

ISSN 2786-5460 (Print)
ISSN 2786-5479 (Online)

МІЖРЕГІОНАЛЬНА АКАДЕМІЯ УПРАВЛІННЯ ПЕРСОНАЛОМ
INTERREGIONAL ACADEMY OF PERSONNEL MANAGEMENT



ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ ТА СУСПІЛЬСТВО

INFORMATION TECHNOLOGY AND SOCIETY

Випуск 2 (17), 2025
Issue 2 (17), 2025



Видавничий дім
«Гельветика»
2025

*Рекомендовано до друку Вченою радою
Міжрегіональної Академії управління персоналом
(протокол № 7 від 29 серпня 2025 року)*

Інформаційні технології та суспільство / [головний редактор О. Попов]. – Київ : Міжрегіональна Академія управління персоналом, 2025. – Випуск 2 (17). – 226 с.

Журнал «Інформаційні технології та суспільство» є науковим рецензованим виданням, в якому здійснюється публікація матеріалів науковців різних рівнів у вигляді наукових статей з метою їх поширення як серед вітчизняних дослідників, так і за кордоном.

Редакційна колегія не обов'язково поділяє позицію, висловлену авторами у статтях, та не несе відповідальності за достовірність наведених даних і посилань.

Головний редактор: **Попов О. О.** – член-кор. НАН України, д-р техн. наук, професор, в.о. директора Центру інформаційно-аналітичного та технічного забезпечення моніторингу об'єктів атомної енергетики Національної академії наук України.

Редакційна колегія:

Василенко М. Д. – д-р фіз.-мат. наук, проф., професор кафедри кібербезпеки, Національний університет «Одеська юридична академія»; **Горбов І. В.** – канд. техн. наук, с.н.с., старший науковий співробітник, Інститут проблем реєстрації інформації НАН України; **Дуднік А. С.** – д-р техн. наук, доц., доцент кафедри мережних та інтернет технологій, Київський національний університет імені Тараса Шевченка; **Євсєєв С. П.** – д-р техн. наук, лауреат національної премії імені Патона 2024 р., професор кафедри кібербезпеки, Національний технічний університет «ХПІ»; **Зибін С. В.** – д-р техн. наук, доц., завідувач кафедри інженерії програмного забезпечення, Національний авіаційний університет; **Кавун С. В.** – д-р екон. наук, канд. техн. наук, проф., завідувач кафедри комп'ютерних інформаційних систем та технологій, Міжрегіональна Академія управління персоналом; **Комарова Л. О.** – д-р техн. наук, с.н.с., директор Навчально-наукового інституту інформаційної безпеки та стратегічних комунікацій, Національна академія Служби безпеки України; **Охріменко Т. О.** – канд. техн. наук, старший науковий співробітник науково-дослідної лабораторії протидії кіберзагрозам в авіаційній галузі, Національний авіаційний університет; **Рудніченко М. Д.** – канд. техн. наук, доц., доцент кафедри інформаційних технологій, Державний університет «Одеська політехніка»; **Скуратовський Р. В.** – канд. фіз.-мат. наук, доц., доцент кафедри обчислювальної математики та комп'ютерного моделювання, Міжрегіональна Академія управління персоналом; **Супрун О. М.** – канд. фіз.-мат. наук, доц., доцент кафедри програмних систем і технологій, Київський національний університет імені Тараса Шевченка; **Табунщик Г. В.** – канд. техн. наук, проф., професор кафедри програмних засобів, Національний університет «Запорізька політехніка»; **Фомін О. О.** – д-р техн. наук, доц., професор кафедри комп'ютеризованих систем управління, професор кафедри прикладної математики та інформаційних технологій, Державний університет «Одеська політехніка»; **Хохлачова Ю. С.** – канд. техн. наук, доц., доцент кафедри безпеки інформаційних технологій, Національний авіаційний університет; **Чолишкіна О. Г.** – канд. техн. наук, доц., доцент кафедри Інтелектуальних технологій, Київський національний університет імені Тараса Шевченка; **Чорний О. П.** – доктор технічних наук, професор, директор Навчально-наукового інституту електричної інженерії та інформаційних технологій, Кременчуцький національний університет імені Михайла Остроградського; **Юдін О. К.** – д-р техн. наук, проф., директор центру кібербезпеки Навчально-наукового інституту інформаційної безпеки та стратегічних комунікацій, Національна академія Служби безпеки України; **Гопеснко Віктор** – dr. sc. ing., проф., проректор з наукової роботи, директор навчальної програми магістратури «Комп'ютерні системи», Університет прикладних наук ISMA (Латвійська Республіка); **Leszczyna Rafal** – dr hab. inż., професор кафедри комп'ютерних наук у менеджменті, Гданський технологічний університет (Республіка Польща).

Реєстрація суб'єкта у сфері друкованих медіа:

Рішення Національної ради України з питань телебачення і радіомовлення № 1173 від 11.04.2024 року.

Ідентифікатор медіа: R30-03890

Суб'єкт у сфері друкованих медіа – Приватне акціонерне товариство «Вищий навчальний заклад «Міжрегіональна Академія управління персоналом» (вул. Фрометівська, буд. 2, м. Київ, 03039, iart@iart.edu.ua, тел. (044) 490-95-00).

Мова видання: українська, англійська, німецька, французька та польська.

Періодичність видання: 4 рази на рік.

Відповідно до Наказу МОН України № 1290 від 30 листопада 2021 року (додаток 3) журнал включено до Переліку наукових фахових видань України (категорія Б) зі спеціальностей F2 – Інженерія програмного забезпечення; F3 – Комп'ютерні науки; F4 – Системний аналіз та наука про дані; F5 – Кібербезпека та захист інформації; F6 – Інформаційні системи і технології; F7 – Комп'ютерна інженерія.

Усі електронні версії статей журналу оприлюднюються на офіційній сторінці видання
<http://journals.maup.com.ua/index.php/it>

Статті у виданні перевірені на наявність плагіату за допомогою програмного забезпечення
StrikePlagiarism.com від польської компанії Plagiat.pl.

**Recommended for publication
by Interregional Academy of Personnel Management
(Minutes No. 7 dated 29 August 2025)**

Information Technology and Society / [chief editor Oleksandr Popov]. – Kyiv : Interregional Academy of Personnel Management, 2025. – Issue 2 (17). – 226 p.

Journal «Information Technology and Society» is a peer-reviewed scientific edition, which publishes materials of scientists of various levels in the form of scientific articles for the purpose of their dissemination both among domestic researchers and abroad.

Editorial board do not necessarily reflect the position expressed by the authors of articles, and are not responsible for the accuracy of the data and references.

Chief editor: Oleksandr Popov – Corresponding Member of NAS of Ukraine, Doctor of Engineering, Professor, Acting Director of the Center for Information-Analytical and Technical Support of Nuclear Power Facilities Monitoring of the National Academy of Sciences of Ukraine.

Editorial Board:

Mykola Vasylenko – Doctor of Physics and Mathematics, Professor, Professor at the Department of Cybersecurity, National University «Odesa Law Academy»; **Ivan Horbov** – PhD in Engineering, Senior Research Associate, Senior Research Fellow, Institute for Information Recording of NAS of Ukraine; **Andrii Dudnik** – Doctor of Engineering, Associate Professor, Senior Lecturer at the Department of Networking and Internet Technologies, Taras Shevchenko National University of Kyiv; **Serhii Yevseiev** – Doctor of Engineering, laureate of the 2024 National Prize of Ukraine named after Borys Paton, Professor at the Department of Cybersecurity, National Technical University “Kharkiv Polytechnic Institute”; **Serhii Zybin** – Doctor of Engineering, Associate Professor, Head of the Department of Software Engineering, National Aviation University; **Serhii Kavun** – Doctor of Economics, PhD in Engineering, Professor, Head of the Department of Computer Information Systems and Technologies Interregional Academy of Personnel Management; **Larysa Komarova** – Doctor of Engineering, Senior Research Scientist, Laureate of State Prize, Director of Educational-Scientific Institute of Information Security and Strategic Communications, National Academy of the Security Service of Ukraine; **Tetiana Okhrimenko** – PhD in Engineering, Senior Research Scientist at the Scientific Research Laboratory for Countering Aviation Cyberthreats, National Aviation University; **Mykola Rudnichenko** – PhD in Engineering, Associate Professor, Senior Lecturer at the Department of Information Technologies, Odessa Polytechnic State University; **Ruslan Skuratovskyi** – PhD in Physics and Mathematics, Associate Professor, Senior Lecturer at the Department of Computational Mathematics and Computer Modeling, Interregional Academy of Personnel Management; **Olha Suprun** – PhD in Physics and Mathematics, Associate Professor, Senior Lecturer at the Department of Software Systems and Technologies, Taras Shevchenko National University of Kyiv; **Halyna Tabunshchyk** – PhD in Engineering, Professor, Professor at the Department of Software Tools, “Zaporizhzhia Polytechnic” National university; **Oleksandr Fomin** – Doctor of Engineering, Associate Professor, Professor at the Department of Computerized Control Systems, Professor at the Department of Applied Mathematics and Information Technologies, Odessa Polytechnic State University; **Yuliia Khokhlachova** – PhD in Engineering, Associate Professor, Senior Lecturer at the Department of Information Technology Security, National Aviation University; **Olha Cholyshkina** – PhD in Engineering, Associate Professor, Associate Professor at the Department of Intellectual Technologies, Taras Shevchenko National University of Kyiv; **Oleksii Chorny** – Doctor of Technical Sciences, Professor, Director of the Educational and Scientific Institute of Electrical Engineering and Information Technologies, Kremenchuk National University named after Mykhailo Ostrogradskyi; **Oleksandr Yudin** – Doctor of Engineering, Professor, Director of the Cybersecurity Center of the Educational-Scientific Institute of Information Security and Strategic Communications, National Academy of the Security Service of Ukraine; **Hopeienko Viktor** – dr. sc. ing., Professor, Vice Rector for Research, Director of the study programme “Computer systems”, ISMA University of Applied Sciences (Republic of Latvia); **Leszczyna Rafal** – dr hab. inż., Profesor, Katedra Informatyki w Zarządzaniu, Politechnika Gdańska (Republic of Poland).

