status`), отримання статистики використання ресурсів (`container.stats(stream=False)`), а також виконання команд усередині контейнера (`container.exec_run()`). Завдяки цьому система може збирати повні дані про навантаження на CPU, оперативну пам'ять та мережу в реальному часі [3].

```

=== Моніторинг контейнерів ===

=== Контейнер: confident_driscoll ===
Статус: running | Image: ubuntu:latest
Пам'ять: 7.00 MB (0.04%) | CPU: 0.01%

=== Контейнер: friendly_edison ===
Статус: running | Image: ubuntu:latest
Пам'ять: 0.95 MB (0.01%) | CPU: 0.01%

```

PID	USER	%CPU	%MEM	COMMAND
22	root	50	0	ps
1	root	0	0	bash

PID	USER	%CPU	%MEM	COMMAND
21	root	42.8	0	ps
1	root	0	0	bash

Рисунок 1 – Результати моніторингу системою 2-х контейнерів

В результаті розроблена система надає можливість отримувати актуальні дані про статус контейнерів і їхнє використання системних ресурсів. Це дозволяє своєчасно виявляти перевантаження, оптимізувати використання ресурсів і підвищувати надійність роботи інфраструктури.

Список використаної літератури

1. Applications Integration in a Semi-Virtualized Environment, 2023. URL: <http://surl.li/oyqvuc>
2. Use containers to Build, Share and Run your applications: <https://www.docker.com/resources/what-container/>
3. Docker SDK for Python: <https://docker-py.readthedocs.io/en/stable/>

НОРМАЛІЗАЦІЯ ГУЧНОСТІ АУДІОСИГНАЛУ ДЛЯ СИСТЕМ БІОМЕТРИЧНОЇ ІДЕНТИФІКАЦІЇ ЗА ГОЛОСОМ

Сафіюліна Я. О., anasafiulina@gmail.com,

Степанова Н. І., nist66@gmail.com

Дніпровський національний університет імені Олеся Гончара

Біометричні методи автентифікації посідають важливе місце в сучасних системах безпеки. Серед них ідентифікація за голосом поєднує природність, швидкість і можливість дистанційного використання. Водночас точність таких систем істотно залежить від якості звукового сигналу. Нестабільність гучності, викликана зміною відстані до мікрофона, характеристиками пристрою чи акустикою приміщення, призводить до зміщення енергетичного спектра та погіршення результатів класифікації. Для усунення цього ефекту застосовується RMS-нормалізація, що вирівнює рівень енергії звуку, роблячи ознаки сигналу узгодженими між різними записами.

Обробка звуку здійснюється за допомогою методів спектрального аналізу. Дискретне перетворення Фур'є (DFT) і його оптимізована форма, швидке перетворення Фур'є (FFT), забезпечують перехід від часової до частотної області та дозволяють досліджувати частотну структуру сигналу [1]. Для коротких фрагментів мовлення використовується короткочасне перетворення Фур'є (STFT), яке дозволяє простежити зміни частотного складу у часі [2]. Важливим етапом є обчислення мел-частотних кепстральних коефіцієнтів (MFCC), що моделюють людське сприйняття частоти й широко застосовуються у системах розпізнавання мовлення.

RMS-нормалізація реалізована як процес обчислення середньої енергії звукового сигналу на коротких часових інтервалах із подальшим масштабуванням амплітуди до еталонного рівня. На відміну від пікової нормалізації, яка реагує лише на максимальні значення, RMS-

вирівнювання відображає середню потужність сигналу, більш близьку до реального відчуття гучності [3].

Реалізацію програмного забезпечення виконано мовою Python із використанням бібліотек NumPy, SciPy, librosa, soundfile, PyQt5 та PyQtGraph. Програма забезпечує повний цикл обробки: зчитування звуку, RMS-нормалізацію, спектральний аналіз, розрахунок MFCC і візуалізацію результатів. Візуальний інтерфейс дає змогу спостерігати вплив нормалізації на форму сигналу та спектрограму в реальному часі.

Проведені експерименти показали, що RMS-нормалізація значно зменшує варіативність енергетичних характеристик сигналу між різними записами одного мовця. Після нормалізації спостерігається стабільність рівня гучності та однорідність спектрограми, що безпосередньо підвищує точність розпізнавання. Аналіз коефіцієнтів MFCC підтвердив зменшення розкиду значень, а отже, і кращу узгодженість ознак між окремими аудіофайлами.

Метод виявився обчислювально ефективним і придатним для роботи у потоковому режимі без відчутних затримок. Отримані результати підтверджують доцільність використання RMS-нормалізації як базового етапу попередньої обробки мовних сигналів. Реалізований підхід може бути інтегрований у будь-які системи реального часу, що використовують мовний сигнал як засіб автентифікації, підвищуючи їхню точність і надійність.

Бібліографічні посилання

1. Луценко Х., Нікулін К. Голосова ідентифікація диктора як один із сучасних біометричних методів ідентифікації особи // Теорія та практика судової експертизи і криміналістики. — 2019. — № 19. — С. 239 – 240.
2. Sturnel N., Daudet L. Signal Reconstruction from STFT magnitude: a state of the art // Proc. of the 14th International Conference on Digital Audio Effects – 2011 – С. 2.
3. Jekal, E., Kim, Y., Ku, J., & Park, H. (2025). Software Development and Application for Sound Wave Analysis.

РОЗРОБКА РЕКОМЕНДАЦІЙНОЇ СИСТЕМИ НА ОСНОВІ ПОВЕДІНКОВИХ ТА КОНТЕНТНИХ ОЗНАК

Селімханов Т.Б., seltim22@gmail.com

Дзюба П.А., avatarr@2282gmail.com

Дніпровський національний університет імені Олеся Гончара

Розвиток сучасних інформаційних технологій зумовлює зростання попиту на інтелектуальні системи, здатні забезпечувати персоналізовану взаємодію користувача з цифровими продуктами. Рекомендаційні системи є ключовим інструментом для підвищення якості користувацького досвіду в електронній комерції, медіасервісах та освітніх платформах. Їх основна мета полягає у формуванні точних і релевантних рекомендацій на основі аналізу попередньої поведінки користувача та властивостей контенту.

У даній роботі розглядається розробка рекомендаційної системи на основі поведінкових і контентних ознак з використанням відкритого набору даних MovieLens 100k. Метою дослідження є підвищення точності рекомендацій шляхом поєднання класичних підходів колаборативної фільтрації з ефективними методами пошуку схожості у високорозмірних просторах. Для цього пропонується застосування бібліотеки FAISS (Facebook AI Similarity Search), яка забезпечує швидкий пошук найближчих сусідів на основі попередньо обчислених векторних подань користувачів і фільмів.

Методологія дослідження включає кілька основних етапів:

1. Попередня обробка даних MovieLens 100k – очищення, нормалізація рейтингів, побудова матриці користувач–фільм.
2. Формування векторних ознак – об'єднання поведінкових (історія переглядів, оцінки, активність) та контентних (жанр, рік випуску, актори) характеристик.
3. Індекссування векторів за допомогою FAISS – використання алгоритмів

IndexFlatIP та HNSW для ефективного пошуку схожих користувачів і фільмів у великих вибірках.

4. Оцінювання точності – проведення експериментів із порівнянням стандартного методу колаборативної фільтрації та запропонованого рішення з FAISS за метриками Precision@k, Recall@k і RMSE.

Результати експериментів показали, що впровадження FAISS дозволило скоротити час пошуку найближчих сусідів у 5–10 разів при збереженні або покращенні точності рекомендацій. Зокрема, для моделей на основі контентних ознак спостерігалось підвищення середнього значення Precision@10 з 0.68 до 0.74, що свідчить про більш релевантні рекомендації для користувачів. Також відзначено покращення узгодженості між поведінковими та контентними підсистемами завдяки спільному простору векторних подань.

Таким чином, розроблена система демонструє підвищену ефективність у формуванні персоналізованих рекомендацій. Застосування FAISS підтвердило свою доцільність у задачах масштабування та оптимізації швидкодії без втрати якості. Отримані результати можуть бути використані для побудови адаптивних рекомендаційних сервісів у сферах електронної комерції, кіноплатформ, онлайн-освіти та цифрового маркетингу.

Список використаних джерел

1. Harper, F. M., & Konstan, J. A. (2016). The MovieLens Datasets: History and Context. *ACM Transactions on Interactive Intelligent Systems (TiiS)*, 5(4), 1–19.
2. Johnson, J., Douze, M., & Jégou, H. (2019). Billion-scale similarity search with GPUs. *IEEE Transactions on Big Data*, 7(3), 535–547.
3. Ricci, F., Rokach, L., & Shapira, B. (2022). *Recommender Systems Handbook*. Springer, Cham.

АЛГОРИТМИ ОПЕРАТОРНОЇ ЕКСТРАПОЛЯЦІЇ ДЛЯ ПСЕВДОМОНОТОННИХ ВАРІАЦІЙНИХ НЕРІВНОСТЕЙ

Семенов В.В., semenov.volodya@gmail.com

Чернорай В.О., vlada2002vlada2001@gmail.com

Київський національний університет імені Тараса Шевченка

Варіаційні нерівності – зручна загальна форма запису різних задач, що виникають в математичній фізиці, дослідженні операцій та машинному навчанні. Розробка алгоритмів розв’язання варіаційних нерівностей та близьких задач (задачі про рівновагу, ігрові задачі) є надзвичайно популярним напрямом обчислювальної математики.

У повідомленні розглянуто варіаційні нерівності в гільбертовому просторі з псевдомонотонними операторами та два методи їх наближеного розв’язання – алгоритм операторної екстраполяції та його адаптивний варіант. Доведено теореми збіжності та оцінки швидкості збіжності алгоритмів. Отримані результати є новими та уточнюють відомі. Щоб отримати більш ефективні алгоритми для задач оптимізації, сідлових задач, варіаційних нерівностей та операторних включень, можна використати прийом змінної метрики або передобумовлення. Розроблено новий варіант алгоритму операторної екстраполяції, що використовує змінну метрику та отримані теореми про слабку та сильну збіжність для варіаційних нерівностей з ліпшицевими та монотонними операторами, що діють в гільбертовому просторі.

Також розглянуто декілька застосувань алгоритму операторної екстраполяції для задач сідлового типу, що зв’язані з наукою про дані та обробкою зашумлених сигналів

Робота виконана за підтримки НАН України (проект «Нові субградієнтні та екстраградієнтні методи для негладких задач регресії», 0124U002162).

СУЧАСНІ ПІДХОДИ ДЕТЕКТУВАННЯ ОБ'ЄКТІВ НА БОРТУ БПЛА**Сизоненко Р.М., syzonenko_r@365.dnu.edu.ua****Клименко С.В., klymenko_s@365.dnu.edu.ua***Дніпровський національний університет імені Олеся Гончара*

Детектування об'єктів – важлива складова безпілотних літальних апаратів (БПЛА) у різних сферах. Комп'ютерний зір на борту БПЛА є ключовим для забезпечення автономності, точності та безпеки в таких задачах, як моніторинг, пошуково-рятувальні операції, сільське господарство та військове застосування. Ця технологія дозволяє БПЛА в реальному часі розпізнавати об'єкти, уникати перешкоди і виконувати складні місії в динамічних умовах. Розвиток комп'ютерного зору сприяє підвищенню ефективності та розширенню можливостей БПЛА та зумовлений потребою в швидкій обробці даних і адаптації до обмежених обчислювальних ресурсів.

Моделі комп'ютерного зору на борту БПЛА потрібно адаптувати до обмежених ресурсів, енергоефективності та роботи у реальному часі. Сучасні підходи базуються на глибокому навчанні, зокрема на одноступеневих нейронних мережах, які поєднують швидкість і точність для одночасного виявлення й класифікації об'єктів.

Програмне забезпечення відіграє важливу роль у їх реалізації. Бібліотеки TensorFlow Lite та OpenVINO дають змогу запускати моделі на міні-комп'ютерах, таких як NVIDIA Jetson або Raspberry Pi. Фреймворки Ultralytics спрощують інтеграцію й налаштування під задачі БПЛА. Хмарні сервіси доповнюють обробку даних у складних сценаріях. Алгоритм YOLO для виявлення об'єктів у реальному часі розглядає задачу детектування об'єктів як регресію, прогножуючи рамки і ймовірності класів за один прохід нейронної мережі [1]. Оригінальна YOLOv1 має 24 згорткові та 2 повнозв'язні шари та опрацьовує зображення розміром 448×448 пікселів. Наступні версії вдосконалюють цю ідею. Спочатку

зображення ділять на сітку розміром, де кожна клітинка відповідає за виявлення об'єктів, центр яких потрапляє до неї. Модель прогнозує кілька обмежувальних рамок для кожної клітинки, визначаючи їхні координати (центр, ширина, висота) та ймовірність наявності об'єкта. Ймовірність належності до певного класу обчислюється на основі умовних ймовірностей для кожної клітинки, що дозволяє класифікувати об'єкти. Для оцінки точності прогнозів використовується показник перетину передбачених і справжніх рамок. Під час навчання модель оптимізує ваги, враховуючи помилки в прогнозуванні координат рамок, ймовірностей наявності об'єктів і класифікації. Вага помилок для координат вища, ніж для порожніх клітинок, щоб покращити локалізацію об'єктів. Даний алгоритм забезпечує високу швидкість обробки, що є вирішальним для БпЛА.

Ці підходи формують надійну основу для автономних БпЛА, забезпечуючи ефективну роботу в умовах обмежених ресурсів. Інтеграція з апаратними та хмарними рішеннями відкриває нові перспективи для безпілотних літальних апаратів. Сучасні підходи до детектування об'єктів на борту БпЛА забезпечують ефективне поєднання швидкості та точності, що критично для автономних систем та дозволяють адаптувати моделі комп'ютерного зору до обмежених ресурсів БпЛА, забезпечуючи їхню роботу в реальному часі.

Бібліографічні посилання

1. Redmon J., Divvala S., Girshick R., Farhadi A. You Only Look Once: Unified, Real-Time Object Detection. 2015. URL: <https://arxiv.org/abs/1506.02640> (дата звернення: 23.10.2025)

ІНТЕГРАЦІЯ ПЕРЕВІРОК БЕЗПЕКИ В CI/CD КОНВЕЄРИ НА БАЗІ GITHUB ACTIONS, DOCKER ТА TRIVY

Сироїд В.П., vitalii.syroid.mitis.2024@lpnu.ua

Басюк Т.М., taras.m.basyuk@lpnu.ua

Національний університет «Львівська політехніка»

Швидкий ритм змін у репозиторіях підвищує імовірність потрапляння відомих вразливостей та конфігураційних помилок у наступні середовища. Перенесення контролю безпеки на пізні етапи більше не забезпечує належного рівня захисту, тому перевірки доцільно вбудовувати безпосередньо у конвеєр безперервної інтеграції. Мета роботи полягає у практичній інтеграції автоматизованих перевірок так, щоб кожен коміт і кожен pull request супроводжувалися однозначним висновком щодо безпечності артефактів без уповільнення циклу збірки.

Запропоноване рішення реалізовано на базі GitHub Actions із використанням контейнеризації Docker і сканера Trivy. Подія push або відкриття pull request автоматично запускає workflow на ізольованому ранері. Сценарій виконує checkout, формує робочу копію, збирає контейнерний образ за Dockerfile і залежностями з файлу requirements.txt, після чого послідовно запускає два сканування Trivy. Перше охоплює файлову систему робочої копії репозиторію та виявляє вразливі залежності й типові помилки у вихідних файлах і конфігураціях. Друге аналізує зібраний контейнерний образ і фіксує вразливості у базовому шарі та встановлених у контейнері пакунках. Результати обох перевірок разом із журналами зберігаються як артефакти виконання, а підсумковий статус автоматично відображається у pull request. Образ не публікується у зовнішній реєстр і весь процес виконується в межах одного ранера, що спрощує контроль доступу та підвищує відтворюваність [1, 2, 3].

Отримані результати показали коректні спрацювання на рівні залежностей і контейнерного шару та прийнятний вплив на час збірки в умовах навчального стенду. Збережені артефакти дають змогу швидко переходити від інтегральної оцінки до конкретного файла чи пакета, який потребує виправлення, а також відтворювати умови запуску в разі повторної перевірки. Така організація процесу дозволяє робити висновок про безпечність артефакту ще на етапі безперервної інтеграції й не переносити прийняття рішень на пізнішу стадію.

Новизна підходу полягає в компактному і відтворюваному шаблоні CI на GitHub Actions із двома послідовними перевітками Trivy в одному проході без винесення проміжних артефактів за межі ранера і без публікації образів у реєстр. Дисципліна зберігання звітів і журналів формує повний слід виконання і спрощує аудит. Рішення не потребує ручних дій, не змінює звичний цикл розробки та переноситься на інші репозиторії зі схожою структурою без значних витрат на налаштування.

Подальший розвиток доречно спрямувати на розширення покриття безпеки без ускладнення процесу. Доцільно додати статичний і динамічний аналіз прикладного коду, перевірки інфраструктури як коду, формування відомості складу програмного забезпечення, підпис артефактів і перевірку походження під час випуску. Запропонований підхід є сумісним з іншими платформами CI/CD і зберігає логіку виконання під час перенесення.

Список використаної літератури:

1. GitHub Actions Documentation. GitHub: офіц. сайт. – [Електронний ресурс]. – Режим доступу: <https://docs.github.com/actions> (дата звернення: 18.10.2025).
2. Docker Documentation. Docker: офіц. сайт. – [Електронний ресурс]. – Режим доступу: <https://docs.docker.com/> (дата звернення: 18.10.2025).
3. Trivy Documentation. Aqua Security: офіц. сайт. – [Електронний ресурс]. – Режим доступу: <https://aquasecurity.github.io/trivy/> (дата звернення: 18.10.2025).

ОПТИМІЗАЦІЯ СПОЖИВАННЯ РЕСУРСІВ У ХМАРНИХ ІНФРАСТРУКТУРАХ ЗА ДОПОМОГОЮ ЕНЕРГОМОДЕЛЕЙ ПРОЦЕСОРІВ

Сімакін С.К., simakin_s@365.dnu.edu.ua

Божуха Л.М., bozhukha_l@365.dnu.edu.ua

Дніпровський національний університет імені Олеся Гончара

Сучасні вебсервіси дедалі частіше реалізуються у вигляді контейнеризованих або мікросервісних систем, розгорнутих у хмарних середовищах. Такі системи забезпечують масштабованість і гнучкість, але водночас спричиняють збільшення енергоспоживання серверних інфраструктур. За даними International Energy Agency (IEA), дата-центри та мережі передачі даних уже споживають понад 1,5% світової електроенергії[1].

Процесори сучасних архітектур (Intel, AMD, ARM) мають вбудовані інтерфейси апаратного моніторингу, зокрема Intel RAPL (Running Average Power Limit), AMD Zen Power та ARM Energy Probe, які дозволяють вимірювати фактичне споживання енергії процесорними ядрами, кешами та контролерами пам'яті. На основі цих даних запропоновано побудувати енергомодель контейнерного середовища, що описує взаємозв'язок між рівнем навантаження, частотою CPU та енергоспоживанням [2].

Сформовано прототип (MVP) енергоорієнтованої системи управління контейнерами. Система фокусується на навантаженні на CPU: дані надходять із Kepler Exporter (якщо наявний) або з апаратних сенсорів *hwmon*; за відсутності сенсорів використовується симуляція (для демонстрації роботи системи у будь-якому середовищі). Аналітика – проста лінійна регресія за CPU-ознаками (% CPU - завантаження процесора, частота, температура), що оцінює потужність і дає короткий прогноз. Оптимізатор динамічно коригує ліміти CPU: у Docker через *cgroups (nano_cpus/cpu_quota)*, у Kubernetes – патчуючи *resources.limits.cpu*

(за namespace/label). Підтримуються режими Observe/Advisory/Active, мертва зона (ϵ) та автоціль як частка від базового рівня, що зменшує коливання й зберігає стабільність. Інтеграція з Prometheus і Grafana забезпечує моніторинг та візуалізацію «до/після». Рішення придатне для хмарних середовищ (AWS, Azure, Kubernetes (K8s) для зниження споживання та вартості за vCPU й слугує базою для подальшого розширення (I/O-метрики, складніші моделі).

Для адаптивного управління ресурсами контейнера використовується простий пропорційний регулятор за відносною похибкою потужності. Поточна (або спрогнозована) потужність CPU порівнюється з цільовим рівнем, після чого ліміт CPU коригується пропорційно цій похибці. Щоб уникнути дрібних коливань, діє «мертва зона» (ϵ), а нове значення ліміту обмежується мінімальним і максимальним порогами. За потреби цільова потужність може визначатися автоматично як частка від базового (спостережного) рівня.

Бібліографічні посилання

1. Energy demand from AI. [Електронний ресурс]. – Режим доступу: <https://www.iea.org/reports/energy-and-ai/energy-demand-from-ai>
2. MDPI Computers. An Experimental Approach to Estimation of the Energy Consumption in Modern Systems Using RAPL. [Електронний ресурс]. – Режим доступу: <https://www.mdpi.com/2073-431X/12/7/139>

АРХІТЕКТУРА BIOMORPHNET ДЛЯ АВТОМАТИЗОВАНОГО ВИЯВЛЕННЯ ІШЕМІЧНОЇ ХВОРОБИ СЕРЦЯ ЗА РЕНТГЕНІВСЬКИМИ ЗНІМКАМИ

Соломатін В.А., mcsol2018@gmail.com, Байбуз О.Г.

Дніпровський національний університет імені Олеся Гончара

У роботі запропоновано нову архітектуру глибокої нейронної мережі BioMorphNet, призначену для автоматизованого виявлення ішемічної хвороби серця (ІХС) за рентгенівськими знімками. Модель поєднує класичні згорткові операції з адаптивними морфологічними перетвореннями, що дозволяє підсилити контраст судинної структури та дрібних анатомічних деталей. Особливістю підходу є інтеграція морфологічних ознак безпосередньо у навчальну архітектуру, що забезпечує її чутливість до структурних змін тканин.

Ішемічна хвороба серця залишається однією з головних причин передчасної смертності у світі. Сучасна діагностика базується на аналізі електрокардіограм, томографічних та рентгенівських досліджень. Проте візуальне виявлення ішемічних проявів на рентгенограмах є складним завданням через низький контраст та слабку вираженість патологічних змін.

Застосування методів глибокого навчання дає змогу автоматизувати процес аналізу медичних зображень, проте більшість існуючих моделей не враховують морфологічні особливості серцево-судинних структур. Це знижує інтерпретованість і точність на складних рентгенівських знімках.

Метою дослідження є розроблення нової архітектури нейронної мережі, яка поєднує біо інспіровані морфологічні операції та згорткові шари, що дозволяє підвищити ефективність розпізнавання ішемічних змін.

Модель BioMorphNet планується бути розробленою, як модульна глибока нейронна мережа, що складається з трьох основних блоків:

На початковому етапі обробки зображення проходить через набір диференційованих морфологічних перетворень — opening, closing, top-hat, black-hat — з різними розмірами структурних елементів.

Отримані карти ознак комбінуються у вигляді навчуваної лінійної суперпозиції, що дозволяє мережі автоматично виділяти ключові анатомічні деталі: контури серця, судини, області зміненої щільності. Використовується послідовність блоків Conv-BatchNorm-ReLU та механізмів уваги (Squeeze-and-Excitation або Channel Attention), що забезпечує вилучення високорівневих локальних ознак.

У порівнянні з класичними CNN, у BioMorphNet перший шар працює з морфологічно збагаченим вхідним простором, що підвищує стабільність навчання та якість розпізнавання на зображеннях низької контрастності.

Після глобального усереднення ознак (Global Average Pooling) застосовується двошаровий повнозв'язний класифікатор із функцією активації sigmoid, який визначає приналежність знімка до класу Ischemia або Healthy.

Архітектура передбачає можливість подальшого розширення — додавання сегментаційної голови (Heatmap Head) для локалізації підозрілих зон.

Архітектура має потенціал для створення інтелектуальної системи підтримки прийняття рішень у кардіологічній діагностиці, яка підвищить об'єктивність оцінювання та зменшить навантаження на лікарів-рентгенологів.

ЛП-ЗАДАЧА ДЛЯ ЗНАХОДЖЕННЯ РОЗРІДЖЕНОГО ВЕКТОРА

Стецюк П.І.^{1,2}, Стороженко А.О.², Хом'як О.М.¹

stetsyukp@gmail.com

¹Інститут кібернетики імені В.М. Глушкова НАН України²Київський академічний університет

Задача про знаходження розрідженого вектора [1] має вигляд

$$\min \sum_{j=1}^n |x_j| \quad \text{за обмежень} \quad \sum_{j=1}^n a_{ij}x_j = b_i, i = 1, \dots, m, \quad (1)$$

де $n \gg m$, а $|\circ|$ – модуль (абсолютна величина числа). Ця задача є апроксимацією задачі знаходження розв'язків з мінімальною кількістю ненульових компонент для системи лінійних рівнянь $Ax = b$.

Якщо ввести нові змінні $y_j = |x_j|$, то задачу (1) можна переформулювати як задачу лінійного програмування (ЛП-задачу) у такому вигляді

$$\min \sum_{j=1}^n y_j \quad \text{за обмежень} \quad \sum_{j=1}^n a_{ij}x_j = b_i, i = 1, \dots, m, \quad (2)$$

$$-y_j \leq x_j \leq y_j, y_j \geq 0, j = 1, \dots, n,$$

де кількість невідомих дорівнює $2n$, кількість обмежень дорівнює $m + 2n$, не включаючи обмежень на невід'ємність компонент вектора y .

В доповіді для різних діапазонів n та m наведемо результати обчислювальних експериментів щодо розв'язання ЛП-задач (2), які отримано за допомогою сучасного програмного забезпечення для розв'язання задач лінійного програмування. Це будуть солвери CPLEX, Gurobi, HiGHS та OOQP, сучасні версії яких доступні на NEOS-сервері [2]. Для опису ЛП-задач використана мова моделювання AMPL.

Список літератури

1. Candes E.J., Tao T. Decoding by linear programming. IEEE. Trans. Inf. Theory, 2005, vol. 51, № 12, pp. 4203–4215.
2. NEOS Solver. <https://neos-server.org/> (accessed: 31.10.2025).

ПРО ПРИСКОРЕНЕ РОЗВ'ЯЗАННЯ ЗАДАЧІ ОПТИМАЛЬНОЇ УПАКОВКИ НЕРІВНИХ КРУГІВ У КВАДРАТ

Стецюк П.І.^{1,2}, Тиводар С.Р.², Кузнєва А.С.²

stetsyukp@gmail.com

¹Інститут кібернетики імені В.М. Глушкова НАН України

²Ужгородський національний університет

Нехай задано сімейство m кругів C_i з радіусами r_i , $i=1, \dots, m$, які потрібно розмістити всередині квадрата K_0 з центром в точці $(0,0)$ та змінною довжиною сторони a . Будемо вважати, що всі радіуси кругів є різними. Оптимальною упаковкою сімейства кругів в квадрат K_0 називають таку їх упаковку, для якої a^* – довжина сторони квадрата K_0 є мінімальною за умов, що круги C_i , $i=1, \dots, m$ знаходяться всередині квадрата K_0 та попарно не перетинаються між собою (можуть торкатися один одного).

Якщо позначити (x_i, y_i) – невідомий центр круга C_i , $i=1, \dots, m$, то задачі оптимальної упаковки сімейства кругів в квадрат K_0 відповідає багатоекстремальна квадратична задача [1], яка має такий вигляд:

$$a^* = \min_{a, x, y} a \quad (1)$$

за обмежень

$$-\frac{a}{2} + r_i \leq x_i \leq \frac{a}{2} - r_i, \quad -\frac{a}{2} + r_i \leq y_i \leq \frac{a}{2} - r_i, \quad i=1, \dots, m, \quad (2)$$

$$(x_i - x_j)^2 + (y_i - y_j)^2 \geq (r_i + r_j)^2, \quad 1 \leq i < j \leq m, \quad (3)$$

де $x = (x_1, \dots, x_m)$, $y = (y_1, \dots, y_m)$.

Для знаходження розв'язку задачі (1) – (3) можна використовувати солвери з NEOS-сервера [2], які призначені для розв'язання задач глобальної оптимізації. Застосуємо відомий солвер BARON [3] для розв'язання п'яти тестових прикладів з веб-сайту

<https://www.packomania.com/>, де $m \in \{5, 6, 7, 8, 9\}$ та $r_i = i$ – радіуси кругів C_i , $i = 1, \dots, m$.

Прискорити розв’язання задачі (1) – (3) можна за рахунок додавання до неї запропонованих в [1] додаткових обмежень

$$x_m \leq 0, y_m \leq 0. \quad (4)$$

Додаткове прискорення можна досягти, якщо до задачі (1) – (4) додати таку лінійну нерівність

$$x_m \leq y_m. \quad (5)$$

Це підтверджують результати в наведеній нижче таблиці.

m	Задача (1)–(3)			Задача (1)–(4)			Задача (1)–(5)		
	N_{itm}	N_{nodes}	t	N_{itm}	N_{nodes}	t	N_{itm}	N_{nodes}	t
5	25	4	0.20	1	1	0.08	1	1	0.05
6	95	7	0.70	17	5	0.39	23	9	0.42
7	8871	78	52.83	953	13	3.31	658	31	2.00
8	66993	295	260.4	20847	208	80.6	22336	292	73.64
9	1358687	6432	7200	528837	1483	2919	276508	936	1358

Вони отримані за п’ять запусків солвера BARON та усереднені. Тут N_{itm} – усереднена кількість ітерацій, N_{nodes} – усереднена максимальна кількість задіяних вершин, t – усереднений час (в секундах). Задача (1) – (3) для $m=9$ не була розв’язана за вибраний максимальний час в 7200 секунд.