Registration of Print media entity:

*Decision of the National Council of Television and Radio Broadcasting of Ukraine: Decision No. 1173 as of 11.04.2024.
Media ID: R30-03890*

Media entity – Private Joint-Stock Company «Higher education institution «Interregional Academy of Personnel Management» (03039, Kyiv, Frometivska str., 2, iapm@iapm.edu.ua, tel. (044) 490-95-00).

*Language of publication: Ukrainian, English, German, French, and Polish.
Periodicity: 4 times a year.*

According to the Decree of MES No. 1290 (Annex 3) dated November 30, 2021, the journal was included in the List of scientific professional publications of Ukraine (category B) in specialties F2 – Software Engineering; F3 – Computer Sciences; F4 – Systems Analysis and Data Science; F5 – Cybersecurity and Data Protection; F6 – Information Systems and Technologies; F7 – Computer Engineering.

All electronic versions of articles in the collection are available on the official website edition
<http://journals.maup.com.ua/index.php/it>

The articles were checked for plagiarism using the software
StrikePlagiarism.com developed by the Polish company Plagiat.pl.

ЗМІСТ

Олександр АРТАМОНОВ, Михайло ЛУЧКЕВИЧ ФОРМАЛІЗАЦІЯ ВИМОГ ДО ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ ПІДТРИМКИ УХВАЛЕННЯ РІШЕНЬ У ЗАДАЧАХ ОПТИМІЗАЦІЇ ХМАРНИХ РЕСУРСІВ.....	8
Микола БІЛИЙ, Євген КРИЛОВ МОДИФІКАЦІЯ МЕТОДУ RECIPROCAL RANK FUSION ДЛЯ ПОЛІПШЕННЯ РЕЗУЛЬТАТІВ ГІБРИДНОГО ПОШУКУ У ІНФОРМАЦІЙНИХ СИСТЕМАХ З ВЕКТОРНИМИ БАЗАМИ ДАНИХ.....	14
Олег ГАРАСИМЧУК, Михайло ЛУЧКЕВИЧ ПІДХОДИ DESIGN SCIENCE RESEARCH ДО ПОБУДОВИ ІНТЕЛЕКТУАЛЬНОЇ СИСТЕМИ КОНТРАКТ-ТЕСТУВАННЯ	20
Mykhailo HNATYSHYN, Oleksii NEDASHKIVSKIY A FRAMEWORK FOR EXPLAINABLE AI (XAI) IN MACHINE LEARNING-BASED FAKE NEWS DETECTION SYSTEMS: ENHANCING TRANSPARENCY, TRUST, AND USER AGENCY.....	24
Галина ГРИГОРЧУК, Любомир ГРИГОРЧУК, Вікторія БАНДУРА АВТОМАТИЗАЦІЯ РЕГУЛЮВАННЯ ВОЛОГОСТІ ПІД ЧАС ОСУШУВАННЯ УТФЕЛЮ.....	30
Олег ЖВАЛІК ДОСЛІДЖЕННЯ ІНТЕГРАЦІЇ ІТ В УПРАВЛІННІ СОЛЯНОЮ ЕНЕРГЕТИКОЮ.....	39
Олександр ЗАДЕРЕЙКО, Світлана ЄЛЕНИЧ, Наталія ЛОГІНОВА, Олена ТРОФИМЕНКО, Володимир ГУРА АНАЛІЗ РЕЛЕВАНТНОСТІ ПОШУКОВОЇ ВИДАЧІ НАУКОМЕТРИЧНИХ БАЗ ДАНИХ.....	47
Володимир КАРАБЧУК МЕТОДИ ВИЯВЛЕННЯ ТА БОРОТЬБИ З ДЕЗІНФОРМАЦІЄЮ ЗА ДОПОМОГОЮ ТЕХНОЛОГІЙ ШТУЧНОГО ІНТЕЛЕКТУ.....	56
Юрій КІШ, Ігор ЛЯХ ЕВОЛЮЦІЯ ПОНЯТЬ РИЗИКУ ТА ЇХ ВЗАЄМОЗВ'ЯЗОК ІЗ ЗАБЕЗПЕЧЕННЯМ ЯКОСТІ ІТ-ПРОДУКТІВ.....	62
Євгеній КЛИМЕНКО, Олена ГЛАЗУНОВА АРХІТЕКТУРА ІНФОРМАЦІЙНОЇ ТЕХНОЛОГІЇ ОСВІТНЬОЇ АНАЛІТИКИ З ВИКОРИСТАННЯМ ІНТЕЛЕКТУАЛЬНОГО АНАЛІЗУ ДАНИХ.....	69
Андрій КОБИЛЮК ПІДВИЩЕННЯ ЕФЕКТИВНОСТІ УПРАВЛІННЯ ХМАРНИМИ РЕСУРСАМИ ЗА ДОПОМОГОЮ РОЗПОДІЛЕНОГО МОНІТОРИНГУ	76
Artem KOLOMYTSEV, Yuliia KUZNETSOVA THE ROLE OF MLOPS IN ADVANCING GREEN ENERGY: AN EVALUATION OF TECHNOLOGIES AND PRACTICES.....	83
Євген ЛАНСЬКИХ, Олексій НАУМОВ СУЧАСНІ МЕТОДИ ЦИФРОВОГО ВОДЯНОГО МАРКУВАННЯ ЗОБРАЖЕНЬ: КЛАСИФІКАЦІЯ, АНАЛІЗ І ТЕНДЕНЦІЇ РОЗВИТКУ	89
Єлизавета ЛИТЮГА, Валерій ЯЦЕНКО РОЗРОБКА TELEGRAM-БОТА ДЛЯ АВТОМАТИЗАЦІЇ ОБЛІКУ КРИПТОТРАНЗАКЦІЙ.....	96
Юрій ЛУК'ЯНЧУК, Лев БОЙКО ІННОВАЦІЙНЕ ПРОГРАМНЕ ЗАБЕЗПЕЧЕННЯ ДЛЯ ФОРМУВАННЯ МОРАЛЬНО-ЕТИЧНИХ ЦІННОСТЕЙ ПІДЛІТКІВ У ЦИФРОВУ ЕПОХУ	103
Ігор МАРТИНЮК, Альона ДЕСЯТКО ЗАГРОЗА АТАКИ 51 % ДЛЯ НИЗЬКОХЕШРЕЙТОВИХ ЦИФРОВИХ ВАЛЮТ	110
Ігнат МИРОШНИЧЕНКО ІНТЕГРОВАНІ НАВІГАЦІЙНІ ТА КОМУНІКАЦІЙНІ ПРОТОКОЛИ ДЛЯ АВТОНОМНИХ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ: АРХІТЕКТУРИ, МОДЕЛІ ТА ВИКЛИКИ РЕАЛІЗАЦІЇ В ДИНАМІЧНИХ СЕРЕДОВИЩАХ	116

Костянтин МІНКОВ ВИКОРИСТАННЯ ACTIVE LEARNING ЯК ІНСТРУМЕНТУ ДЛЯ УТОЧНЕННЯ МІТОК У ЗАДАЧАХ, ЗАСНОВАНИХ НА ПРИНЦИПАХ СЛАБКОГО НАВЧАННЯ	123
Артем МОІСЕЄНКО, Юлія КУЗНЕЦОВА АНАЛІЗ ПРОБЛЕМНИХ ПИТАНЬ ВПРОВАДЖЕННЯ СУЧАСНИХ ХМАРНИХ ТЕХНОЛОГІЙ У СИСТЕМИ УПРАВЛІННЯ НАВЧАННЯМ	130
Борис ПАНАСЮК, Наталя БАБЮК КОНТРАКТНО-ОРІЄНТОВАНЕ МОДЕЛЮВАННЯ МІКРОСЕРВІСНОЇ АРХІТЕКТУРИ ВИСОКОНАВАНТАЖЕНИХ СИСТЕМ У UML-ПОДІБНОМУ СЕРЕДОВИЩІ	136
Олексій ПІСКАРЬОВ, Данило АБРАМОВИЧ, Владіслав ПЛУГІН ЗАСТОСУВАННЯ НЕЙРОМЕРЕЖ ТА ГІБРИДНИХ СХОВИЩ ПРИ СТВОРЕННІ КОМП'ЮТЕРНИХ СИСТЕМ	141
Денис ПОКИДЬКО, Олена СКЛЯРЕНКО АЛГОРИТМИ ВИЯВЛЕННЯ ТА ЗАПОБІГАННЯ МАНІПУЛЯЦІЯМ У ГЕЙМІФІКОВАНИХ ОСВІТНІХ СЕРЕДОВИЩАХ	150
Олександр ПОПОВ, Валерія КОВАЧ, Євген ПИЛИПЧУК, Євгенія КОЧЕЛАБ, Володимир АРТЕМЧУК, Теодозія ЯЦИШИН, Володимир КУЦЕНКО ІННОВАЦІЙНІ ПІДХОДИ ДО СТВОРЕННЯ ГІБРИДНИХ МАТЕРІАЛІВ НА ОСНОВІ БІОПОЛІМЕРІВ ДЛЯ ЗАХИСТУ ВІД ІОНІЗУЮЧОГО ВИПРОМІНЮВАННЯ	157
Олексій РИХАЛЬСЬКИЙ АРХІТЕКТУРА ІНТЕЛЕКТУАЛЬНОЇ СИСТЕМИ ПІДТРИМКИ РІШЕНЬ НА ОСНОВІ СЕМАНТИЧНОГО УЗГОДЖЕННЯ МОДЕЛЕЙ	166
Inna ROZLOMII, Serhii NAUMENKO, Volodymyr SYMONYUK, Vitalii PTASHENCHUK, Vitalii ZAZHOMA SIMPLIFIED CRYPTOGRAPHY FOR SECURE VIBRATION PARAMETERS IN POST-PROCESSING OF 3D-PRINTED PARTS	175
Oleh SYPIAHIN, Dmytro POLTAVSKYI, Roman MARTYENKO EXPLORING THE ADVANTAGES AND LIMITATIONS OF THE FEATURE-SLICED DESIGN ARCHITECTURAL METHODOLOGY IN COMMERCIAL FRONTEND PROJECTS	183
Khrystyna TERLETSKA ENHANCING REAL-TIME DATA REPLICATION EFFICIENCY THROUGH OPTIMIZATION OF CHANGE DATA CAPTURE METHODS	188
Данило ТРОЩІЙ, Юлія КУЗНЕЦОВА ОГЛЯД ТА АНАЛІЗ АСПЕКТІВ ПОВТОРЮВАНOSTІ РУТИННИХ ЗАВДАНЬ У ПРОЦЕСІ РОЗРОБЛЕННЯ МОБІЛЬНИХ ДОДАТКІВ	195
Tamara FRANCHUK, Dmytro TYSHCHENKO, Alona DESIATKO, Dmytro HNATCHENKO ASSESSMENT OF THE USE OF CLOUD TECHNOLOGIES IN ACCOUNTING AS AN INNOVATIVE TOOL	200
Олексій ХОНІН, Олена СКЛЯРЕНКО АЛГОРИТМИ ОПТИМАЛЬНОЇ МАРШРУТИЗАЦІЇ ДЛЯ СЛУЖБ ЕКСТРЕНИХ РЕАГУВАНЬ	206
Maksym CHEREMNOV TRANSFORMATION OF INTERNATIONAL CYBERSECURITY STANDARDS AS A RESPONSE TO TECHNOLOGICAL AND GEOPOLITICAL CHALLENGES	211
Sofia SHVETS INTEGRATION OF LARGE LANGUAGE MODELS INTO CHATBOT-BASED CUSTOMER CALL PROCESSING SYSTEMS	216

CONTENTS

Oleksandr ARTAMONOV, Mykhailo LUCHKEVYCH FORMALISATION OF REQUIREMENTS FOR INTELLIGENT DECISION SUPPORT SYSTEMS IN CLOUD RESOURCE OPTIMISATION TASKS	8
Mykola BILYI, Ievgen KRYLOV MODIFICATION OF THE RECIPROCAL RANK FUSION METHOD TO IMPROVE HYBRID SEARCH RESULTS IN INFORMATION SYSTEMS WITH VECTOR DATABASES	14
Oleh HARASYMCHUK, Mykhailo LUCHKEVYCH DESIGN SCIENCE RESEARCH APPROACHES TO BUILDING AN INTELLIGENT CONTRACT TESTING SYSTEM	20
Mykhailo HNATYSHYN, Oleksii NEDASHKIVSKIY A FRAMEWORK FOR EXPLAINABLE AI (XAI) IN MACHINE LEARNING-BASED FAKE NEWS DETECTION SYSTEMS: ENHANCING TRANSPARENCY, TRUST, AND USER AGENCY	24
Galina GRYGORCHUK, Lyubomir GRYGORCHUK, Viktoriia BANDURA AUTOMATION OF MOISTURE CONTROL DURING MASSECUITE DRYING	30
Oleh ZHVALIK ADVANCED INTEGRATION OF IT IN THE MANAGEMENT OF SOLAR ENERGY INDUSTRY	39
Olexander ZADEREYKO, Svitlana YELENYCH, Nataliia LOGINOVA, Olena TROFYMENKO, Volodymyr HURA ANALYSIS OF THE RELEVANCE OF SEARCH RESULTS IN SCIENTIFIC METRICS DATABASES	47
Volodymyr KARABCHUK METHODS FOR DETECTING AND COMBATING DISINFORMATION USING ARTIFICIAL INTELLIGENCE TECHNOLOGIES	56
Yurii KISH, Ihor LIAKH EVOLUTION OF RISK CONCEPTS AND THEIR RELATIONSHIP TO IT PRODUCT QUALITY ASSURANCE	62
Yevhenii KLYMENKO, Olena HLAZUNOVA ARCHITECTURE OF EDUCATIONAL ANALYTICS INFORMATION TECHNOLOGY USING DATA MINING METHODS	69
Andrii KOBLYIUK INCREASING THE EFFICIENCY OF CLOUD RESOURCE MANAGEMENT WITH THE HELP OF DISTRIBUTED MONITORING	76
Artem KOLOMYTSEV, Yuliia KUZNETSOVA THE ROLE OF MLOPS IN ADVANCING GREEN ENERGY: AN EVALUATION OF TECHNOLOGIES AND PRACTICES	83
Yevhen LANSKYKH, Oleksii NAUMOV MODERN METHODS OF DIGITAL IMAGE WATERMARKING: CLASSIFICATION, ANALYSIS AND DEVELOPMENT TRENDS	89
Yelyzaveta LYTIUHA, Valerii YATSENKO USE OF CHATBOTS IN THE INTERNAL FINANCIAL CONTROL OF A CRYPTO EXCHANGE	96
Iurii LUKIANCHUK, Lev BOIKO INNOVATIVE SOFTWARE FOR DEVELOPING MORAL AND ETHICAL VALUES AMONG TEENAGERS IN THE DIGITAL AGE	103
Ihor MARTYNIUK, Alona DESIATKO THREAT OF 51 % ATTACK FOR LOW-HASHRATE CRYPTOCURRENCIES	110
Ihnat MYROSHNYCHENKO INTEGRATED NAVIGATION AND COMMUNICATION PROTOCOLS FOR AUTONOMOUS INTELLIGENT SYSTEMS: ARCHITECTURES, MODELS AND IMPLEMENTATION CHALLENGES IN DYNAMIC ENVIRONMENTS	116