Список літератури

1. Тиводар С.Р., Стецюк П.І. Використання солвера BARON для знаходження оптимальної упаковки нерівних кругів у квадрат. Матеріали XXVII Міжнародного науково-практичного семінару «Комбінаторні конфігурації та їхні застосування» (Запоріжжя–Кропивницький–Київ, 4-6 червня 2025 року). Запоріжжя: НУ «Запорізька політехніка», 2025. С. 131–136.
2. NEOS Solver. <https://neos-server.org/> (accessed: 28.10.2025).
3. Sahinidis N.V., BARON 21.1.13: Global Optimization of Mixed-Integer Nonlinear Programs, User's manual, 2021.

r-АЛГОРИТМ ДЛЯ ЗАДАЧ ЕНТРОПІЙНО-ЛІНІЙНОГО ПРОГРАМУВАННЯ

Стовба В.О., vik.stovba@gmail.com

Жмуд О.О., zhmud17@gmail.com

Інститут кібернетики імені В.М. Глушкова НАН України

Задача ентропійно-лінійного програмування (ЕЛП) [1] в одній зі своїх найпоширеніших постановок має такий вигляд:

$$f^* = f(x^*) = \max_{x \geq 0} \left\{ f(x) = - \sum_{i=1}^n x_i \ln x_i + \sum_{i=1}^n c_i x_i \right\} \quad (1)$$

за обмежень

$$\sum_{i=1}^n a_{ji} x_i = b_j, \quad j = \overline{1, m}. \quad (2)$$

Вектор $c = \{c_i\}_{i=1}^n$, матриця $A = \{a_{ji}\}_{j,i=1}^{m,n}$ та вектор $b = \{b_j\}_{j=1}^m$ задані, а невід'ємні елементи вектора $x = \{x_i\}_{i=1}^n$ – невідомі. Передбачається, що задача (1) – (2) має єдиний додатний розв'язок $x^* = \{x_i^* > 0\}_{i=1}^n$.

Двоїста ЕЛП-задача до задачі (1) – (2) має такий вигляд:

$$\psi^* = \psi(u^*) = \min_{u \in \mathbb{R}^m} \left\{ \psi(u) = - \sum_{j=1}^m b_j u_j + \sum_{i=1}^n \exp \left(c_i + \sum_{j=1}^m a_{ji} u_j - 1 \right) \right\}, \quad (3)$$

де $u = \{u_j\}_{j=1}^m$ – невідомі множники Лагранжа, що відповідають обмеженням (2). Для ЕЛП-задач характерна умова $n \gg m$, тому для швидкого розв'язання задачі (1) – (2) необхідно ефективно розв'язати задачу (3), в якій кількість змінних набагато менша, ніж у прямої задачі. Для цього можна використати стандартні солвери для розв'язання задач нелінійного програмування з розділу «Nonlinearly Constrained Optimization» NEOS-сервера. Однак кращих результатів можна досягти, використовуючи двоїсті алгоритми, які враховують специфіку прямої та

двоїстої задач, причому ефективність таких алгоритмів залежить від обраного методу розв'язання задачі (3).

Нижче наведено результати розв'язання задач (1) – (2) та (3) при $n = 4000$ та $m = 400$ за допомогою солверів NEOS-сервера та задачі (3) за допомогою програми `ralgb5`, яка реалізує $r(\alpha)$ -алгоритм зі сталим коефіцієнтом розтягу простору та адаптивним регулюванням кроку [2]. Величини A , b та c генерувались з використанням генератора випадкових чисел, рівномірно розподілених на інтервалі $(0,1)$ [3].

№	Ipopt		KNITRO		MOSEK		ralgb5
	t_p	t_d	t_p	t_d	t_p	t_d	t_d
t_{best}	45.6	14.3	6.1	13.4	25.5	114.0	1.22
t_{avg}	47.5	15.2	10.2	16.9	35.7	132.4	1.24

В таблиці наведено результати роботи солверів, які продемонстрували найкращі результати з 9 наявних солверів, які підтримують мову AMPL. Для кожного солвера вказано найкращий та середній час розв'язання задачі з серії з 10 прямих та двоїстих задач, а для програми `ralgb5` – серії з 7 задач. Результати показують, що час на розв'язання задачі, витрачений $r(\alpha)$ -алгоритмом, є суттєво меншим, ніж час, витрачений солверами для розв'язання як прямої, так і двоїстої задач. Також з $r(\alpha)$ -алгоритмом були проведені додаткові тести з розв'язання задач зі значно більшими значеннями m та n , результати яких будуть представлені на доповіді.

Подяки. Робота підтримана грантом НАН України №02/01-2024(5) для дослідницьких груп/лабораторій молодих вчених.

Список літератури

1. Fang S.-C., Rajasekera J.R., Tsao H.-S.J. Entropy optimization and mathematical programming. Kluwer's International Series, 1997. 343 p.
2. Shor N.Z. Minimization Methods for Non-Differentiable Functions. Berlin: Springer-Verlag, 1985. 178 p.
3. Стецюк П.И., Гасников А.В. NLP-программы и r -алгоритм в задаче энтропийно-линейного программирования. Теория оптимальных решений. 2015. С. 73–78.

ПРОЄКТУВАННЯ МУЛЬТИМОДАЛЬНОЇ СТЕГANOГРАФІЧНОЇ СИСТЕМИ ДЛЯ ПРИХОВУВАННЯ ІНФОРМАЦІЇ В ГЕТЕРОГЕННИХ ЦИФРОВИХ КОНТЕЙНЕРАХ (ЗОБРАЖЕННЯ ТА АУДІО)

Стружко В.Р., Антоненко С.В.

Дніпровський національний університет імені Олеся Гончара

Традиційні стеганографічні системи, як правило, обмежені роботою з одним типом цифрового контейнера. Такий підхід має суттєві обмеження: він знижує загальну пропускну здатність і полегшує виявлення прихованих даних для стегоаналізатора, якому достатньо зосередитись на одному типі носія. Актуальною задачею є створення інформаційної технології, яка б розглядала різні типи файлів як єдиний ресурс для приховування даних.

Метою дослідження є архітектура мультимодальної стеганографічної системи, що може дозволити гнучко розподіляти секретне повідомлення між кількома гетерогенними контейнерами (цифровими зображеннями та аудіосигналами) з використанням оптимального набору алгоритмів.

Для досягнення мети були поставлені наступні задачі:

1. Проаналізувати сильні та слабкі сторони класичних методів стеганографії: LSB, DCT, Patchwork (для зображень) та Phase-mod, Spread-spectrum, Echo (для аудіо).
2. Спроекувати модульну архітектуру системи, що інтегрує ці алгоритми в єдину інформаційну технологію.
3. Розробити логіку модуля управління, який приймає рішення про стратегію вбудовування даних (розподіл або дублювання) на основі вимог користувача (пріоритет ємності, непомітності чи робастності). Наприклад, при необхідності більшої ємності, система вибере LSB для PNG-файлу, робастність до стиснення, система обере DCT для JPEG-файлу, а для робастності до шуму система вибере Spread-spectrum для WAV-файлу.

Запропонована система складається з чотирьох ключових модулів:

1. Модуль аналізу — приймає на вхід секретне повідомлення та набір потенційних контейнерів (наприклад, photo.jpg, logo.png, music.wav). Модуль оцінює характеристики контейнерів: розмір, формат, ступінь

стиснення, рівень шуму (для аудіо), наявність текстур (для зображень).

2. Бібліотека стеганографічних методів — містить реалізації досліджуваних алгоритмів. Кожен метод характеризується трьома параметрами: пропускна здатність, робастність (стійкість до атак типу стиснення, фільтрації) та непомітність (рівень внесених спотворень).
3. Модуль стратегічного планування (МСП) — ядро системи. На основі аналізу контейнерів, розміру повідомлення та обраного користувачем пріоритету, МСП обирає одну зі стратегій:
 - *розподіл*: повідомлення розбивається на частини. Кожна частина ховається в окремому контейнері з використанням найбільш підходящого методу (наприклад, 30% в photo.jpg методом DCT, 70% в music.wav методом Spread-spectrum). Це збільшує загальну ємність;
 - *дублювання*: повідомлення повністю або частково дублюється в різних контейнерах (наприклад, в logo.png методом LSB і в music.wav методом Echo). Це кардинально підвищує робастність — дані можна відновити, навіть якщо один з контейнерів буде пошкоджений або виявлений.
4. Модуль вбудовування/вилучення — технічно виконує операції приховування та вилучення даних відповідно до плану, створеного МСП.

Розроблена архітектура мультимодальної стеганографічної системи демонструє переваги комплексного підходу до приховування даних. Використання гетерогенних контейнерів та гнучкий вибір стратегії (розподіл або дублювання) дозволяє ефективно балансувати між трьома ключовими вимогами стеганографії: ємністю, непомітністю та робастністю. В той же час запропонована інформаційна технологія створює значно складніші умови для стегоаналізу порівняно з класичними одноконтейнерними системами.

НЕЙРОННІ МЕРЕЖІ У ЖАНРОВІЙ КЛАСИФІКАЦІЇ МУЗИЧНОГО КОНТЕНТУ

Сулейманов Є.С., suleimanov_e@365.dnu.edu.ua

Байбуз О.Г., baibuz_o@365.dnu.edu.ua

Зайцева Т.А., zaitseva_t@365.dnu.edu.ua

Дніпровський національний університет імені Олеся Гончара

Нейронні мережі набули широкої популярності при вирішенні найрізноманітніших класів задач, забезпечуючи високу результативність завдяки здатності автоматично виявляти приховані закономірності у вхідних даних. Проте, якість навчання суттєво залежить як від обраної архітектури мережі з відповідно налаштованими гіперпараметрами, так і від якості набору вхідних ознак та розміру тренувальної вибірки.

Музичні дані мають складний характер та значний обсяг, тому для ефективного навчання важливо ретельно виокремлювати ознаки, що характеризують різні аспекти композицій: гармонічні (співвідношення між акордами, тональностями та інтервалами), ритмічні (повторювані структури темпу, тривалості нот та акцентів), мелодичні (типові послідовності висот звуків або мотивів), темброві (відмінності у спектральних характеристиках інструментів або голосів), динамічні (зміни гучності, енергії та інтенсивності звуку у часі), структурні (розпізнавання куплетів та приспівів), стилістичні (характерні особливості різних жанрів або композиторських стилів) та частотні (співвідношення гармонік у спектрі звуку, які формують певний тембр).

Для жанрової класифікації застосовуються різні архітектури мереж: багатошаровий перцептрон (MLP), рекурентні (RNN) та згорткові (CNN), які при використанні мел-частотних кепстральних коефіцієнтів (MFCCs), що моделюють тембральні характеристики аудіосигналу, наближаючи розподіл частот до психоакустичного сприйняття людини, у якості вхідних даних демонструють точність у 60%, 67% та 74% відповідно [1].

На додачу до класичних архітектур, суттєвий прогрес у точності класифікації демонструють гібридні модифікації CNN+RNN завдяки використанню переваг згорткових мереж у знаходженні візуальних патернів та залучення властивостей рекурентних мереж для аналізу довготривалих залежностей [2].

Значний вплив на результативність навчання має вибір функції активації, що застосовується на кожному з шарів для додавання нелінійності. Незважаючи на широку популярність ReLU, кращі результати демонструє Leaky ReLU, дозволяючи зберігати від'ємні значення градієнта, масштабуючи їх із малим коефіцієнтом замість занулення [3].

Особливістю музичних даних є їх часова тривалість, що дозволяє збільшити обсяг вхідних даних шляхом поділу кожного музичного твору на відповідні часові фрагменти однакової довжини. При цьому особливу увагу слід приділяти тому, щоб усі елементи кожної композиції входили лише до однієї вибірки: тренувальної, валідаційної або тестової для запобігання витоку інформації. Поєднання результатів жанрової класифікації кожного з фрагментів для всього твору із залученням методів ансамблювання дозволяє суттєво збільшити точність передбачень [4].

Бібліографічні посилання

1. Jain S., Yadav S., Prabir P., Sundar S. Music information retrieval and classification using deep learning. International Research Journal of Engineering and Technology (IRJET), 06 June 2021, Vol. 08, P. 1059-1066.
2. Umale A., Mehul, Bhandw P., Bagwan S., Patil S.M. Music genre classification. International Journal of Advanced Research in Science, Communication and Technology (IJARSCT), 13 May 2023, Vol. 3, P. 414-425.
3. Hu Y., Mogos G. Music genres classification by deep learning. Indonesian Journal of Electrical Engineering and Computer Science, [S.l.] , February 2022, Vol. 25, No. 2, P. 1186-1198. DOI: <https://doi.org/10.11591/ijeecs.v25.i2.pp1186-1198>
4. Drotar P., Gazda M., Vokorokos L. Ensemble feature selection using election methods and ranker clustering. Information Sciences, 2019, Vol. 480, P. 365-380.

БАЗИ ДАНИХ СИСТЕМ ЕКОЛОГІЧНОГО МОНІТОРИНГУ ЩОДО СТАНУ ЕВТРОФУВАННЯ ВОД ПІВНІЧНО-ЗАХІДНОЇ ЧАСТИНИ ЧОРНОГО МОРЯ ТА ВИЗНАЧЕННЯ ТЕНДЕНЦІЙ ЙОГО ЗМІНИ

Тітяпкин А.С., tityapkin@ukr.net

*Український науковий центр екології моря
Інститут морської біології НАН України*

Евтрофікація є однією із головних проблем північно-західної частини Чорного моря (ПнЗЧМ), яка зазнає суттєвого впливу від стоку річок Дунай, Дністер, Дніпро, Південний Буг та інтенсивної господарської діяльності у прибережній зоні. Надлишкове надходження біогенних речовин (азоту, фосфору, кремнію) призводить до цвітіння води, дефіциту кисню, зниженню прозорості, погіршенню умов існування морських екосистем. Для ефективного контролю цих процесів необхідно використовувати узгоджені просторово-часові дані різних типів: супутникові, модельні та in-situ спостереження [1].

Основним джерелом супутникових і модельних даних сьогодні є Copernicus Marine Environment Monitoring Service (CMEMS), який забезпечує вільний доступ до багаторічних рядів спостережень за фізичними та біогеохімічними параметрами Чорного моря [2]. Так, дані з хлорофілу та біогенних речовин дозволяють оцінювати їх динаміку, визначати зони інтенсивного розвитку фітопланктону, оцінювати тенденції евтрофікації. Додатково, супутникові дані можна отримати з таких ресурсів, як NASA Ocean Color та Copernicus Data Space Ecosystem.

Важливою інформаційною складовою є вимірювання in-situ. Окрім національних та регіональних баз даних, наприклад, BS e-DataPlatform, Argo, система Держводагенства України, це дані, узагальнені в базах EMODnet та SeaDataNet. Ці системи містять результати багаторічних спостережень концентрацій біогенних речовин, розчиненого кисню, хлорофілу та інших параметрів, необхідних для валідації супутникових

продуктів, а також для аналізу вертикальних профілів і підтверджень наявності зон гіпоксії в придонних шарах ПнЗЧМ.

Для визначення зовнішнього навантаження важливі гідрологічні дані Global Runoff Data Centre (GRDC), що містять багаторічні ряди даних річкового стоку, а також дані стоку та антропогенного навантаження національних гідрометеорологічних служб, Дунайської комісії, тощо. Поєднання цих даних із супутниковими спостереженнями дозволяє кількісно оцінити вплив річкового притоку біогенних речовин на екологічний стан прибережних морських вод.

Аналіз багаторічних змін проводиться із застосуванням методів статистичної обробки часових рядів. Просторовий аналіз трендів за даними CMEMS виконується піксельно, що дозволяє виділяти зони зростання або зниження концентрації та значень гідролого-гідрохімічних та біологічних показників та оцінювати стабільність цих процесів у часі.