Kostiantyn MINKOV THE USE OF ACTIVE LEARNING FOR REFINING LABELS IN WEAK SUPERVISION BASED TASKS	123
Artem MOISEIENKO, Yuliia KUZNETSOVA ANALYSIS OF IMPLEMENTATION CHALLENGES OF MODERN CLOUD TECHNOLOGIES IN LEARNING MANAGEMENT SYSTEMS	130
Borys PANASIUK, Natalia BABIUK CONTRACT-ORIENTED MODELLING OF MICROSERVICE ARCHITECTURE FOR HIGH-LOAD SYSTEMS IN A UML-LIKE ENVIRONMENT	136
Oleksiy PISKAROV, Danilo ABRAMOVICH, Vladislav PLUHIN APPLICATION OF NEURAL NETWORKS AND HYBRID STORAGE IN THE CREATION OF COMPUTER SYSTEMS	141
Denys POKYDKO, Olena SKLIARENKO ALGORITHMS FOR DETECTION AND PREVENTION OF MANIPULATIONS IN GAMIFIED EDUCATIONAL ENVIRONMENTS	150
Oleksandr POPOV, Valeriia KOVACH, Ievhen PYLYPCHUK, Yevheniia KOCHELAB, Volodymyr ARTEMCHUK, Teodoziia YATSYSHYN, Volodymyr KUTSENKO INNOVATIVE APPROACHES TO CREATING HYBRID MATERIALS BASED ON BIOPOLYMERS FOR PROTECTION AGAINST IONISING RADIATION	157
Oleksiy RYKHALSKYY ARCHITECTURE OF AN INTELLIGENT DECISION SUPPORT SYSTEM BASED ON SEMANTIC MODEL MATCHING	166
Inna ROZLOMII, Serhii NAUMENKO, Volodymyr SYMONYUK, Vitalii PTASHENCHUK, Vitalii ZAZHOMA SIMPLIFIED CRYPTOGRAPHY FOR SECURE VIBRATION PARAMETERS IN POST-PROCESSING OF 3D-PRINTED PARTS	175
Oleh SYPIAHIN, Dmytro POLTAVSKYI, Roman MARTYENKO EXPLORING THE ADVANTAGES AND LIMITATIONS OF THE FEATURE-SLICED DESIGN ARCHITECTURAL METHODOLOGY IN COMMERCIAL FRONTEND PROJECTS	183
Khrystyna TERLETSKA ENHANCING REAL-TIME DATA REPLICATION EFFICIENCY THROUGH OPTIMIZATION OF CHANGE DATA CAPTURE METHODS	188
Danylo TROSHCHYI, Yuliya KUZNETSOVA REVIEW AND ANALYSIS OF REPETITIVE ROUTINE TASKS IN THE MOBILE APPLICATION DEVELOPMENT PROCESS	195
Tamara FRANCHUK, Dmytro TYSHCHENKO, Alona DESIATKO, Dmytro HNATCHENKO ASSESSMENT OF THE USE OF CLOUD TECHNOLOGIES IN ACCOUNTING AS AN INNOVATIVE TOOL	200
Oleksii KHONIN, Olena SKLIARENKO OPTIMAL ROUTING ALGORITHMS FOR EMERGENCY RESPONSE SERVICES	206
Maksym CHEREMNOV TRANSFORMATION OF INTERNATIONAL CYBERSECURITY STANDARDS AS A RESPONSE TO TECHNOLOGICAL AND GEOPOLITICAL CHALLENGES	211
Sofiia SHVETS INTEGRATION OF LARGE LANGUAGE MODELS INTO CHATBOT-BASED CUSTOMER CALL PROCESSING SYSTEMS	216

УДК 004.415.3

DOI <https://doi.org/10.32689/maup.it.2025.2.1>

Олександр АРТАМОНОВ

аспірант кафедри інформаційних систем та мереж,
Національний університет «Львівська політехніка»,
oleksandr.o.artamonov@lpnu.ua
ORCID: 0009-0000-8295-733X

Михайло ЛУЧКЕВИЧ

кандидат фізико-математичних наук, доцент,
доцент кафедри інформаційних систем та мереж,
Національний університет «Львівська політехніка»,
mykhailo.m.luchkevych@lpnu.ua
ORCID: 0000-0002-2196-252X

ФОРМАЛІЗАЦІЯ ВИМОГ ДО ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ ПІДТРИМКИ УХВАЛЕННЯ РІШЕНЬ У ЗАДАЧАХ ОПТИМІЗАЦІЇ ХМАРНИХ РЕСУРСІВ

Анотація. У контексті сучасних хмарних обчислювальних середовищ зростає необхідність впровадження автоматизованих механізмів управління ресурсами з метою забезпечення належного рівня продуктивності та оптимізації витрат. Інтелектуальні системи підтримки прийняття рішень виступають ключовим компонентом таких механізмів, оскільки забезпечують аналіз динамічних метрик і прогнозування змін навантаження для адаптивного масштабування.

Мета. Розробка формальної моделі вимог до інтелектуальних систем підтримки прийняття рішень у задачах оптимізації ресурсів хмарної інфраструктури.

Методологія. У межах дослідження здійснено аналіз існуючих підходів до автоскейлінгу та інтелектуальних рішень для керування хмарними середовищами. Застосовано методи системного аналізу для виокремлення категорій вимог та побудови онтологічної моделі, яка включає такі сутності: *CloudResource*, *MonitoringAgent*, *MLModel*, *ScalingPolicy* та *DecisionRecord*. Експериментальна валідація виконана шляхом реалізації прототипу в середовищі *Kubernetes* із застосуванням *Prometheus* для збору метрик.

Наукова новизна. Наукова новизна роботи полягає у комплексній формалізації вимог до інтелектуальних систем підтримки прийняття рішень з урахуванням як функціональних, так і нефункціональних аспектів. Зокрема, до функціональних вимог віднесено моніторинг різномірних метрик, механізми прогнозування, генерацію рекомендацій щодо масштабування та аудит ухвалених рішень.

Висновки. Отримані результати свідчать, що формалізація вимог є необхідною умовою створення надійних і прозорих інтелектуальних систем автоскейлінгу в хмарних середовищах. Запропонована модель забезпечує ефективний баланс між продуктивністю та витратами на інфраструктуру, сприяючи оптимальному використанню ресурсів без погіршення якості обслуговування. У перспективі доцільно розширити онтологію для мультихмарних і безсерверних сценаріїв, розробити методи оперативного оновлення моделей (*Concept Drift*) та вдосконалити механізми *Explainable AI* для забезпечення більш прозорого обґрунтування рекомендацій.

Ключові слова: Інтелектуальні системи підтримки ухвалення рішень, оптимізація хмарних ресурсів, формалізація вимог, автоскейлінг, онтологічна модель.

Oleksandr ARTAMONOV, Mykhailo LUCHKEVYCH. FORMALISATION OF REQUIREMENTS FOR INTELLIGENT DECISION SUPPORT SYSTEMS IN CLOUD RESOURCE OPTIMISATION TASKS

Abstract. In the context of modern cloud computing environments, there is an increasing demand for automated resource management systems to ensure optimal performance and cost efficiency. Intelligent decision support systems are a vital part of these mechanisms, providing analysis of dynamic metrics and forecasting of load changes for adaptive scaling.

Objective. The objective is to develop a formal model of the requirements for intelligent decision support systems for optimising cloud infrastructure resources.

Methodology. This study analyses existing approaches to autoscaling and intelligent solutions for managing cloud environments. System analysis methods were employed to identify categories of requirements and construct an ontological model comprising the following entities: *CloudResource*, *MonitoringAgent*, *MLModel*, *ScalingPolicy* and *DecisionRecord*. Experimental validation was performed by implementing a prototype in a *Kubernetes* environment and using *Prometheus* to collect metrics.

Scientific novelty. This work is scientifically novel due to its comprehensive formalisation of requirements for intelligent decision support systems, considering both functional and non-functional aspects. The functional requirements include monitoring heterogeneous metrics, forecasting mechanisms, generating scaling recommendations and auditing decisions.

Conclusion. The results obtained show that the formalisation of requirements is a prerequisite for the creation of reliable and transparent intelligent autoscaling systems in cloud environments. The proposed model strikes an effective balance

© О. Артамонов, М. Лучкевич, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

between performance and infrastructure costs, ensuring the optimal use of resources without compromising service quality. Future work should focus on extending the ontology to cover multi-cloud and serverless scenarios, developing methods for rapidly updating models (concept drift) and improving Explainable AI mechanisms to provide more transparent justification for recommendations.

Key words: Intelligent decision support systems, cloud resource optimisation, requirements formalisation, autoscaling, ontological model.

Постановка проблеми. У сучасних інформаційних системах хмарні сервіси дедалі частіше виступають основою для розгортання додатків, зберігання даних та виконання обчислень. Завдяки еластичності та високій доступності, платформи IaaS, PaaS і SaaS [10] забезпечують можливість динамічного масштабування ресурсів відповідно до змін навантаження. Водночас зростання кількості користувачів і ускладнення бізнес-процесів породжує низку завдань щодо оптимального розподілу віртуальних машин, контейнерів або безсерверних функцій з метою мінімізації витрат, збереження належного рівня продуктивності та гарантування необхідних показників якості обслуговування.

Проблема набуває значної ваги як у прикладному, так і в теоретичному аспектах. З одного боку, організації прагнуть зменшити фінансові витрати на утримання хмарної інфраструктури, уникнути надмірного резервування обчислювальних потужностей і скоротити час простою критичних сервісів. З іншого боку, наукова спільнота стикається з рядом невирішених питань: яким чином побудувати модель поведінки хмарної платформи за умов непередбачуваної динаміки навантаження; як інтегрувати алгоритми машинного навчання й штучного інтелекту для адаптивного налаштування конфігурації; яким чином формалізувати вимоги до інтелектуальних систем підтримки ухвалення рішень (ІСПУР) [5], що виконуватимуть функцію «інтелектуального ядра» й забезпечуватимуть аналіз поточного стану системи та генерацію оптимальних рекомендацій щодо резервування й масштабування.

Наявність нечітко окреслених або неповних вимог ускладнює процес розробки ефективних ІСПУР: це призводить до фрагментарного застосування алгоритмів, недостатньої інтеграції з існуючою інфраструктурою й, як наслідок, до зниження загальної ефективності оптимізації хмарних ресурсів. Отже, актуальним завданням є чітке визначення ключових характеристик таких систем, формулювання функціональних та нефункціональних вимог, а також встановлення критеріїв якості, що забезпечать надійну, адаптивну й прозору роботу ІСПУР в умовах динамічних хмарних середовищ.

Аналіз останніх досліджень і публікацій. У останні роки спостерігається значне поглиблення досліджень, присвячених розробці інтелектуальних рішень для управління ресурсами хмарних середовищ. У науковій літературі можна виокремити такі основні напрями досліджень.

Перший напрям стосується методів автоскейлінгу, які базуються на PID-регулюванні та адаптивних алгоритмах зворотного зв'язку [2]. Застосування класичних PID-контролерів дозволяє відносно просто реалізувати механізми автоскейлінгу, однак ефективність таких підходів значною мірою залежить від точного налаштування параметрів контролера [6]. У більшості випадків ця процедура здійснюється вручну або з використанням евристичних методів, що обмежує масштабованість і автоматизацію управління.

Другий напрям пов'язаний із моделями прогнозування завантаження, у яких переважно використовуються методи аналізу часових рядів (зокрема ARIMA [4], Prophet та інші). Ці підходи демонструють обнадійливі результати щодо прогнозування майбутнього навантаження на основі історичних даних, проте їхня точність залежить від стабільності користувацьких патернів та змін зовнішніх умов. Виділення коректних трендів і сезонних складових у даних часто ускладнене неоднорідністю навантаження та наявністю різких, непередбачуваних стрибків у запитах.

Третій напрям охоплює створення інтелектуальних систем підтримки ухвалення рішень із застосуванням методів машинного навчання [1; 7]. У таких системах переважно використовують нейромережеві архітектури (LSTM, GRU, CNN) для прогнозування попиту на ресурси та оптимізації політик розподілу. Окремі дослідники пропонують гібридні рішення, які поєднують евристичні правила з методами глибокого навчання, що дозволяє поєднати точність прогнозування з необхідною швидкістю при ухваленні рішень.

Четвертий напрям присвячений розробці фреймворків для створення інтелектуальних систем підтримки ухвалення рішень у хмарних середовищах [8]. Прикладами таких рішень є OpenStack Telemetry (Ceilometer) та Kubernetes HPA (Horizontal Pod Autoscaler), які надають базові можливості моніторингу та автоскейлінгу. Проте вони не завжди забезпечують достатню гнучкість для реалізації складних сценаріїв ухвалення рішень, тому в сучасних роботах дедалі частіше обговорюють необхідність інтеграції багаторівневого аналізу й методів штучного інтелекту з базовими метриками продуктивності.

П'ятий напрям стосується формалізації вимог до інтелектуальних систем підтримки ухвалення рішень [3]. Деякі публікації пропонують розробку спеціалізованих мов (наприклад, домен-специфічних

мов) для опису політик автоскейлінгу [9], однак питання комплексної формалізації всіх функціональних і нефункціональних вимог ІСПУР досі залишається недостатньо висвітленим.

Аналіз сучасного стану досліджень показує, що більшість робіт зосереджена переважно або на розробці та вдосконаленні алгоритмів прогнозування та адаптивного контролю, або на реалізації ІСПУР у рамках конкретних фреймворків хмарних платформ. Водночас системного узагальнення вимог до таких систем, яке б поєднувало потреби як наукової, так і прикладної спільнот, поки що недостатньо. Саме це обґрунтовує необхідність подальшого поглибленого вивчення та формалізації вимог до інтелектуальних систем підтримки ухвалення рішень у задачах оптимізації хмарних ресурсів.

Метою статті є розробка формальної моделі вимог до інтелектуальних систем підтримки ухвалення рішень у контексті оптимізації хмарних ресурсів. Зокрема, передбачається виконати такі завдання: по-перше, здійснити класифікацію та ієрархізацію функціональних і нефункціональних вимог до ІСПУР для хмарної інфраструктури; по-друге, запропонувати формалізований опис цих вимог із застосуванням онтологічного підходу, що охоплюватиме моделі ресурсів, політик та сценаріїв ухвалення рішень; по-третє, визначити набір базових критеріїв якості (ключових показників ефективності, KPI), які дозволяють оцінювати роботу ІСПУР у режимі реального часу; по-четверте, продемонструвати практичне застосування розробленої моделі вимог шляхом створення прототипу ІСПУР із використанням методів машинного навчання для прогнозування навантаження. Отже, завдання дослідження полягають не лише в теоретичній побудові моделі вимог, а й у перевірці її коректності через прототипну реалізацію та експериментальну валідацію в середовищі Kubernetes з використанням інструментів телеметрії Prometheus.

Виклад основного матеріалу. Інтелектуальна система підтримки ухвалення рішень для оптимізації хмарних ресурсів реалізується у вигляді трирівневої архітектури. Перший рівень – «Шар збору та агрегації метрик» (Monitoring Layer) – забезпечує збір даних за допомогою агентів, які моніторять показники ефективності: завантаження CPU, споживання оперативної пам'яті, пропускну здатність мережі, використання дискового простору, а також метрики рівня додатків (затримка відповіді, довжина черги запитів тощо). Отримана інформація надходить до централізованого сховища часових рядів (Time-Series Database), де зберігається для подальшої обробки.

Другий рівень – «Шар аналітики й прогнозування» (Analytics & Prediction Layer) – відповідає за підготовку даних та побудову прогностичних моделей. Даний шар включає модулі попередньої обробки, які здійснюють нормалізацію, фільтрацію та агрегацію метрик відповідно до заданих часових інтервалів. На очищених даних навчання проходять алгоритми машинного навчання (наприклад, LSTM або Prophet), що передбачають майбутнє навантаження у визначені горизонти прогнозу.

Третій рівень – «Шар ухвалення рішень і застосування політик» (Decision & Control Layer) – генерує рекомендації щодо масштабування та конфігурацій ресурсів на основі прогнозів попереднього рівня. Зокрема, на цьому рівні визначається доцільність горизонтального (збільшення чи зменшення кількості екземплярів) або вертикального (зміни параметрів CPU/memory) масштабування, політики розміщення контейнерів і віртуальних машин, а також адаптивні параметри балансувальників навантаження. Сформовані рішення передаються до компонента оркестрування (наприклад, Kubernetes Controller), який безпосередньо виконує зміни конфігурації інфраструктури.

Вимоги до ІСПУР поділяються на дві основні групи (рис. 1): функціональні та нефункціональні.

Система має забезпечувати збір метрик із різномірних джерел, зокрема з рівня гіпервізора, гостьової операційної системи, контейнерів та додатків, із подальшим агрегуванням даних на різних часових інтервалах (секундному, хвилинному, годинному). Отримані показники зберігаються в історичній базі даних для проведення аналітичної обробки та навчання прогностичних моделей. Інтеграція з фреймворками машинного навчання (наприклад, TensorFlow, PyTorch, scikit-learn) дозволяє здійснювати навчання моделей на історичних даних з урахуванням сезонних та трендових змін, а також адаптивно оновлювати їх у режимі реального часу або потокової обробки нових метрик.

Механізми масштабування реалізуються у двох вимірах: горизонтальному (зміна кількості реплік) та вертикальному (регулювання параметрів CPU і пам'яті). При формуванні рекомендацій щодо масштабування необхідно враховувати існуючі квоти на ресурси – обмеження облікових записів та платіжних планів. Крім того, має бути визначено політики динамічного ребалансування навантаження між обчислювальними ресурсами та забезпечено можливість гібридного режиму ухвалення рішень із участю «людини в циклі» (human-in-the-loop).

Для аудиту та трасування всіх змін інфраструктури передбачено логування відповідних рекомендацій і дій, що застосовуються до цільових об'єктів. Має бути забезпечена можливість відстеження мотивів ухвалення кожного рішення на основі принципів Explainable AI, а також збереження стану системи до й після виконання рекомендацій щодо масштабування.