Застосування поєднаних супутникових, модельних і in-situ даних забезпечує комплексне розуміння динаміки евтрофікаційних процесів у північно-західній частині Чорного моря. Такий підхід сприяє формуванню науково-обґрунтованих рекомендацій щодо управління якістю морського середовища та моніторингу впливу антропогенних чинників у рамках міжнародних екологічних ініціатив – зокрема, Рамкової директиви ЄС щодо Морської Стратегії (MSFD) та програм Black Sea Commission.

Бібліографічні посилання:

1. Fedoronchuk N. et al. (2022). *Available Databases and FAIR Data/Metadata of Black Sea Environmental Parameters for Assessing the State of the Marine Environment and Predicting Geohazards*. DOI: [10.3997/2214-4609.2022580273](https://doi.org/10.3997/2214-4609.2022580273)
2. Зайченко М. Д., Тітяпкин А. С., Український В. В. (2022). *Огляд продуктів Служби моніторингу морського середовища Copernicus для оцінки стану та якості вод Чорного моря за показниками евтрофованості*. Євроінтеграція екологічної політики України: матеріали IV всеукр. наук.-практ. конф., Одеса: ОДЕКУ, с. 187–193.

**ЧИСЕЛЬНЕ МОДЕЛЮВАННЯ ДИНАМІЧНИХ ПРОЦЕСІВ
ПЕРЕНОСУ НА ОСНОВІ ОДНОВИМІРНИХ МОДЕЛЕЙ****Тіщук А.О., emorrior@gmail.com, Тонкошур І.С.***Дніпровський національний університет імені Олеся Гончара*

Дослідження процесів переносу має важливе значення в багатьох галузях науки й техніки. Ефективне чисельне моделювання цих процесів, особливо у нелінійних випадках, де можливе утворення розривів розв'язку, вимагає застосування стійких і точних чисельних схем.

У роботі розглядаються одновимірні моделі процесів переносу, в основі яких лежить скалярний закон збереження. Математична модель описується рівнянням у частинних похідних першого порядку:

$$\frac{\partial u}{\partial t} + \frac{\partial F}{\partial x} = 0,$$

де t – час, x – просторова координата, $u(x,t)$ – консервативна змінна, а $F(u)$ – функція потоку. Рівняння доповнюється відповідними початковими та крайовими умовами.

Для розв'язання цієї задачі застосовано скінченно-різницеві методи першого і другого порядків. Проведено порівняльний аналіз класичних різницевих схем, зокрема методу Лакса та методу МакКормака. Оскільки для схем другого порядку точності можливі нефізичні осциляції поблизу розривів, особливу увагу приділено методам їх пригнічення [1]. З цією метою використано TVD-обмежувачі, що дозволяє коректно відтворювати розв'язок задачі поблизу розривів або в областях з різкою зміною шуканих величин.

Програмну реалізацію алгоритмів та обчислювальні експерименти виконано у середовищі MATLAB. Отримано чисельні розв'язки наступних задач: про розповсюдження ударної хвилі, про транспортний потік на автомагістралі та про течію рідкої плівки по твердій поверхні.

1. Toro E.F. Riemann solvers and numerical methods for fluid dynamics: a practical introduction. Berlin: Springer-Verlag, 2009, 724 p.

ЧИСЕЛЬНИЙ АНАЛІЗ АЕРОДИНАМІЧНИХ ХАРАКТЕРИСТИК НЕКРУГОВИХ КОНУСІВ

Тонкошкур І.С., tonkoshkuris@i.ua

Дніпровський національний університет імені Олеся Гончара

Однією з важливих задач сучасної аеродинаміки є дослідження надзвукових течій навколо неосесиметричних несущих тіл, що рухаються в повітряному середовищі під деяким кутом атаки. Такі тіла можуть мати менший хвильовий опір, ніж тіла обертання той же довжини та об'єму. Але, внаслідок більшої площі поверхні, сили тертя, що діють на них, можуть істотно збільшити повний опір.

Розглядається задача про надзвукове обтікання гострого конуса з гладким поперечним перерізом потоком в'язкого газу. Припускається, що вплив в'язкості зосереджено в тонкому шарі поблизу твердої поверхні і всю область збуреної течії між ударною хвилею та поверхнею конуса можна розділити на нев'язкий потік і примежовий шар. Для розв'язання задачі в області нев'язкої течії застосовувався стаціонарний аналог метода Годунова, в примежовому шарі – скінченно-різницеви́й метод Петухова.

В результаті розв'язання вказаних вище крайових задач можуть бути знайдені поля швидкостей і тиску в примежовому шарі та в нев'язкій області течії, а також локальні коефіцієнти тиску і тертя. За локальними коефіцієнтами обчислюються інтегральні аеродинамічні характеристики: коефіцієнти лобового опору і підйимальної сили, а також аеродинамічна якість та коефіцієнт поздовжнього моменту.

По описаній методиці проведені розрахунки аеродинамічних характеристик конічних тіл з поперечним перерізом у вигляді круга, еліпса, а також трикутника або квадрата з округленими крайками. Проаналізовано вплив геометричних і фізичних параметрів на інтегральні характеристики течії. Показано, що форма поперечного перерізу суттєво впливає на величину аеродинамічних коефіцієнтів конуса.

ІНСТРУМЕНТАЛЬНІ ЗАСОБИ РОЗРОБКИ ЕКСПЕРТНИХ СИСТЕМ

Топольськов Є.О., y.topolskov@knu.ua,

Москаленко Н.В., moskalenkonv@ukr.net

Київський національний університет імені Тараса Шевченка

Експертна система є інтелектуальною системою, яка емулює дії людини-експерта в процесі пошуку та прийняття рішень. Експертні системи успішно вирішують класичні завдання інтелекту, що програмується. Вперше визначення експертної системи зробив професор Стенфордського університету Едвард Фейгенбаум. Він визначив експертну систему як інтелектуальну комп'ютерну програму, яка використовує знання та процедури виведення для вирішення проблем, що є досить складними і потребують використання людських знань. Основою експертної системи є база знань з набором правил та механізм виведення, що формує висновки на основі знань [1, 2].

Робота зі знаннями є основною задачею сучасних інтелектуальних технологій, а системи, ядром яких є база знань чи модель предметної галузі, вважаються інтелектуальними. Можна сказати, що експертна система є прикладною інтелектуальною системою, що функціонує в певній предметній області. Склад програмних компонентів і взаємодія між ними в експертних системах є доволі сталою, але задачі, які вирішують ці системи, потребують застосування відмінних від традиційних технологій програмування. В результаті узагальнення та аналізу розробок експертних систем була запропонована технологія створення цього виду інтелектуальних систем, що включає наступні етапи: ідентифікація, витягання знань, структуризація знань, формалізація, реалізація, тестування, дослідна експлуатація [2].

Розробка та створення сучасних експертних систем відбувається з використанням спеціальних інструментальних засобів, до яких відносяться [2, 3]:

1. Мови програмування, орієнтовані на створення експертних систем Lisp, CLIPS, Prolog та їх різновиди, Smalltalk, FRL, InterLisp, а також традиційні мови, такі як: C++, Java, Python тощо;

2. Середовища програмування, що автоматизують проєктування та розроблення експертних систем (Kee, BABYLON, ART, TEIRESIAS, MIKE, TIMM та інш.);

3. Оболонки або порожні фреймворки для створення експертних системи, такі як Jess (Java Expert System Shell), WindExS, RT-EXPERT та інші.

Оболонки можна розділити на дві групи в залежності від їхньої функціональності та спеціалізації:

- *Повні середовища розробки (Complete Environment)*. Ці оболонки є комплексними інструментами, які використовуються на всіх етапах життєвого циклу експертних систем: від безпосередньої розробки та тестування до використання і супроводу готової програми. Вони надають потужні засоби для моделювання знань та надійної роботи системи. Прикладами статичних експертних систем, створених за допомогою таких засобів є Nexpert Object, Leonardo, ProKappa, ART*Enterprise, Level 5 Object тощо.

- *Проблемно-орієнтовані засоби* – це спеціалізовані оболонки, сфокусовані на конкретних типах задач або предметних областях. Вони у свою чергу поділяються на дві підгрупи:

- *Орієнтовані на клас задач (Problem-Specific)*. Ці засоби мають у своєму складі набір альтернативних функціональних модулів, спеціально розроблених для вирішення певного класу проблем, що пов'язані з пошуком, управлінням, плануванням або прогнозуванням. Використання таких модулів значно спрощує реалізацію машини виводу для вузького кола задач.

- *Предметно-орієнтовані засоби (Domain-Specific)*. Ці оболонки містять вбудовані знання про певні типи предметних областей (наприклад, технічна діагностика чи медичний скринінг). Включення цих початкових знань дозволяє значно скоротити час розробки кінцевої бази знань.

Сучасний ринок інструментальних засобів дуже різноманітний і коректний їх підбір дозволяє значно скоротити як терміни і трудомісткість розробки експертної системи, так і підвищити їх якість.

Бібліографічні посилання

1. Булгакова О.С., Зосімов В.В., Поздєєв В.О. Методи та системи штучного інтелекту: теорія та практика: Навчальний посібник.- Херсон.:Олді плюс, 2020. 341 с.
2. Нікітіна Л.О. Експертні системи: навчальний посібник. – Харків: НТУ «ХПІ», 2023, 210 с.
3. Мирончук Д.Б. Програмні інструменти для автоматизації побудови експертних систем. – [Електронний ресурс]. URL: <https://www.inforum.in.ua/conferences/24/77/559>

ЗАСТОСУВАННЯ АЛГОРИТМУ ДЕЙКСТРИ ДЛЯ УЗАГАЛЬНЕНОЇ ЗАДАЧІ ПРО МІНІМАЛЬНІ ШЛЯХИ

Турчина В. А., Драниця М.В., nikitadranitsa@gmail.com.

Дніпровський національний університет імені Олеся Гончара

Алгоритм Дейкстри є класичним методом пошуку найкоротших шляхів у зважених графах, який забезпечує знаходження оптимального маршруту від заданої вершини до всіх інших [1]. Його особливістю є робота з одним типом ваги, що відображає певну міру відстані між вершинами. У стандартному варіанті ця вага найчастіше відповідає довжині шляху або часу у дорозі, і завдання полягає в мінімізації лише цього показника.

Проте у практичних задачах, зокрема у транспортних мережах, одного параметра недостатньо для адекватного моделювання ситуації. Рух між двома точками може характеризуватися не тільки довжиною шляху, а й додатковими факторами: середньою швидкістю руху, наявністю платних ділянок, витратами пального. Тому в практичних випадках звичайна постановка задачі пошуку найкоротшого шляху перетворюється на багатокритеріальну оптимізаційну задачу, де необхідно знайти маршрут, оптимальний за декількома критеріями.

У нашій задачі було розглянуто критерії пов'язані з пошуком економних та швидких маршрутів.

1. Економний маршрут відповідає мінімізації загальної вартості подорожі. В цьому випадку вага ребра визначається не лише довжиною, а й вартістю проїзду та витратами, обчисленими як функція відстані та вартості пального.
2. Швидкий маршрут має на меті мінімізацію загального часу в дорозі. Тут вага ребра обчислюється з урахуванням середньої швидкості на конкретній ділянці дороги.

Таким чином, узагальнення алгоритму Дейкстри у даному випадку полягає в тому, що вага ребра стала багатовимірною. Для кожного сценарію оптимізації ми будемо окрему функцію вартості: у першому випадку це грошові витрати, у другому — час. При цьому сам механізм роботи алгоритму Дейкстри залишається незмінним: на кожному кроці вибирається вершина з мінімальним значенням поточної функції вартості, і далі оновлюються оцінки для суміжних вершин.

Такий підхід дозволяє адаптувати класичний алгоритм до практичних задач, у яких необхідно враховувати кілька параметрів. Фактично відбувається перехід від одномірного до багатокритеріального аналізу, де користувач може обирати, що для нього важливіше — економія коштів чи мінімізація часу. Саме така гнучкість робить алгоритм Дейкстри придатним для використання в сучасних транспортних системах, логістичних мережах і навігаційних сервісах коли необхідно враховувати декілька критеріїв.

Окрім базових переваг, варто зазначити, що подібні алгоритми можуть бути основою для побудови інтелектуальних систем аналізу та прогнозування. Якщо у процес пошуку найкоротших шляхів інтегрувати дані з історичних спостережень або машинного навчання, можна отримати гібридну модель, здатну адаптуватися до змін у режимі реального часу. Наприклад, система може автоматично оновлювати ваги ребер, якщо змінюється швидкість руху на певній ділянці чи зростає вартість пального. Такий підхід дає змогу прогнозувати оптимальний маршрут не лише в поточний момент, а й з урахуванням майбутніх змін.

Перелік використаних джерел

1. Рудик О. «Пошук найкоротшого шляху у зважених графах». Інформатика та інформаційні технології в навчальних закладах, 2007, № 6, с. 104-107.

ПОРІВНЯННЯ ВНУТРІШНІХ ІНДЕКСІВ ЯКОСТІ КЛАСТЕРИЗАЦІЇ В ЗАДАЧІ ГРУПУВАННЯ ОЗНАК

Фунтиков М.К., nuntikov@gmail.com

Мацуга О.М., olga.matsuga@gmail.com

Дніпровський національний університет імені Олеся Гончара

Знаходження оптимальної кількості кластерів є актуальною задачею аналізу даних. Один з найбільш універсальних підходів до її розв'язання полягає у переборі розбиттів з різною кількістю кластерів і виборі оптимального на основі внутрішніх індексів якості кластеризації. Існує велика кількість таких індексів, проте питання вибору найефективнішого з них досі залишається відкритим. У цій роботі поставлено за мету експериментально порівняти внутрішні індекси якості для вибору оптимальної кількості кластерів у задачі групування ознак набору даних.

Для порівняння обрані 13 індексів [1]:

- індекси, для яких оптимальним є мінімальне значення: Davies-Bouldin, C-Index, Gplus, Xie-Beni;
- індекси, для яких оптимальним є максимальне значення: Silhouette, Ratkowsky-Lance, Calinski-Harabasz, Dunn, Point-biserial, Cubic Clustering Criterion, PBM, SF1, SF2.

Для кожного індексу додатково були реалізовані дві модифікації, що базуються на лівій (ldiff) та правій (rdiff) скінченних різницях першого порядку. Нехай задано функцію $f(k)$, яка описує залежність значення індексу від кількості кластерів k . Тоді модифікації ldiff та rdiff визначаються відповідно:

$$\nabla f(k) = f(k) - f(k - 1),$$

$$\Delta f(k) = f(k + 1) - f(k).$$

Експерименти були проведені на 20 тестових наборах даних, що містили групи ознак із вираженою лінійною кореляцією між ними. Для групування ознак застосовувалася агломеративна ієрархічна кластеризація

з мірою зв'язку average, яка виконувалася на матриці парних коефіцієнтів кореляції Пірсона. Кластеризація проводилася за різної кількості кластерів, оптимальна кількість визначилася за обраними індексами і порівнювалася з відомою правильною кількістю.

Для експериментів були використані Python-бібліотеки scikit-learn, pandas, scipy та ClusterValid [2]. В останній реалізовані усі зазначені індекси кластеризації.