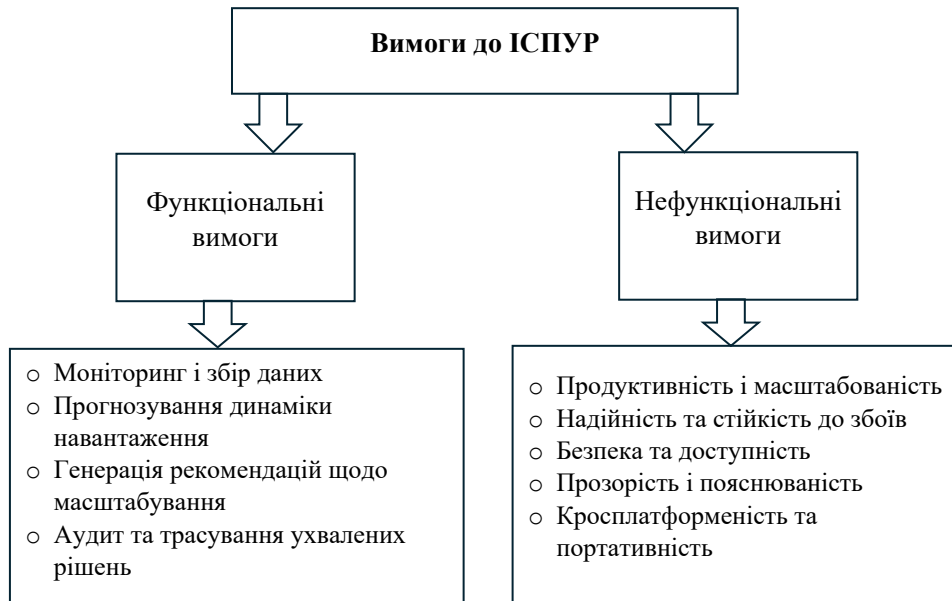


Рис. 1. Вимоги до інтелектуальних систем підтримки ухвалення рішень

Затримка обробки нових метрик і генерації рекомендацій не повинна перевищувати встановленого рівня сервісу (наприклад, одну хвилину з моменту отримання останньої метрики). Підсистеми моніторингу й аналітики повинні бути відмовостійкими та здатними автоматично масштабуватися у відповідь на збільшення навантаження або кількості клієнтських вузлів.

Для забезпечення надійності передбачено збереження консистентності даних у разі втрати зв'язку з окремими вузлами та впровадження механізмів резервування критично важливих компонентів (зокрема кластерного зберігання бази даних і резервних сервісів машинного навчання). Безпеку системи гарантують аутентифікація та авторизація користувачів для доступу до інтерфейсів моніторингу й управління політиками, шифрування каналів зв'язку між агентами моніторингу та сервером збору даних за допомогою TLS, а також ізоляція критичних компонентів з метою запобігання несанкціонованому доступу або виконанню шкідливого коду.

Прозорість ухвалених рішень забезпечується наданням доступу до історії рекомендацій із детальним описом факторів, що впливали на їх формування, а також інтеграцією з інструментами візуалізації (Grafana, Kibana) для аналізу трендів метрик та взаємозв'язків між показниками ресурсів і згенерованими рекомендаціями. Крім того, рішення повинні підтримувати мультихмарні та гібридні сценарії розгортання (AWS, Azure, Google Cloud, on-premises) і забезпечувати уніфікований інтерфейс взаємодії з різними хмарними провайдерами через API-адаптери, що спрощує інтеграцію та гарантує портативність системи.

Для уніфікованого опису вимог до ІСПУР використовується онтологічна модель, яка включає основні сутності: хмарний ресурс, агент моніторингу, модель машинного навчання, політика масштабування та запис ухваленого рішення. Хмарний ресурс характеризується типом (віртуальна машина, контейнер або функція) та відповідними метриками (наприклад, CPUUtilization, MemoryUsage, NetworkIO). Агент моніторингу містить налаштування збору метрик, інтервал опитування та формат передачі даних. Модель машинного навчання описується архітектурою (LSTM, RandomForest, Prophet), набором гіперпараметрів і періодом оновлення. Політика масштабування визначає умови триггеру (наприклад, CPU > 70 % протягом п'яти хвилин), тип масштабування (горизонтальне або вертикальне) та межі кількості ресурсів (мінімальні та максимальні значення). Запис ухваленого рішення зберігає позначку часу (timestamp), набір вхідних метрик, прогнозоване значення навантаження, рекомендовану дію та статус виконання.

Кожна з цих сутностей має власні атрибути та зв'язки. Так, хмарний ресурс містить атрибути resourceID, resourceType і currentMetrics і пов'язаний із агентом моніторингу через зв'язок monitoredBy. Зі свого боку, політика масштабування асоціюється з моделлю машинного навчання за допомогою зв'язку usesModel і визначає перелік цільових ресурсів через зв'язок appliesTo → CloudResource. Така структура моделі дозволяє формалізувати не лише структурні, а й функціональні аспекти доменної логіки: у разі виявлення аномалії навантаження за допомогою відповідної моделі-аномалії (ANOMALY_DETECTOR) система може автоматично активувати резервну політику, залучаючи альтернативні набори ресурсів.

Для об'єктивної оцінки діяльності інтелектуальної системи підтримки ухвалення рішень пропонується застосовувати низку ключових показників ефективності. Першим таким показником є точність прогнозу (Accuracy of Forecast), яка обчислюється на основі середньоквадратичної помилки (RMSE) або середньої абсолютної відносної помилки (MAPE) між передбаченими та фактичними значеннями навантаження. Цей індикатор дозволяє оцінити, наскільки коректно модель машинного навчання формує прогнози й у такий спосіб впливає на якість ухвалених рішень.

Другий показник визначає затримку ухвалення рішення (Decision Latency), що вимірюється як інтервал часу між моментом отримання останнього пакета метрик і часом генерації рекомендації щодо масштабування. Значення цього показника безпосередньо впливає на здатність системи реагувати на зміну навантаження в реальному часі.

Третім індикатором виступає покращення використання ресурсів (Resource Utilization Improvement). Його оцінка здійснюється шляхом порівняння середнього коефіцієнта завантаження CPU та оперативної пам'яті до і після впровадження інтелектуальної системи. Даний показник демонструє ефективність балансування ресурсів та зменшення надмірного резервування.

Четвертим важливим критерієм є зниження витрат (Cost Reduction), яке визначається як процентне співвідношення зменшення витрат на хмарну інфраструктуру у порівнянні з режимом статичного масштабування протягом контрольного періоду. Висока ефективність ухвалених рекомендацій має сприяти зменшенню загальних витрат на обчислювальні ресурси.

Останнім зазначається відповідність рівням обслуговування (Service Level Compliance), що обчислюється як частка випадків перевищення допустимого значення SLA, наприклад, затримки відповіді сервісу понад 200 мс. Моніторинг цього показника дозволяє контролювати якість кінцевого користувацького досвіду й своєчасно виявляти потенційні проблеми.

Постійний супровід зазначених показників дає змогу своєчасно коригувати конфігурацію прогностичних моделей та політик масштабування, а у разі критичних відхилень – оперативно інформувати адміністраторів про необхідність втручання.

Для демонстрації практичної доцільності запропонованої формалізованої моделі вимог розроблено прототип інтелектуальної системи підтримки ухвалення рішень із використанням платформи контейнерної оркестрації Kubernetes, механізму збору та збереження часових рядів метрик Prometheus, засобів візуалізації Grafana і сервера машинного навчання з попередньо навченими моделями LSTM і Prophet для прогнозування навантаження.

Висновки. У статті пропонується формалізація вимог до інтелектуальних систем підтримки ухвалення рішень у задачах оптимізації хмарних ресурсів. Виконано класифікацію функціональних і нефункціональних вимог, визначено ключові критерії якості та розроблено онтологічну модель для формального опису основних сутностей ІСПУР.

Водночас слід зазначити певні обмеження: обрана онтологія потребує адаптації для мультихмарних і гібридних середовищ, оскільки не всі API провайдерів уніфіковано подають метрики. Крім того, експериментальні результати вказують на необхідність розробки механізмів для онлайн-адаптації моделей з урахуванням концептуального зсуву, оскільки змінні патерни навантаження можуть зменшувати точність прогнозування. Також актуальною залишається проблема забезпечення прозорості рекомендацій, що сприятиме довірі з боку адміністраторів, і питання інтеграції безсерверних середовищ (Serverless), де метрики функціонування відрізняються від класичних контейнерних сервісів.

Отже, формалізація вимог до інтелектуальних систем підтримки ухвалення рішень є важливим кроком для створення надійних і прозорих рішень у сфері оптимізації хмарних ресурсів. Подальші зусилля доцільно спрямувати на вдосконалення адаптивного навчання моделей і розширення онтологічної бази з урахуванням специфіки мультихмарних та безсерверних інфраструктур.