Результати експериментів показали:

- найточнішими серед базових індексів (без урахування модифікацій) були: Dunn (16 правильних відповідей з 20), Point-biserial (14/20), Silhouette (13/20), Gplus та Cubic Clustering Criterion (10/20);

- найменш точними виявилися такі базові індекси: PBM (0/20), Davies-Bouldin та SF1 (2/20), C-Index (5/20);

- найточнішими серед модифікацій були: Dunn ldiff та rdiff (17/20), Calinski-Harabasz ldiff, Cubic Clustering Criterion ldiff, Ratkowsky-Lance rdiff (13/20), Xie-Beni rdiff (12/20);

- модифікування покращило точність таких індексів: Davies-Bouldin ldiff (2/20 → 8/20), Ratkowsky-Lance rdiff (7/20 → 13/20), SF1 ldiff (2/20 → 7/20);

- загалом, найкращі результати продемонстрували модифікації індексу Dunn на основі скінченних різниць (ldiff та rdiff), які забезпечили 17 правильних визначень із 20.

Список використаної літератури

1. Todeschini R., Ballabio D., Termopoli V., Consonni V. Extended multivariate comparison of 68 cluster validity indices. A review. Chemometrics and Intelligent Laboratory Systems. 2024. Vol. 251. 105117. <https://doi.org/10.1016/j.chemolab.2024.105117>.

2. ClusterValid – New library for clustering evaluation in Python URL: <https://github.com/VladGTT/ClusterValid> (дата звернення: 30.10.2025)

**МАТЕМАТИЧНЕ МОДЕЛЮВАННЯ ТА ОБЧИСЛЕННЯ
ПРОСТОРОВОЇ ДИНАМІКИ ТРАЄКТОРІЇ ПОЛЬОТУ
АРТИЛЕРІЙСЬКОГО СНАРЯДУ**

Циба В. В., Пасічник А. М., aksel1455@gmail.com

Дніпровський державний технічний університет

В сучасній балістиці моделювання та обчислення траєкторій польоту артилерійських снарядів із врахуванням їх обертального руху навколо своєї осі є достатньо складною математичною задачею через необхідність врахування нелінійних сил та моментів. Для стабілізації польоту снаряду задають початкову швидкість обертання, що породжує складні динамічні ефекти такі як гіроскопічні ефекти, сили тертя і сила Магнуса та зумовлює необхідність застосування шести ступеневої моделі вільності, яка описує повну просторову динаміку і відома як модель 6-DOF.

Для точного розрахунку похідних аеродинамічних сил та моментів необхідно проводити експериментальні дослідження, такі як випробування в аеродинамічній трубі або іскровий метод, які несуть в собі значні обмеження щодо врахування всіх умов та відповідно пов'язаних витрат. За останні кілька десятиліть було зроблено величезний прорив у застосуванні обчислювальної гідрогазодинаміки CFD для обчислення аеродинамічних характеристик літальних апаратів, включаючи і динаміку польоту снаряда. У комбінації з методом «rigid body motion» алгоритм застосування CFD/RBD дозволяє здійснювати «віртуальний політ» снаряду і інтегровано прогнозувати всі пов'язані з ним аеродинамічні характеристики вільного польоту включаючи похідні 6 DOF-моделі. Крім того, CFD/RBD дозволяє розраховувати такі динамічні похідні, як демпфування тангажу і відхилення, а також коефіцієнти сил і моментів Магнуса, які складно отримати експериментально в аеродинамічній трубі.

Застосування налаштованої CFD-модель дозволяє проводити комплексну оптимізацію аеродинамічної форми снарядів в широкому

діапазоні зміни параметрів конструкції та початкових умов і дальності їх польоту. CFD моделювання дозволяє більш точно обчислювати вплив зміни щільності, температури, тиску та швидкості звуку, що є суттєвими факторами визначення динаміки траєкторії польоту снаряда при стрільбі з далекобійної артилерії.

Активний розвиток паралельних обчислень за допомогою GPU-прискорювачів дозволив значно збільшити швидкість обчислень при використанні CFD моделювання на кілька порядків, що дозволяє використовувати більш точні моделі. При цьому необхідно враховувати, що застосування методу великих вихорів (LES) також дозволяє покращити точність обчислення коефіцієнтів сили та моменту Магнуса, як показано в роботі [1] для обчислення траєкторії польоту 203-мм снаряда.

За результатами проведених досліджень алгоритм застосування підходу CFD/RBD показав себе достатньо ефективним методом створення повного аеродинамічного опису динаміки траєкторії польоту снаряда. Постійно зростаюча комп'ютерна потужність, доступна для обчислень, дозволяє застосовувати більш досконалі моделі та методи, такі як гібридний метод DES і підхід великого вихрового моделювання (LES), що забезпечує більш точний розрахунок складних нестационарних аеродинамічних характеристик.

1. Chughtai FA, Masud J, Akhtar S. Unsteady aerodynamics computation and investigation of magnus effect on computed trajectory of spinning projectile from subsonic to supersonic speeds. *The Aeronautical Journal*. 2019. № 123(1264). p.863-889. doi:10.1017/aer.2019.32.

МОДЕЛІ І АЛГОРИТМИ ОПТИМІЗАЦІЇ РУХУ АВТОНОМНИХ РОБОТІВ У ЗАМКНЕНОМУ ПРОСТОРІ

Чаша М. Ю., cchashka20@gmail.com.

Дзюба П. А., avatarr@2282gmail.com.

Дніпровський національний університет імені Олеся Гончара

Робота була присвячена темі моделювання та оптимізації руху автономних роботів у замкненому просторі. Автономні роботи застосовуються для автоматизації складів, транспортування матеріалів у лікарнях і навіть у космічних місіях, де вони мають ефективно планувати траєкторії, уникаючи перешкод. Задача пошуку оптимального шляху в обмеженому просторі з перешкодами є однією з центральних у робототехніці, адже від її розв'язання залежить швидкість, безпека та економія ресурсів, таких як час і енергія. Основна мета полягала у створенні математичної моделі, алгоритмічних методів та програмної реалізації системи, яка дозволяє ефективно планувати траєкторію руху автономного робота в умовах наявності перешкод [1].

Було проведено аналітичний огляд сучасних підходів до автономної навігації, включаючи графові методи, машинне навчання та системи локалізації. Особливу увагу приділено алгоритму A*, який поєднує точність і швидкість, а також його перевагам над алгоритмами Дейкстри та D* Lite [2, 3]. Розглянуто можливість інтеграції методів машинного навчання для підвищення адаптивності навігаційних систем.

Створено математичну модель дискретного простору, де кожна клітинка може бути вільною або зайнятою перешкодою. Розроблений алгоритм A* дозволяє знаходити найкоротший шлях від старту до цілі, обходячи перешкоди [4]. Модель включає механізм регенерації простору, який гарантує існування шляху навіть у складних умовах.

Програмна реалізація виконана мовою Python із використанням бібліотек NumPy і Matplotlib [5]. Система забезпечує генерацію простору,

анімовану візуалізацію руху робота та логування процесу пошуку шляху. Передбачено можливість автоматичної регенерації середовища у випадку відсутності доступного маршруту.

У ході експериментів було проведено серію тестів із різними параметрами середовища. Результати показали, що алгоритм A* забезпечує високу швидкість і точність пошуку шляху. Ефективність оцінювалася за довжиною маршруту, коефіцієнтом ефективності та часом виконання. Отримані результати підтвердили правильність математичної моделі та адекватність програмної реалізації.

У цій роботі розроблено повноцінну систему для моделювання руху автономного робота у замкненому просторі. Досягнуто всі поставлені завдання: створено математичну модель, реалізовано алгоритм пошуку шляху, проведено тестування.

Бібліографічні посилання

1. Corke P. Robotics, Vision and Control: Fundamental Algorithms in Python. – 3rd ed. – Springer, 2023. – 822 p.
2. LaValle S.M. Planning Algorithms. – Cambridge: Cambridge University Press, 2006. – 846 p.
3. Thrun S., Burgard W., Fox D. Probabilistic Robotics. – MIT Press, 2005. – 672 p.
4. Koenig S., Likhachev M. D* Lite. – Proceedings of the AAAI Conference on Artificial Intelligence, 2002. – pp. 476–483.
5. Matplotlib Development Team. Matplotlib: Visualization with Python [Електронний ресурс]. – Режим доступу: <https://matplotlib.org/stable/index.html>.

МЕТОДИ МАШИННОГО НАВЧАННЯ ДЛЯ АДАПТИВНОЇ ПЕРСОНАЛІЗАЦІЇ ВЕБ ІНТЕРФЕЙСІВ У РЕАЛЬНОМУ ЧАСІ

Червоняк А.Б., Полягушко Л.Г., Сліпченко В.Г.,

chervoniak.andrii@lil.kpi.ua

*Національний технічний університет України «Київський політехнічний
інститут імені Ігоря Сікорського» (<https://kpi.ua>)*

Сучасні веб-системи стикаються з викликом персоналізації досвіду користувачів у реальному часі. Традиційні підходи на основі статичних правил недостатні для динамічної адаптації інтерфейсів. Персоналізовані веб-інтерфейси на основі машинного навчання підвищують конверсію на 15-40% та зменшують відмови на 25-30%.

Актуальною задачею є вибір оптимального методу машинного навчання для адаптивної персоналізації в реальному часі з урахуванням точності та швидкодії [1]. Проблема включає дисбаланс класів у даних про поведінку. Метою роботи є узагальнення методів машинного навчання та оцінка їхньої ефективності [2].

Експеримент проводився на датасеті "Online Shoppers Purchasing Intention" (12 330 сесій) [3]. Датасет розділено на тренувальну (80%) та тестову (20%) вибірки. Для дисбалансу застосовувався SMOTE. Порівняно шість алгоритмів: MLP (10 нейронів), SVM (лінійне та RBF), Random Forest (100 дерев), Naive Bayes, KNN (k=5). Оцінювання за Accurasy, F1, AUC, часом інференсу та пропускну здатністю.

MLP показав найкращий баланс: 86.70% accurasy, 63.15% F1, 89.68% AUC, інференс <0.1 мс, ~34 500 прогнозів/с. Naive Bayes також швидкий (~29 500 прогнозів/с), але з нижчою точністю. SVM та Random Forest конкурентні, але повільніші.

Результати порівняння ефективності методів представлені на таблиці 1.

Таблиця 1. Порівняння ефективності методів

Метод	Точність	F1-міра	AUC	Інференс (мс)	Пропускна здатність (прог./с)
MLP	86.70%	63.15%	89.68%	<0.1	~34 500
SVM Linear	87.39%	61.27%	87.11%	0.1	~8 200
SVM RBF	86.90%	61.50%	88.09%	0.3	~3 400
Random Forest	82.85%	56.88%	86.96%	1.3	~760
Naive Bayes	81.47%	52.94%	80.86%	<0.1	~29 500
KNN	81.27%	54.97%	83.98%	0.3	~3 500

MLP є оптимальним для адаптивної персоналізації завдяки балансу точності та швидкодії. Підхід застосовний у веб-системах електронної комерції. Перспективи: гібридні моделі, глибоке навчання (RNN, Transformers) та оптимізація для edge-пристроїв.

ДЖЕРЕЛА

1. Zhang S., Yao L., Sun A., Tay Y. Deep learning based recommender system: A survey and new perspectives. ACM computing surveys (CSUR). 2019. Vol. 52, no. 1. P. 1–38. DOI: <https://doi.org/10.1145/3285029>
2. Naumov M. et al. Deep learning recommendation model for personalization and recommendation systems. arXiv preprint arXiv:1906.00091. 2019. URL: <https://arxiv.org/abs/1906.00091>
3. Online Shoppers Purchasing Intention Dataset / C. Sakar et al. UCI Machine Learning Repository. 2018. URL: <https://archive.ics.uci.edu/ml/datasets/Online+Shoppers+Purchasing+Intention+Dataset>

КОМПОЗИТНЕ ГЕОМЕТРИЧНЕ МОДЕЛЮВАННЯ ДЛЯ БРАХІТЕРАПІЇ З АНІЗОТРОПНИМ ВИПРОМІНЮВАННЯМ: ОПТИМІЗАЦІЙНИЙ ПІДХІД ДО РОЗМІЩЕННЯ КАПСУЛ

Чугай А.М.^{1,2}, Яськов Г.М.^{1,3}, Старкова О.В.², Щербина М.О.¹
chugay.andrey80@gmail.com, yaskov@ukr.net, olha.starkova@hneu.net,
maxshcherbyna247@gmail.com

¹*Інститут енергетичних машин і систем ім. А.М. Підгорного НАНУ*

²*Харківський національний економічний університет імені Семена Кузнеця*

³*Харківський національний університет радіоелектроніки*

Брахітерапія — це метод внутрішньої променевої терапії, при якому джерела випромінювання розміщуються безпосередньо в пухлині або поруч із нею. Такий підхід забезпечує високу точність покриття випромінювання та мінімальний вплив на здорові тканини. Особливий інтерес становить брахітерапія з анізотропним випромінюванням, коли джерело розташоване на одному торці капсули, а випромінювання спрямоване переважно в один бік [1]. Це дозволяє захистити критичні органи та зменшити токсичність лікування.

Запропоновано математичну модель для планування лікування за допомогою брахітерапії. Вона ґрунтується на класичній задачі пакування однакових об'єктів у обмежену область (Identical Item Packing Problem, ІІРР [3]). У нашому випадку об'єктами є циліндричні капсули, які необхідно розмістити в межах цільової області (опуклий поліедр), забезпечуючи відсутність перетинів між капсулами, дотримання мінімальних допустимих відстаней, контроль орієнтації для уникнення паралельного розміщення капсул.

Для аналітичного опису розміщення капсул використано метод нормалізованих Ф-функцій, який дає змогу формалізувати геометричні обмеження.

Запропоновано композитну модель капсули, яка складається з циліндра (фізичний імплант) та кулі, розташованої на торці циліндра (зона впливу джерела випромінювання), що відображає анізотропність випромінювання.

Розроблено гібридну стратегію оптимізації, яка поєднує дві стратегії пакування: онлайн та офлайн.

Пакування онлайн передбачає послідовне додавання капсул у зменшеному масштабі для швидкої генерації початкового розміщення. Пакування офлайн коригує положення та орієнтацію капсул із поступовим відновленням розмірів шляхом розв'язання задачі нелінійного програмування. Такий підхід забезпечує адаптивність до складних анатомічних умов і баланс між обчислювальною ефективністю та клінічною точністю.

Для оцінки ефективності алгоритму проведено серію обчислювальних експериментів із різною кількістю капсул, геометричними параметрами та обмеженнями. Результати показали, що метод здатний адаптивно формувати допустимі конфігурації навіть за складних просторових умов, забезпечуючи контроль орієнтації та дозової відповідності. Алгоритм успішно працює як у випадках без кутових обмежень, так і при введенні додаткових геометричних та анатомічних вимог, що підтверджує його гнучкість і придатність для клінічного застосування.

Запропонована модель та гібридний алгоритм оптимізації забезпечують точне розміщення капсул у складних анатомічних областях, враховують анізотропність випромінювання та геометричні обмеження, а також демонструють масштабованість і можливість інтеграції з інтелектуальними системами планування.

Подальші дослідження передбачають використання більш точних моделей анізотропного випромінювання, розробку удосконалених Ф-функцій для циліндрів та інтеграцію з роботизованими системами багатокритеріального планування.

Aima M., DeWerd L. A., Mitch M. G., Hammer G. G., Culberson W. S. Dosimetric characterization of a new directional low-dose rate brachytherapy source. *Medical Physics*. 2018. Vol. 45, no. 6. doi:10.1002/mp.12994.

Wäscher G., Haußner H., Schumann H. An improved typology of cutting and packing problems. *European Journal of Operational Research*. 2007. Vol. 183, no. 3. Pp. 1109–1130. doi:10.1016/j.ejor.2005.12.047.

ЗМІСТ

1.	Akhmetshina L.G., Yegorov A.A., Nikishyna O.Y. AN IMPROVED METHOD OF GRAYSCALE IMAGES CONTRAST ENHANCEMENT BASED ON POWER TRANSFORMATIONS	3
2.	Bidyuk P.I., Tertychnyi R.V., Chuprin D.S. FORECASTING NONLINEAR NON-STATIONARY PROCESSES IN FINANCE	5
3.	Dzhenkova M., Sheveleva A. ESTIMATING TASKS IN SCRUM PROJECTS BASED ON FUZZY LOGIC	7
4.	Forkert P. P., Ivanchenko M. G. IMPLEMENTING DEFAULT PARAMETER VALUES IN GO PROGRAMMING LANGUAGE DIALECT	8
5.	Huk K., Sheveleva A. INTELLIGENT CONTROL OF AIR QUALITY BASED ON MASS BALANCE MODEL AND EXPERT RULES	10
6.	Huk N., Yehoshkin D. ADAPTIVE APPROACH TO THE DEVELOPMENT OF EXPERT SYSTEMS BASED ON WEIGHTED RULES AND PATTERN ANALYSIS OF SYSTEM STATES	12
7.	Ivanchenko M., Bondarenko B. MIXED REALITY WORKSPACES: REDEFINING REMOTE/HYBRID PRODUCTIVITY VIA XR HEADSETS	14
8.	Kiseleva O., Prytomanova O., Zhurava D. SOFTWARE IMPLEMENTATION FOR VISUALIZATION OF OPTIMAL PARTITIONING SET RESULTS ON MAPS	16
9.	Kiselyova O.M., Stroieva V.O., Tarasiuk O.S. OPTIMIZATION OF TECHNOLOGICAL DESIGN TASKS OF COMPLEX SYSTEMS	17
10.	Krasnoshapka D.V., Zolotko K.Y. WIRELESS COMPUTER NETWORKS SIMULATION	19
11.	Letuta N., Dubinskyi, V., Romanova T., Gutierrez Luis GEOMETRIC DESIGN OF NANOCOMPLEXES: MATHEMATICAL AND COMPUTER MODELING	21
12.	Levchenko Y., Khyzha O. PROVING SORTING ALGORITHMS IN THE DAFNY VERIFICATION SYSTEM	22
13.	Lushnya K.K., Yehoshkin D.I. OPTIMIZATION OF LOGISTICS ROUTES USING GRAPH THEORY ALGORITHMS	24
14.	Lutsenko O.M., Trofimov O.V. BUILDING A MATHEMATICAL MODEL OF AGENT INTERACTION IN COLLECTIVE DECISION-MAKING TASKS	27
15.	Nikolenko K., Moroz V. EXPERIMENTAL EVALUATION OF SEMANTIC SEGMENTATION WITH MULTI-SCALE ADAPTIVE ATTENTION ON THE CAMVID DATASET	29
16.	Novikova O.V., Suima I.P. INTEGRATING ARTIFICIAL INTELLIGENCE INTO ENGLISH LANGUAGE TEACHING	31
17.	Pankratova N.D., Musiienko D.I. COMMUNITY ANALYSIS AND MODULARITY ASSESSMENT IN COGNITIVE MODELLING OF COMPLEX SYSTEMS	33

18.	Pankratova N.D., Pankratov V.A., Golinko I.M. DIGITAL TWINS IN DECENTRALIZED UAV SWARM CONTROL	35
19.	Poslaiko N.I. METHOD OF CALCULATING PROBABILITY OF STATES OF A SINGLE MASS SERVICE SYSTEM WITH UNRELIABLE SERVICE	37
20.	Radiuk P.M., Barmak O.V., Krak I.V. EXPLAINABLE AI VIA TRANSITION-MATRIX MAPPING FROM DEEP TO INTERPRETABLE FEATURES	39
21.	Rakova K.O., Sheveleva A.E. APPLICATION OF BAYESIAN NETWORKS IN RECOMMENDER SYSTEMS	42
22.	Shcherbak R., Sheveleva A. MODELING DISCRETE PRE-FRACTURE GAPS AND ELECTROMECHANICAL RESPONSE IN DISSIMILAR PIEZOELECTRIC BIMATERIAL SYSTEMS	44
23.	Shekhovtsov S., Romanova T. OPTIMIZED PACKING OF IRREGULAR OBJECTS	46
24.	Stroieva V.O., Zhyhaleva S.P. MATHEMATICAL MODELING OF ADAPTIVE LEARNING SYSTEMS	47
25.	Suima I.P., Novikova O.V. AI AS A TEACHING ASSISTANT: REDEFINING THE ROLE OF THE ENGLISH LANGUAGE TEACHER	48
26.	Sukovenko K. DECISION SUPPORT INFORMATION TECHNOLOGY IN VIDEO SURVEILLANCE AND MONITORING TASKS	50
27.	Sushentsev N., Zaikin F., Barrett T., Blyuss O. AN EXPLAINABLE ARTIFICIAL INTELLIGENCE FRAMEWORK FOR PREDICTING PROSTATE CANCER PROGRESSION IN MEN UNDER ACTIVE SURVEILLANCE	52
28.	Symonov D., Symonov Y., Zaika B., Kutova I., Fesenko M. INTELLIGENT MULTI-AGENT SYSTEMS FOR ADAPTIVE MANAGEMENT UNDER UNCERTAINTY	54
29.	Trotsenko A.G., Kuzenkov O.O. MATHEMATICAL MODELLING OF COMPLEX NATURAL SYSTEMS: A PROBABILISTIC FIRE SPREAD MODEL WITH ENVIRONMENTAL FACTORS	56
30.	Tsukanova A. O. LINEAR ALGEBRA, COMPUTER METHODS, AND OPTIMIZATION APPARATUS	58
31.	Vovk S.M. MATHEMATICAL MODELS OF DATA PROCESSING BASED ON THE FUNCTIONAL OF QUASI-EXTENT	60
32.	Yakobson O.M., Nakonechna T.V. IN RESEARCH OF SIMPLEST CONDITIONS FOR CLUSTERING OF A SET OF AGENTS WITH MEMORY	62
33.	Yuriev M.V., Matsuga O.M. PERFORMANCE EVALUATION OF ALGORITHMS FOR FITTING SEGMENTED LINEAR REGRESSION WITH AUTOMATIC BREAKPOINT DETERMINATION	64
34.	Zora I., Lishchynska L. INTELLIGENT SYSTEMS AS A TOOL FOR SOLVING THE PROBLEM OF SUBJECTIVITY IN INPUT DATA	66

35.	Анісімов О.А., Білозьоров В.Є. МАТЕМАТИЧНЕ ТА ПРОГРАМНЕ ЗАБЕЗПЕЧЕННЯ СИСТЕМИ СЕГМЕНТАЦІЇ ДІЛЯНОК ГОЛОВНОГО МОЗКУ ТА АНАЛІЗУ ЇХ АТРОФІЇ ДЛЯ ВИЯВЛЕННЯ ОЗНАК ХВОРОБИ АЛЬЦГЕЙМЕРА	69
36.	Антонов В.С., Турчина В. А. ПРОГНОЗУВАННЯ ЧАСУ ВИКОНАННЯ ФРАГМЕНТІВ КОДУ	71
37.	Афанасьєв Д.К., Зайцев В.Г. СИНТЕЗ СИСТЕМ КЕРУВАННЯ ПРОСТОРОВИМ РУХОМ ЛІТАЛЬНОГО АПАРАТА	73
38.	Балейко А.С., Михальчук Г.Й. УЗАГАЛЬНЕНА ЗАДАЧА ПРО ПРИЗНАЧЕННЯ ЯК ІНСТРУМЕНТ ОПТИМІЗАЦІЇ ВИРОБНИЧИХ СИСТЕМ	75
39.	Бобко О.Л., Ліщинська Л.Б. ВИКОРИСТАННЯ ІНФОРМАЦІЙНИХ СИСТЕМ У МЕТОДАХ РОЗПАРАЛЕЛЕННЯ ПРОЦЕСУ ФОРМУВАННЯ ТРИВИМІРНИХ ГРАФІЧНИХ ЗОБРАЖЕНЬ	77
40.	Бовкун М.Ф., Мацуга О.М. АНАЛІЗ ЕФЕКТИВНОСТІ АЛГОРИТМІВ ГЛИБОКОГО НАВЧАННЯ З ПІДКРІПЛЕННЯМ У 2D-ІГРОВОМУ СЕРЕДОВИЩІ	79
41.	Богданович Б.М., Козакова Н.Л. ЗАСТОСУВАННЯ ГЕНЕТИЧНОГО АЛГОРИТМУ ДЛЯ РОЗВ'ЯЗАННЯ ЗАДАЧІ КОМІВОЯЖЕРА З УРАХУВАННЯМ ЧАСУ ТА ВІДСТАНІ	81
42.	Богомаз В.М., Богомаз О.В. ПРО КОМПЛЕКСИ НАВЧАЛЬНИХ ТРЕНАЖЕРІВ ІНЖЕНЕРНОЇ ТЕХНІКИ ДЛЯ ПІДГОТОВКИ ВІЙСЬКОВИХ ФАХІВЦІВ	83
43.	Бодю К., Вовк С. ІНТЕРАКТИВНИЙ ВЕБ-РЕСУРС З ВИВЧЕННЯ ІНФОРМАТИКИ В КОНТЕКСТІ НОВОЇ УКРАЇНСЬКОЇ ШКОЛИ	85
44.	Божуха Д.І., Байбуз О.Г. ПРО ОПТИМІЗАЦІЮ НАВАНТАЖЕННЯ НА ІНФРАСТРУКТУРУ МЕТОДАМИ МАШИННОГО НАВЧАННЯ ТА АГЕНТНИМИ ТЕХНОЛОГІЯМИ	87
45.	Бондар О.О., Козакова Н.Л. ЗАСТОСУВАННЯ МЕТОДІВ НЕЧІТКОЇ ЛОГІКИ ДЛЯ ОПТИМІЗАЦІЇ МІСЬКИХ ПЕРЕВЕЗЕНЬ В УМОВАХ НЕВИЗНАЧЕНОСТІ	89
46.	Бочек М.Ю., Шишканова Г.А., Коротунова О.В. МАТЕМАТИЧНІ ОСНОВИ ПРОЕКТУВАННЯ АПАРАТНОГО ПРИСКОРЮВАЧА ДЛЯ ЗГОРТКОВОГО ШАРУ НЕЙРОННОЇ МЕРЕЖІ	91
47.	Вакульчик С.О., Байбуз О.Г. ОПТИМІЗАЦІЯ ВИБОРУ ФУНКЦІЙ НАЛЕЖНОСТІ ТА МЕТОДУ ДЕФУЗЗИФІКАЦІЇ ДЛЯ ПІДВИЩЕННЯ НАДІЙНОСТІ FANP-СИСТЕМИ	93
48.	Васько А. О., Вовк С. М. ВИКОРИСТАННЯ АУГМЕНТАЦІЇ ДАНИХ ДЛЯ ПІДВИЩЕННЯ СТІЙКОСТІ НЕЙРОННИХ МЕРЕЖ	95
49.	Вишенський В.І., Чикрій А.О. МОДЕЛЮВАННЯ РОЗВ'ЯЗКУ ІГРОВИХ ЗАДАЧ ПЕРЕХОПЛЕННЯ КЕРОВАНИХ ЦІЛЕЙ	97
50.	Віт Р.В., Мазурець О.В. ПІДХІД ДО ВИЯВЛЕННЯ ЦИФРОВОЇ ВТОМИ ЗА ПОВІДОМЛЕННЯМИ ІЗ ВИЗНАЧЕННЯМ СЕГМЕНТІВ СПІЛКУВАННЯ	99

51.	Войцехов М.О., Ізмайлова М.К. РОЗРОБЛЕННЯ ТА ОПТИМІЗАЦІЯ КРОСПЛАТФОРМНОЇ 2D-ГРИ ЖАНРУ ПАСЬЯНС ІЗ СЮЖЕТНОЮ СКЛАДОВОЮ ТА ДОДАТКОВОЮ МІНІ-ГРОЮ «ПОШУК ВІДМІННОСТЕЙ»	101
52.	Волк А.В., Саранча В.О., Сірик С.Ф. ВПЛИВ АНТИВІРУСНИХ ПРОГРАМ НА ПРОДУКТИВНІСТЬ ОПЕРАЦІЙНОЇ СИСТЕМИ	103
53.	Волошин Д.А., Книш Л.І. НЕЙРОМЕРЕЖЕВЕ МОДЕЛЮВАННЯ ДЛЯ ЗАДАЧ РОЗПІЗНАВАННЯ ДОКУМЕНТІВ	105
54.	Воробйов Я.Ю., Волошко В.Л. АВТОМАТИЗАЦІЯ ОЦІНЮВАННЯ ТЕСТІВ РІЗНИХ ТИПІВ В СИСТЕМІ ОФІС 365	106
55.	Гавриленко М.О., Мацуга О.М. АНАЛІЗ ЕФЕКТИВНОСТІ МЕТОДІВ-ОБГОРТОК ПІД ЧАС СТАБІЛЬНОГО ВІДБОРУ ОЗНАК У ЗАДАЧІ ПРОГНОЗУВАННЯ ОФТАЛЬМОЛОГІЧНИХ ПОКАЗНИКІВ	108
56.	Гайдуков А.С., Степанова Н.І. РОЗПІЗНАВАННЯ ОБ'ЄКТІВ У ЗОБРАЖЕННЯХ ЗА ДОПОМОГОЮ МОДЕЛІ YOLO	110
57.	Гамбаров М.М., Волошко В.Л. ШИФРУВАННЯ ТА КОДУВАННЯ ІНФОРМАЦІЇ : МАТЕМАТИЧНІ ОСНОВИ ТА ПРАКТИЧНЕ ЗАСТОСУВАННЯ	112
58.	Гарт Л.Л., Бугасенко А.О. ПРО АПРОКСИМАЦІЙНИЙ ПІДХІД ДО РОЗВ'ЯЗАННЯ ЗАДАЧІ ОПТИМАЛЬНОГО КЕРУВАННЯ РОЗПОДІЛОМ ТЕПЛА У НАПІВОБМЕЖЕНОМУ СТРИЖНІ	114
59.	Гарт Л.Л., Гарт А.В., Петров І.С. АНАЛІЗ ДИНАМІЧНИХ МОДЕЛЕЙ ОПТИМАЛЬНОГО ІНВЕСТИВАННЯ	117
60.	Гладуш А.С. АВТОМАТИЗАЦІЯ ПРОЦЕСІВ ПРИЙНЯТТЯ РІШЕНЬ У БАНКІВСЬКІЙ СФЕРІ: МОЖЛИВОСТІ, РИЗИКИ ТА ПЕРСПЕКТИВИ РОЗВИТКУ	119
61.	Глушко О.С., Золотько К.Є. ENCODER-ONLY АРХІТЕКТУРА TRANSFORMER У ЗАДАЧАХ КЛАСИФІКАЦІЇ ТЕКСТУ	121
62.	Гончаров Я.А. ВЕРИФІКАЦІЯ ТА ВІЗУАЛІЗАЦІЯ КОНТАКТНИХ РОЗРАХУНКІВ ІЗ ВИКОРИСТАННЯМ ANSYS, MATLAB ТА IN-HOUSE ІНСТРУМЕНТІВ ПОСТПРОЦЕСИНГУ ДАНИХ	123
63.	Горбачук В.М., Ніколенко Д.І., Осипенко Д.О. ПЛАТФОРМА ДЛЯ ТОРГІВ ЕЛЕКТРОЕНЕРГІЄЮ	124
64.	Грегорович М.Ю., Зайцев В.Г. СИНЕРГЕТИЧНЕ КЕРУВАННЯ АВТОНОМНИМ МОБІЛЬНИМ РОБОТОМ	126
65.	Давидов О.О., Давидов О.С. ЗАДАЧА СИЛЬВЕСТРА ТА ЇЇ УЗАГАЛЬНЕННЯ	129
66.	Деменко І.О. МУЛЬТИАГЕНТНІ ПІДХОДИ ДО АВТОМАТИЗОВАНОГО ТЕСТУВАННЯ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ	131
67.	Демиденко В.С., Волошко В.Л. РОЗВ'ЯЗУВАННЯ ЗАДАЧ ДРОБОВО-ЛІНІЙНОГО ПРОГРАМУВАННЯ З ТРАНСЦЕНДЕНТНИМИ ФУНКЦІЯМИ	133