Список використаних джерел:

1. Ali R. et al. Intelligent decision support systems – An analysis of machine learning and multicriteria decision-making methods. *Applied Sciences*. 2023. V. 13. №. 22. P. 12426. <https://doi.org/10.3390/app132212426>
2. Burroughs S. et al. Towards autoscaling with guarantees on kubernetes clusters. 2021 IEEE International Conference on Autonomic Computing and Self-Organizing Systems Companion (ACSOS-C). IEEE, 2021. P. 295–296. doi: 10.1109/ACSOS-C52956.2021.00073.
3. Hejase M., Katis A., Mavridou A. Design, formalization, and verification of decision making for intelligent systems. *AIAA SCITECH 2024 Forum*. 2024. P. 2409. <https://doi.org/10.2514/6.2024-2409>
4. Joshi N. S. et al. ARIMA-PID: container auto scaling based on predictive analysis and control theory. *Multimedia Tools and Applications*. 2024. V. 83. №. 9. P. 26369–26386. <https://doi.org/10.1007/s11042-023-16587-0>
5. Onwujekwe G., Weistroffer H. R. Intelligent Decision Support Systems: An Analysis of the Literature and a Framework for Development. *Information Systems Frontiers*. 2025. P. 1–32. <https://doi.org/10.1007/s10796-024-10571-1>

6. Sabuhi M., Mahmoudi N., Khazaei H. Optimizing the performance of containerized cloud software systems using adaptive PID controllers. *ACM Transactions on Autonomous and Adaptive Systems (TAAS)*. 2021. V. 15. №. 3. P. 1–27. <https://doi.org/10.1145/3465630>
7. Saeed A. Machine Learning Models for Intelligent Decision Support Systems. *Journal of AI Range*. 2024. V. 1. №. 1. P. 54–66. URL: <https://www.researchcorridor.org/index.php/jair/article/view/268>
8. Simeone A., Zeng Y., Caggiano A. Intelligent decision-making support system for manufacturing solution recommendation in a cloud framework. *The International Journal of Advanced Manufacturing Technology*. 2021. V. 112. № 3. P. 1035–1050. <https://doi.org/10.1007/s00170-020-06389-1>
9. Tari M. et al. Auto-scaling mechanisms in serverless computing: A comprehensive review. *Computer Science Review*. 2024. V. 53. P. 100650. <https://doi.org/10.1016/j.cosrev.2024.100650>
10. Younis R. et al. A Comprehensive Analysis of Cloud Service Models: IaaS, PaaS, and SaaS in the Context of Emerging Technologies and Trend. 2024 International Conference on Electrical, Communication and Computer Engineering (ICECCE). IEEE, 2024. P. 1–6. doi: 10.1109/ICECCE63537.2024.10823401

Дата надходження статті: 09.06.2025

Дата прийняття статті: 15.06.2025

Опубліковано: 23.09.2025

УДК 004.85:004.912:519.8

DOI <https://doi.org/10.32689/maup.it.2025.2.2>

Микола БІЛИЙ

аспірант кафедри інформаційних систем та технологій,

Національний технічний університет України

«Київський політехнічний інститут імені Ігоря Сікорського», mykola.o.b@gmail.com

ORCID: 0009-0004-0639-3658

Євген КРИЛОВ

кандидат технічних наук, доцент кафедри інформаційних систем та технологій,

Національний технічний університет України

«Київський політехнічний інститут імені Ігоря Сікорського», ekrylov1964@gmail.com

ORCID: 0000-0003-4313-938X

**МОДИФІКАЦІЯ МЕТОДУ RECIPROCAL RANK FUSION ДЛЯ ПОЛІПШЕННЯ РЕЗУЛЬТАТІВ
ГІБРИДНОГО ПОШУКУ У ІНФОРМАЦІЙНИХ СИСТЕМАХ З ВЕКТОРНИМИ БАЗАМИ ДАНИХ**

Анотація. Метою даної роботи є удосконалення методу Reciprocal Rank Fusion (RRF) для підвищення точності об'єднання результатів гібридного пошуку в інформаційних системах, що використовують векторні бази даних. Гібридний пошук передбачає інтеграцію результатів, отриманих за допомогою різних стратегій пошуку – лексичних, семантичних, візуальних тощо. Однак традиційна формула RRF не враховує ступінь релевантності документів, що може призводити до ситуацій, коли низько релевантні результати однієї пошукової стратегії впливають на високо релевантні результати іншої. У результаті відбувається зміщення бажаного порядку документів у фінальному списку результатів.

Методологія дослідження ґрунтується на аналізі оцінок релевантності, отриманих у процесі лексичного та семантичного пошуку. Запропоновано підхід до класифікації результатів за групами релевантності на основі поетапної деградації початкового запиту. Далі запропоновано модифіковану формулу RRF, в якій ранги документів коригуються з урахуванням групи релевантності через експоненціальну функцію. Це дозволяє зменшити вплив результатів, що мають низьку релевантність, на документи з високою відповідністю. Для експериментальної перевірки було використано датасет MS MARCO, що містить реальні пошукові запити й ручні оцінки релевантності. Порівняння класичного та модифікованого RRF здійснювалося за метрикою MRR@10.

Наукова новизна роботи полягає у використанні динамічного, контекстно-залежного підходу до формування груп релевантності без необхідності в повторних зверненнях до бази даних чи індексації даних. Запропоноване рішення обмежується лише текстом запиту, що забезпечує його ефективність і не вимагає відносно великих обчислювальних витрат. На відміну від існуючих модифікацій RRF, запропонований підхід дозволяє гнучко адаптувати ваги ранжування на основі семантичного та лексичного динамічного аналізу.

Висновки дослідження підтверджують, що запропонована модифікація методу RRF покращує позицію релевантного документа у результатах гібридного пошуку. Модифікований метод демонструє вищу метрику MRR@10 (0.1880 проти 0.1718 у класичному RRF), що свідчить про зменшення негативного впливу нерелевантних результатів.

Ключові слова: семантичний пошук, лексичний пошук, гібридний пошук, комбінований пошук, векторні бази даних, інформаційні системи.

**Mykola BILYI, Ievgen KRYLOV. MODIFICATION OF THE RECIPROCAL RANK FUSION METHOD TO IMPROVE
HYBRID SEARCH RESULTS IN INFORMATION SYSTEMS WITH VECTOR DATABASES**

Abstract. The aim of this study is to improve the Reciprocal Rank Fusion (RRF) method to enhance the accuracy of result merging in hybrid search systems that utilize vector databases. Hybrid search involves the integration of results obtained through various strategies, including lexical, semantic, and visual search. However, the traditional RRF formula does not take into account the degree of document relevance, which may lead to situations where low-relevance results from one retrieval strategy affect the high-relevance results from another. As a consequence, the desired ranking of documents in the final result list may be distorted.

The methodology of this research is based on analyzing relevance scores obtained through lexical and semantic search. A relevance-based classification approach is proposed, where search results are grouped according to their level of relevance using a stepwise degradation of the original query. A modified RRF formula is introduced, in which document ranks are adjusted based on relevance group membership using an exponential function. This approach helps reduce the influence of low-relevance results on highly relevant ones. For experimental validation, the MS MARCO dataset was used, which contains real user queries and manually annotated relevance judgments. The effectiveness of the original and modified RRF methods was compared using the MRR@10 metric.

The scientific novelty of this work lies in the use of a dynamic, context-sensitive approach to forming relevance groups without requiring repeated access to the database or index. The proposed solution operates solely on the query text, ensuring

© М. Білий, Є. Крилов, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

computational efficiency and minimizing resource consumption. Unlike existing RRF modifications, the proposed method allows for flexible adjustment of ranking weights based on lexical and semantic analysis.

The findings of the study confirm that the proposed modification of the RRF method improves the ranking position of relevant documents in hybrid search results. The modified method achieves a higher MRR@10 score (0.1880 compared to 0.1718 for the classical RRF), indicating a reduction in the negative impact of irrelevant results.

Key words: semantic search, lexical search, hybrid search, combined search, vector databases, information systems.

Постановка проблеми. Векторні бази даних є сучасним різновидом інформаційних систем, що забезпечують зберігання, індексацію та пошук даних у вигляді числових векторів, сформованих за допомогою моделей машинного навчання. Кожен вектор відображає об'єкт у вигляді набору числових ознак, які кількісно характеризують його властивості. Наприклад, зображення автомобіля може бути перетворене у вектор виду [0.85, 0.6, 0.3], де компоненти описують, відповідно, яскравість кузова, округлість силуету та інтенсивність світлових елементів. Такий підхід використовується для зберігання та пошуку об'єктів різного типу, представлених у вигляді векторів ознак [6].