68.	Дерев'янюк С.О., Мацуга О.М. АНАЛІЗ ПРОДУКТИВНОСТІ ОБРОБКИ ВЕЛИКИХ ДАНИХ В АРАСНЕ SPARK ТА POLARS	135
69.	Діденко Д.К., Козакова Н.Л. СИСТЕМНИЙ ПІДХІД ДО СЕГМЕНТАЦІЇ КЛІЄНТІВ ЕЛЕКТРОННОЇ КОМЕРЦІЇ НА ОСНОВІ ПОВЕДІНКОВИХ ХАРАКТЕРИСТИК	137
70.	Долотов І. О., Гук Н.А. СУЧАСНІ МЕТОДИ КЛАСИФІКАЦІЇ ВЕБСТОРИНОК	139
71.	Дробахін О.О., Олевський О.В. ПОКРАЩЕННЯ ОЦІНОК ПАРАМЕТРІВ СИГНАЛУ ЗА ДОПОМОГОЮ МЕТОДУ ПУЧКА МАТРИЦЬ ЗАВДЯКИ ПРОПУСКАННЮ ЗАШУМЛЕНИХ ТОЧОК	141
72.	Єгер М.Д., Лефтеров О.В., Стецюк П.І., Федосєєв О.І. ОПТИМАЛЬНИЙ ВИННИЙ МАРШРУТ ВІДЕНЬ – ВЕНЕЦІЯ	143
73.	Єлі М.Я., Байбуз О.Г. ВДОСКОНАЛЕННЯ АЛГОРИТМУ ДИНАМІЧНОГО РОЗПОДІЛУ РЕСУРСІВ (DARA) НА ОСНОВІ ГЛИБОКОГО НАВЧАННЯ З ПІДКРІПЛЕННЯМ ДЛЯ БАГАТОКРИТЕРІАЛЬНОХ ОПТИМІЗАЦІЇ	145
74.	Єфремов М.С., Ляшко А.В., Крак Ю.В., Стеля О.Б. ДО ТОЧНОГО ВИЯВЛЕННЯ R-ЗУБЦІВ В ЕЛЕКТРОКАРДІОГРАМАХ НА ОСНОВІ ПОРОГОВОГО МЕТОДУ	147
75.	Жушман В.В., Зайцева Т.А. ЗАСТОСУВАННЯ ТЕХНОЛОГІЇ DIGITAL TWINS ДЛЯ АНАЛІЗУ НАПРУЖЕНО-ДЕФОРМІВНОГО СТАНУ КОНТАКТУЮЧИХ ТІЛ	149
76.	Задорожний Б.О., Стецюк П.І. УДОСКОНАЛЕННЯ ЕВРИСТИЧНОГО АЛГОРИТМУ УПАКОВКИ НЕРІВНИХ КРУГІВ В КРУГ МІНІМАЛЬНОГО РАДІУСУ	150
77.	Земляний О.Д., Байбуз О.Г. ПІДВИЩЕННЯ ЯКОСТІ МОДЕЛЕЙ КЛАСИФІКАЦІЇ ТА ПРОГНОЗУВАННЯ ШЛЯХОМ ІМПУТАЦІЇ ПРОПУСКІВ	152
78.	Золотько К. Є., Шаповаленко Н.Р. МОДЕЛІ І АЛГОРИТМИ СТИСКАННЯ ЗОБРАЖЕНЬ У СИСТЕМІ ОБМІНУ ФАЙЛАМИ	155
79.	Інкін О. А., Білозьоров В. Є. БЕГ У ПЛОЩИНІ ХАОТИЧНОЇ ДИНАМІКИ	157
80.	Каманцев А.С., Гарт Л.Л. ТЕХНОЛОГІЇ АВТОНОМНОЇ НАВІГАЦІЇ ДЛЯ ЛЮДЕЙ ІЗ ВАДАМИ ЗОРУ	159
81.	Кармазін К.В. ІНТЕЛЕКТУАЛЬНІ АЛГОРИТМИ ПІДВИЩЕННЯ ЯКОСТІ КОРИСТУВАЦЬКОГО ДОСВІДУ В АДАПТИВНИХ ВІДЕОСИСТЕМАХ СИТУАЦІЙНОГО ЦЕНТРУ	162
82.	Ківа М.В., Мацуга О.М. ДОСЛІДЖЕННЯ ЕФЕКТИВНОСТІ НЕЙРОННИХ МЕРЕЖ У ЗАДАЧІ КЛАСИФІКАЦІЇ РЕАЛЬНИХ ТА ЗГЕНЕРОВАНИХ ЗОБРАЖЕНЬ ОБЛИЧ	164
83.	Кісельова О.М., Кузенков О.О., Закутній Д. ДИНАМІЧНІ ЗАДАЧІ ОПТИМАЛЬНОГО РОЗБИТТЯ МНОЖИН: МАТЕМАТИЧНІ МОДЕЛІ, МЕТОДИ, АЛГОРИТМИ, ПРАКТИЧНЕ ЗАСТОСУВАННЯ	166

84.	Кісельова О.М, Притоманова О.М., Лебедєв Д.М. МЕТОД РОЗВ'ЯЗАННЯ НЕЧІТКОЇ ДВОЕТАПНОЇ ЗАДАЧІ ОПТИМАЛЬНОГО РОЗБИТТЯ МНОЖИН	167
85.	Клецьков О.М., Шевельова А.Є., Лобода В.В. ВПЛИВ ПОВЕРХНЕВОЇ ПРУЖНОСТІ НА НАНОТРИЩИНУ В ІЗОТРОПНОМУ МАТЕРІАЛІ	168
86.	Клименко О.Д. ІНФОРМАЦІЙНА ТЕХНОЛОГІЯ ОБРОБКИ ВИМІРЮВАНЬ ТА ЇХ АНАЛІЗУ З ВИКОРИСТАННЯМ ПАРАМЕТРИЧНИХ ТА НЕПАРАМЕТРИЧНИХ КРИТЕРІЇВ	169
87.	Козін І.В., Сардак О.В. ФРАГМЕНТАРНІ МОДЕЛІ ДЛЯ ЗАДАЧ ПОКРИТТЯ МНОЖИНИ ТОЧОК	171
88.	Колесніков Д.С. БАГАТОКРИТЕРІЙНА ОПТИМІЗАЦІЯ ІНВЕСТИЦІЙНОГО ПОРТФЕЛЯ	173
89.	Компан С.В. РЕКУРСИВНІ ЗАПИТИ В SQL: СЕМАНТИКА, УМОВИ ЗАВЕРШУВАНOSTІ ТА ОПТИМІЗАЦІЯ	175
90.	Корабльов М.М. ВИКОРИСТАННЯ R-АЛГОРИТМУ ДЛЯ ЗНАХОДЖЕННЯ ПАРАМЕТРІВ ДВОКЛАСОВОЇ SVM МОДЕЛІ	176
91.	Костюченко А.Д., Герасимов В.В. ВИКОРИСТАННЯ МОДЕЛЕЙ ГЛИБОКОГО НАВЧАННЯ ДЛЯ КЛАСИФІКАЦІЇ РИТМУ СЕРЦЯ ЗА РЕЗУЛЬТАТАМИ ЕКГ	178
92.	Кочержинська А.Д., Сірик С.Ф., Лисиця Н.М., Шишканова Г.А. ЦИФРОВИЙ ЖИВОПИС ЯК ЕВОЛЮЦІЙНА ФОРМА ОБРАЗОТВОРЧОГО МИСТЕЦТВА В ЦИФРОВУ ЕПОХУ	180
93.	Кошель В.В., Єгошкін Д.І., Сірик С.Ф. АНАЛІЗ ВАРІАБЕЛЬНОСТІ СЕРЦЕВОГО РИТМУ	182
94.	Красюк М.В., Наконечна Т.В. ЗАСТОСУВАННЯ НЕЙРОННИХ МЕРЕЖ ДЛЯ ПРОГНОЗУВАННЯ ВАРТІСНИХ ПОКАЗНИКІВ ОБ'ЄКТІВ	184
95.	Кривцов О.С., Шевельова А.Є. ПОРІВНЯЛЬНИЙ АНАЛІЗ ЕВРИСТИЧНИХ АЛГОРИТМІВ РОЗВ'ЯЗАННЯ ЗАДАЧІ ОПТИМАЛЬНОГО ПРОЄКТУВАННЯ РЕЗЕРВУАРІВ ПІД ТИСКОМ	186
96.	Кузенков О.О., Козакова Н.Л., Григоренко О., Балейко Н.В. МАТЕМАТИЧНЕ МОДЕЛЮВАННЯ ЗАДАЧ ОПТИМАЛЬНОГО РОЗБИТТЯ ТА ПЛАНУВАННЯ В УМОВАХ НЕВИЗНАЧЕНОСТІ	188
97.	Кузенков О.О., Козакова Н.Л., Лупинський С., Балейко Н.В. АНАЛІЗ ТА КЕРУВАННЯ БІФУРКАЦІЯМИ В МОДЕЛЯХ ПРИРОДНИХ ПРОЦЕСІВ РОЗМІЩЕННЯ-РОЗБИТТЯ	189
98.	Кузенков О.О., Козакова Н.Л., Шведов В., Сірик С.Ф. МЕТОДИ РОЗВ'ЯЗАННЯ ДИНАМІЧНИХ ЗАДАЧ ОПТИМАЛЬНОГО РОЗМІЩЕННЯ-РОЗБИТТЯ	190
99.	Кузенков О.О., Лозовський А.В. ДИНАМІЧНІ ЗАДАЧІ ОПТИМАЛЬНОГО РОЗБИТТЯ МНОЖИН ЗІ ЗМІННИМ ПОПИТОМ	191
100.	Кузнєцов В.О., Крак Ю.В., Куляс А.І., Кудін Г.І. ДО РОЗРОБКИ АЛГОРИТМІВ КЛАСИФІКАЦІЇ ЖЕСТОВОЇ МОВИ НА ОСНОВІ ПСЕВДООБЕРНЕННЯ І ГРУПУВАННЯ ОЗНАК	193

101.	Кузьменко В.І., МЕТОДИ УМОВНОЇ ОПТИМІЗАЦІЇ У МОДЕЛЯХ РОЗШАРУВАННЯ ГІРСЬКИХ ПОРІД	195
102.	Куліш С.П., Ткаченко О.М. СЕМАНТИЧНИЙ АНАЛІЗ РЕКЛАМНИХ ВІДЕО ЗА ДОПОМОГОЮ ВЕЛИКИХ МОВНИХ МОДЕЛЕЙ	196
103.	Лавушник Я.П., Степанова Н.І. АЛГОРИТМИ РОЗВ'ЯЗАННЯ СЛАЙДІНГ-ПАЗЛІВ	198
104.	Лебєдєва Т.Т., Семенова Н.В.,Сергієнко Т.І. СТІЙКІСТЬ ТА РЕГУЛЯРИЗАЦІЯ ЗА ОБМЕЖЕННЯМИ ВЕКТОРНИХ ЗАДАЧ ЦІЛОЧИСЛОВОЇ ОПТИМІЗАЦІЇ	200
105.	Ленський М.М., Михальчук Г.Й. ТОЧНИЙ АЛГОРИТМ ДЛЯ ЗАДАЧІ CVRP НА ОСНОВІ МЕТОДУ BRANCH-PRICE-AND-CUT З GPU-ПРИСКОРЕННЯМ	202
106.	Лісняк В.С., Волошко В.Л. ЗАСТОСУВАННЯ МЕТОДІВ ОПТИМІЗАЦІЇ НУЛЬОВОГО ТА ПЕРШОГО ПОРЯДКІВ ДЛЯ РОЗВ'ЯЗУВАННЯ ЕКОНОМІЧНИХ ЗАДАЧ	204
107.	Логвин Д.А., Божуха Л.М. МЕТРИКИ ОЦІНЮВАННЯ ЕФЕКТИВНОСТІ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ ДЛЯ ОБРОБКИ ПРИРОДНОЇ МОВИ	206
108.	Луг О.О., Кузенков О.О. ЗАСТОСУВАННЯ ГЛИБИННОГО НАВЧАННЯ ДЛЯ СЕМАНТИЧНОГО АНАЛІЗУ ТЕКСТІВ ВАКАНСІЙ	208
109.	Ляшко А.В., Стеля О.Б., Куляс А.І., Єфремов М.С., Крак Ю.В. ПРОГРАМНИЙ ІНСТРУМЕНТАРІЙ ДЛЯ АНАЛІЗУ І ВІЗУАЛІЗАЦІЇ ХАРАКТЕРИСТИЧНИХ ОЗНАК В ЕЛЕКТРОКАРДІОГРАМАХ	209
110.	Магас О.С., Гук Н.А. ОГЛЯД ЗАСТОСУВАННЯ МЕТОДІВ МАШИННОГО НАВЧАННЯ ДО ОБЕРНЕНОЇ ЗАДАЧІ ІДЕНТИФІКАЦІЇ ПАРАМЕТРІВ ТОНКОСТІННОЇ ОБОЛОНКИ	211
111.	Мажара К.О., Трофімов О.В. ВИКОРИСТАННЯ МЕТОДУ МОНТЕ-КАРЛО ДЛЯ ГЕНЕРАЦІЇ ПОЧАТКОВИХ УМОВ У ПРЯМІЙ ЗАДАЧІ ДЛЯ БАГАТОШАРОВОГО ПАКЕТУ НА ТВЕРДІЙ ОСНОВІ	213
112.	Максимів І.А., Басюк Т.М. ПРОЄКТ ІНФОРМАЦІЙНОЇ СИСТЕМИ НАДАННЯ РЕКОМЕНДАЦІЙ З ФОРМУВАННЯ ОСОБИСТОГО СТИЛЮ	215
113.	Мелашенко О.П., Романова Т.Є., Дюрягина З.А., Кравченко О.В. ОПТИМІЗАЦІЯ КОМПОНУВАННЯ М'ЯКИХ ЕЛІПСОЇДІВ	217
114.	Миргородський А.В., Ліщинська Л.Б. ЗАСТОСУВАННЯ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ В РОЗПОДІЛЕНИХ БАЗАХ ДАНИХ	219
115.	Мірзаєв Т. Р., Михальчук Г. Й. ДВОЕТАПНИЙ МЕТОД ОПТИМІЗАЦІЇ МАРШРУТІВ ДЛЯ ГЕТЕРОГЕННИХ ТРАНСПОРТНИХ СИСТЕМ	221
116.	Моїсєєнков Д.В., Зайцев В.Г. МАТЕМАТИЧНЕ МОДЕЛЮВАННЯ ДИНАМІКИ ЗВЕРНЕНЬ ДО МЕДИЧНОГО ЦЕНТРУ З ЕЛЕМЕНТАМИ АДАПТИВНОГО ПРОГНОЗУВАННЯ ТА ВИЯВЛЕННЯ АНОМАЛІЙ	223
117.	Мойсєєнко В.М., Антоненко С.В. ВИКОРИСТАННЯ ВЕЛИКИХ МОВНИХ МОДЕЛЕЙ (LLM) ДЛЯ ОПТИМІЗАЦІЇ БІЗНЕС-ПРОЦЕСІВ ПІДПРИЄМСТВА	225

118. Молчанов А.О. ЗБАГАЧЕННЯ БАЗИСУ В МЕТОДІ ФУНДАМЕНТАЛЬНИХ РОЗВ'ЯЗКІВ ПОТЕНЦІАЛАМИ ПОДВІЙНОГО ШАРУ З В-СПЛАЙНОВОЮ ЩІЛЬНІСТЮ	227
119. Мороз В.В., Кулик Д.В. МОДИФІКОВАНИЙ АЛГОРИТМ WINDOW-EMD ІЗ ННТ ДЛЯ РОЗПІЗНАВАННЯ КОРИСНИХ ТА ШУМОВИХ КОМПОНЕНТ У МОВНИХ АУДІОСИГНАЛАХ	229
120. Морозов Ю.С., Зайцева Т.А. НЕЙРОМЕРЕЖЕВІ МЕТОДИ ПІДВИЩЕННЯ ТОЧНОСТІ ТА ШВИДКОДІЇ МСЕ У ЗАДАЧАХ КОНТАКТНОЇ МЕХАНІКИ	231
121. Наконечна Т.В. ПОБУДОВА МОДЕЛІ СТРУКТУРИ ТА ОПИСУ СИСТЕМИ ДИСТАНЦІЙНОГО НАВЧАННЯ	233
122. Новіков С.О., Антоненко С.В., Ізмайлова М.К. ДОСЛІДЖЕННЯ МЕТОДІВ ПЕРЕДАЧІ ДАНИХ ЧЕРЕЗ ІНТЕРНЕТ У РОЗПОДІЛЕНИХ СИСТЕМАХ З МОБІЛЬНИМ КЛІЄНТОМ	235
123. Овсов М. В., Верба О. В., Сафронова І. А., Зайцева Т.А. МАТЕМАТИЧНЕ МОДЕЛЮВАННЯ ТА АВТОМАТИЗОВАНА СИСТЕМА РОЗРАХУНКУ НАВЧАЛЬНОГО НАВАНТАЖЕННЯ ЗАКЛАДУ ВИЩОЇ ОСВІТИ	237
124. Овчаренко Д.О., Михальчук Г.Й. ПРОЄКТУВАННЯ ТА РЕАЛІЗАЦІЯ СЕРВЕРНОЇ ЧАСТИНИ ВЕБПЛАТФОРМИ ДЛЯ ПІДТРИМКИ НАВЧАЛЬНОГО ПРОЦЕСУ	239
125. Овчарук О.М., Мазурець О.В. ПІДХІД ДО ОЦІНЮВАННЯ ПОВЕДІНКОВИХ РИЗИКІВ КОРИСТУВАЧІВ З ПТСР ЗА ЦИФРОВИМ ПРОФІЛЕМ ЗАСОБАМИ NLP	241
126. Орлов С.К., Наконечна Т.В. АВТОМАТИЗОВАНЕ РОЗТАШУВАННЯ МОДЕЛЕЙ НА ЦИЛІНДРИЧНІЙ ПЛАТФОРМІ 3D-ПРИНТЕРА ТИПУ FUGO	243
127. Павлів М.Я., Басюк Т.М. ПРОЄКТ ІНФОРМАЦІЙНОЇ СИСТЕМИ ДЛЯ АВТОМАТИЗОВАНОГО УПРАВЛІННЯ ЗАМОВЛЕННЯМИ ТА ДОСТАВКИ ФЕРМЕРСЬКОЇ ПРОДУКЦІЇ	245
128. Павлюк Д.І., Байбуз О.Г. МУЛЬТИАГЕНТНА СИСТЕМА З RAG ДЛЯ МОДЕРАЦІЇ КОРИСТУВАЦЬКИХ ПОВІДОМЛЕНЬ	247
129. Панасюк Б.Ю., Ліщинська Л.Б. ІНТЕЛЕКТУАЛЬНА НОРМАЛІЗАЦІЯ ПОВІДОМЛЕНЬ У МУЛЬТИФОРМАТНИХ ПОДІЄВИХ СИСТЕМАХ	249
130. Панкратов Є.В. ПРОЄКТУВАННЯ ЦИФРОВИХ МОДУЛЬНИХ СИСТЕМ ДЛЯ АВТОМАТИЗОВАНОГО КОНТРОЛЮ ТЕХНОЛОГІЧНИХ ПРОЦЕСІВ	251
131. Пасічник А.М., Пасічник В.А. ДИНАМІКА РОЗВИТКУ ІНФОРМАЦІЙНИХ ТЕХНОЛОГІЙ І СТВОРЕННЯ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ	253
132. Пашенко А.Ю., Книш Л.І. НЕЙРОМЕРЕЖЕВЕ МОДЕЛЮВАННЯ ПРОЦЕСУ КОМП'ЮТЕРНОЇ ГРИ	255
133. Перемітько М.В., Надригайло Т.Ж. РЕАЛІЗАЦІЯ ОБМІНУ ДАНИМИ МІЖ МОБІЛЬНИМ ДОДАТКОМ ТА ЕЛЕКТРОННИМ ТОНОМЕТРОМ	257

134. Петров А.Д., Трофімов О.В. МЕТОД СПРЯЖЕНОГО СТАНУ В АНАЛІЗІ ОБЕРНЕНИХ ЗАДАЧ ДЛЯ БАГАТОШАРОВИХ ОСНОВ	259
135. Пилипенко О.О., Шишканова Г.А., Коротунова О.В. РОЛЬ МАТЕМАТИЧНОГО ЗАБЕЗПЕЧЕННЯ ПРИ ПРОЕКТУВАННІ МЕМС	261
136. Подвинцев В.В., Козакова Н.Л. ПРОГНОЗУВАННЯ ПРИБУТКУ ВІД КРЕДИТНИХ ПРОДУКТІВ НА ОСНОВІ МЕТОДІВ МАШИННОГО НАВЧАННЯ	263
137. Політов А.І., Вовк С.М. ВПЛИВ ЗМАГАЛЬНИХ АТАК НА СТІЙКІСТЬ НЕЙРОННИХ МЕРЕЖ	265
138. Попов О. В., Павлюк А.В. ВАРІАЦІЙНІ МЕХАНІЗМИ АДАПТИВНОГО КЕРУВАННЯ ПРОДУКТИВНІСТЮ В ГІБРИДНИХ СИСТЕМАХ	267
139. Пригода О.О., Золотько К. Є. МОДЕЛЮВАННЯ ДИНАМІКИ ПАНДЕМІЙ ЗА ДОПОМОГОЮ АТТРАКТОРІВ І МЕТОДУ SINDY	269
140. Ріпка Є.В., Сафронова І.А. ПІДВИЩЕННЯ МАСШТАБОВАНOSTІ КОНТРОЛЮ ЦІЛІСНОСТІ ФАЙЛІВ У ХМАРНИХ СИСТЕМАХ НА ОСНОВІ BLAKE2B ТА ДЕРЕВА МЕРКЛА	271
141. Рябоволенко В.А., Байбуз О.Г. РОЗРОБКА СИСТЕМИ МОНІТОРИНГУ DOCKER-КОНТЕЙНЕРІВ	273
142. Сафіюліна Я.О., Степанова Н.І. НОРМАЛІЗАЦІЯ ГУЧНОСТІ АУДИОСИГНАЛУ ДЛЯ СИСТЕМ БІОМЕТРИЧНОЇ ІДЕНТИФІКАЦІЇ ЗА ГОЛОСОМ	275
143. Селімханов Т.Б., Дзюба П.А. РОЗРОБКА РЕКОМЕНДАЦІЙНОЇ СИСТЕМИ НА ОСНОВІ ПОВЕДІНКОВИХ ТА КОНТЕНТНИХ ОЗНАК	277
144. Семенов В.В., Чернорай В.О. АЛГОРИТМИ ОПЕРАТОРНОЇ ЕКСТРАПОЛЯЦІЇ ДЛЯ ПСЕВДОМОНОТОННИХ ВАРІАЦІЙНИХ НЕРІВНОСТЕЙ	279
145. Сизоненко Р.М., Клименко С.В. СУЧАСНІ ПІДХОДИ ДЕТЕКТУВАННЯ ОБ'ЄКТІВ НА БОРТУ БПЛА	280
146. Сироїд В.П., Басюк Т.М. ІНТЕГРАЦІЯ ПЕРЕВІРОК БЕЗПЕКИ В CI/CD КОНВЕЄРИ НА БАЗІ GITHUB ACTIONS, DOCKER ТА TRIVY	282
147. Сімакін С.К., Божуха Л.М. ОПТИМІЗАЦІЯ СПОЖИВАННЯ РЕСУРСІВ У ХМАРНИХ ІНФРАСТРУКТУРАХ ЗА ДОПОМОГОЮ ЕНЕРГОМОДЕЛЕЙ	284
148. Соломатін В.А., Байбуз О.Г. АРХІТЕКТУРА BIOMORPHNET ДЛЯ АВТОМАТИЗОВАНОГО ВИЯВЛЕННЯ ІШЕМІЧНОЇ ХВОРОБИ СЕРЦЯ ЗА РЕНТГЕНІВСЬКИМИ ЗНІМКАМИ	286
149. Стецюк П.І., Стороженко А.О., Хом'як О.М. ЛП-ЗАДАЧА ДЛЯ ЗНАХОДЖЕННЯ РОЗРІДЖЕНОГО ВЕКТОРА	288
150. Стецюк П.І., Тиводар С.Р., Кузнєва А.С. ПРО ПРИСКОРЕНЕ РОЗВ'ЯЗАННЯ ЗАДАЧІ ОПТИМАЛЬНОЇ УПАКОВКИ НЕРІВНИХ КРУГІВ У КВАДРАТ	289
151. Стовба В.О., Жмуд О.О. R-АЛГОРИТМ ДЛЯ ЗАДАЧ ЕНТРОПІЙНО-ЛІНІЙНОГО ПРОГРАМУВАННЯ	291

152. Стружко В.Р., Антоненко С.В. ПРОЄКТУВАННЯ МУЛЬТИМОДАЛЬНОЇ СТЕГANOГРАФІЧНОЇ СИСТЕМИ ДЛЯ ПРИХОВУВАННЯ ІНФОРМАЦІЇ В ГЕТЕРОГЕННИХ ЦИФРОВИХ КОНТЕЙНЕРАХ (ЗОБРАЖЕННЯ ТА АУДІО)	293
153. Сулейманов Є.С., Байбуз О.Г., Зайцева Т.А. НЕЙРОННІ МЕРЕЖІ У ЖАНРОВІЙ КЛАСИФІКАЦІЇ МУЗИЧНОГО КОНТЕНТУ	295
154. Тітяпкин А.С. БАЗИ ДАНИХ СИСТЕМ ЕКОЛОГІЧНОГО МОНІТОРИНГУ ЩОДО СТАНУ ЕВТРОФУВАННЯ ВОД ПІВНІЧНО-ЗАХІДНОЇ ЧАСТИНИ ЧОРНОГО МОРЯ ТА ВИЗНАЧЕННЯ ТЕНДЕНЦІЙ ЙОГО ЗМІНИ	297
155. Тіщук А.О., Тонкошкур І.С. ЧИСЕЛЬНЕ МОДЕЛЮВАННЯ ДИНАМІЧНИХ ПРОЦЕСІВ ПЕРЕНОСУ НА ОСНОВІ ОДНОВИМІРНИХ МОДЕЛЕЙ	299
156. Тонкошкур І.С. ЧИСЕЛЬНИЙ АНАЛІЗ АЕРОДИНАМІЧНИХ ХАРАКТЕРИСТИК НЕКРУГОВИХ КОНУСІВ	300
157. Топольськов Є.О., Москаленко Н.В. ІНСТРУМЕНТАЛЬНІ ЗАСОБИ РОЗРОБКИ ЕКСПЕРТНИХ СИСТЕМ	301
158. Турчина В.А., Драниця М.В. ЗАСТОСУВАННЯ АЛГОРИТМУ ДЕЙКСТРИ ДЛЯ УЗАГАЛЬНЕНОЇ ЗАДАЧІ ПРО МІНІМАЛЬНІ ШЛЯХИ	303
159. Фунтиков М.К., Мацуга О.М. ПОРІВНЯННЯ ВНУТРІШНІХ ІНДЕКСІВ ЯКОСТІ КЛАСТЕРИЗАЦІЇ В ЗАДАЧІ ГРУПУВАННЯ ОЗНАК	305
160. Циба В.В., Пасічник А.М. МАТЕМАТИЧНЕ МОДЕЛЮВАННЯ ТА ОБЧИСЛЕННЯ ПРОСТОРОВОЇ ДИНАМІКИ ТРАЄКТОРІЇ ПОЛЬОТУ АРТИЛЕРІЙСЬКОГО СНАРЯДУ	307
161. Чаша М.Ю., Дзюба П.А. МОДЕЛІ І АЛГОРИТМИ ОПТИМІЗАЦІЇ РУХУ АВТОНОМНИХ РОБОТІВ У ЗАМКНЕНОМУ ПРОСТОРІ	309
162. Червоняк А.Б., Полягушко Л.Г., Сліпченко В.Г., МЕТОДИ МАШИННОГО НАВЧАННЯ ДЛЯ АДАПТИВНОЇ ПЕРСОНАЛІЗАЦІЇ ВЕБ ІНТЕРФЕЙСІВ У РЕАЛЬНОМУ ЧАСІ	311
163. Чугай А.М., Яськов Г.М., Старкова О.В., Щербина М.О. КОМПОЗИТНЕ ГЕОМЕТРИЧНЕ МОДЕЛЮВАННЯ ДЛЯ БРАХІТЕРАПІЇ З АНІЗОТРОПНИМ ВИПРОМІНЮВАННЯМ: ОПТИМІЗАЦІЙНИЙ ПІДХІД ДО РОЗМІЩЕННЯ КАПСУЛ	313

Підп. до друку 12.11.2025 р. Формат 60х84/16. Друк цифровий.

Папір офсетний. Гарнітура Times. Ум.-друк. арк. 20,25.

Наклад 100 прим. Зам. № 173

ПП «Ліра ЛТД»

49107, м. Дніпро, вул. Наукова, 5.

Свідоцтво про внесення до Державного реєстру ДК № 6042 від 26.02.2018 р.