У порівнянні з традиційними реляційними базами даних, у яких пошук здійснюється через точну відповідність значень, у векторних базах даних реалізується пошук за принципом схожості між векторами. Для оцінювання подібності між векторними поданнями використовуються метричні та квазіметричні функції, зокрема евклідова відстань та косинусна схожість [1]. Наприклад, якщо два вектори зображень різних моделей взуття розташовані досить близько у векторному просторі і розрахована відстань між цими векторами відносно невелика, система інтерпретує їх як подібні об'єкти.

У сучасних інформаційних системах, зокрема в електронній комерції, векторні бази даних застосовуються для зберігання мультимодальних даних – текстових описів товарів, їхніх зображень, аудіо або відео. Для прикладу, опис товару включає як текстову інформацію (назву, характеристики, опис), так і візуальне представлення. Усі ці елементи конвертуються у відповідні вектори за допомогою моделей глибокого навчання (наприклад, CLIP або BERT), що дозволяє здійснювати гібридний (комбінований) пошук [10].

Гібридний пошук дозволяє ефективніше обробляти запити користувачів поєднуючи кілька стратегій пошуку, наприклад лексичний, семантичний та пошук за зображенням. Лексичний пошук орієнтований на фільтрацію результатів за ключовими словами, які містяться в текстовому описі об'єкта. Семантичний підхід дозволяє знаходити об'єкти, подібні за змістом, навіть якщо у формулюванні запиту використовуються інші слова, ніж у базі даних. Наприклад, запит «взуття для гірських походів» може повернути результати, де згадуються «трекінгові черевики» або «трекінгові кросівки». Пошук за зображенням забезпечує виявлення візуально схожих об'єктів, навіть за відсутності повного або часткового текстового опису. Наприклад, пошук «футболка команди Арсенал» дозволить знайти всі схожі футболки з емблемою команди, навіть якщо в описах не згадуються конкретні бренди чи клуби чи навіть слова із пошукового запиту.

У задачах гібридного пошуку актуальною проблемою є завдання узгодженого об'єднання результатів, отриманих за допомогою різних стратегій пошуку, з метою підвищення якості остаточного списку відповідей. Одним із широко використовуваних методів такого об'єднання є Reciprocal Rank Fusion (RRF) [5]. Його математичне визначення задається формулою:

$$RRF(d) = \sum_{s \in S} \frac{1}{k + rank_s(d)} \quad (1)$$

де $RRF(d)$ – підсумкова оцінка документа d ; S – множина пошукових стратегій; $rank_s(d)$ – позиція документа d у списку результатів пошуку s ; k – постійна для згладжування (зазвичай $k = 60$).

Суть методу полягає в присвоєнні кожному результату оцінки, яка залежить від зворотного значення його позиції у відповідному списку. Після цього всі отримані оцінки підсумовуються, що дозволяє виводити на початок об'єднаного списку ті результати, які посідають високі місця у кількох пошукових видачах, навіть якщо не очолюють жодну з них.

Основна складність цього підходу полягає в тому, що системи ранжування завжди повертають певну кількість результатів, незалежно від того, чи мають вони релевантний зв'язок із запитом. Це особливо характерно для пошуку за векторною схожістю, де система повертає найближчі об'єкти, навіть якщо вони не мають суттєвого змістового збігу з запитом.

Наприклад, у списку результатів пошуку за зображенням, перше зображення може бути дійсно релевантним і посідає перше місце, тоді як наступні (друге, третє тощо), з точки зору людини, не мають змістовного зв'язку з інформаційним запитом. Незважаючи на це, система ранжування зображень все одно присвоює їм високі позиції (ранги 2, 3 тощо). У разі використання методу RRF з параметром $k = 60$, перше зображення отримає оцінку 0,01639, а друге – 0,01611. Таким чином, навіть нерелевантне

зображення, що випадково опинилося на другій позиції, отримує оцінку, яка лише не значно поступається оцінці справді релевантного документа. Це знижує точність об'єднання результатів, оскільки високі оцінки нерелевантних об'єктів можуть зменшити вплив більш точних результатів з інших стратегій пошуку під час інтеграції списків.

Проілюструємо принцип об'єднання результатів семантичного та лексичного пошуку на основі методу RRF, використовуючи приклад, наведений у (табл. 1). З метою спрощення аналізу припустимо, що результати лексичного та семантичного пошуку нормалізовані в межах від 0 до 1.

Таблиця 1

Вплив низько релевантних оцінок на високо релевантні оцінки

Номер документа	Оцінка семантичн. пошуку	Ранг згідно семантичн. пошуку	Оцінка лексичн. пошуку	Ранг згідно лексичн. пошуку	RRF	Ранг згідно RRF
1	0.95	1	0.1	5	0.031778	1
2	0.94	2	0.2	3	0.031754	3
3	0.93	3	0.25	3	0.031746	5
4	0.92	4	0.29	2	0.031754	4
5	0.91	5	0.3	1	0.031778	2

Високе значення оцінки з певного типу пошуку свідчить про сильну релевантність документа щодо запиту, навіть якщо інші типи пошуку присвоюють цьому ж документу низькі оцінки. Значення семантичного пошуку в цьому прикладі вказують на те, що знайдені документи добре відповідають змісту користувацького запиту. Водночас низькі оцінки, отримані з лексичного пошуку, вказують на слабку текстову відповідність за формальними метриками, такими як BM25 [11] або TF-IDF [7]. На основі обчислення RRF для кожного документа було отримано наступний порядок у зведеному результаті пошуку: документ 1, документ 5, документ 2, документ 4, документ 3. Як видно, отримане впорядкування не відповідає первісному порядку за семантичним пошуком, що пояснюється впливом низьких оцінок з боку лексичного пошуку на високі семантичні оцінки. Зокрема, документ 3, незважаючи на його високу релевантність у семантичному пошуку, був переміщений на останню позицію через недостатню лексичну відповідність. Таким чином, слабкі оцінки лексичного аналізу можуть вносити додаткові похибки в інтегровану оцінку.

Аналіз останніх досліджень і публікацій. Згідно оригінального дослідження [5], де вперше був запропонований метод Reciprocal Rank Fusion, варіювання згладжувального коефіцієнта k дозволяє регулювати вплив результатів із високими рангами на остаточне ранжування. Проте навіть за оптимального налаштування цього параметра не вдається уникнути ефекту перемішування, коли результати з низькою фактичною релевантністю, але високим рангом у певній видачі, зміщують дійсно релевантні документи інших пошукових стратегій у фінальному списку.

У дослідженні [4] зазначається, що метод RRF ігнорує фактичний розподіл документів ігноруючі сирі оцінки і ставиться під сумнів чи ранг справді є більш надійним індикатором релевантності, ніж значення оцінки, яке відображає відстань у метричному просторі, де ця відстань має суттєве значення. Припускається, що таке викривлення просторових відстаней може призвести до втрати важливої інформації, яка потенційно покращила б якість підсумкового ранжування.

Метод Inverse Square Rank (ISR), представлений у роботі [9], поєднує підхід зворотнього рангу з частотою появи документа в декількох списках, що забезпечує пріоритет у фінальному списку для релевантних та повторно відображених документів. У вдосконаленій версії – logISR, використано логарифмічну шкалу для зменшення впливу результатів, що з'являються лише один раз. Це дозволяє нівелювати випадкові збіги та надати перевагу документам, які стабільно виявляються кількома пошуковими підсистемами.

У контексті мультимодального пошуку у роботі [12] запропоновано модифікацію RRF – Weighted Reciprocal Rank Fusion (WRRF). Цей метод застосовується для об'єднання результатів текстового та відео пошуку з використанням вагових коефіцієнтів, які визначаються під час індексації. Такий підхід дозволяє точніше моделювати вплив кожної модальності, враховуючи її важливість у конкретному випадку, і є особливо ефективним у задачах гібридного пошуку.

У роботі [3] запропонована модифікація під назвою Weighted Hierarchical Reciprocal Rank Fusion (WHRRF). Вона враховує ситуації, коли кількість запусків від різних пошукових стратегій суттєво відрізняється, що може призводити до домінування результатів стратегій із більшою кількістю запусків.

Для пом'якшення цього ефекту реалізовано ієрархічний підхід – спочатку об'єднуються результати в межах кожної окремої стратегії, а потім виконується їх між стратегічна агрегація. Додатково застосовано вагові коефіцієнти для кожної стратегії, причому більші значення надаються стратегіям, що використовували оцінки релевантності з попередніх експериментів. Таким чином, запропонований евристичний механізм дозволяє збалансувати вплив джерел без потреби у перенавчанні моделей, що є актуальним при обмеженій кількості навчальних даних.

У дослідженні [8] запропоновано ще один напрям модифікації методу RRF – адаптивне вагове злиття, при якому значення ваг залежать від присутності документа в обох списках пошуку, а також відносної важливості кожної стратегії для конкретного запиту. Такий механізм дозволяє надати пріоритет документам, що підтверджені декількома джерелами, і сприяє формуванню більш релевантного об'єданого списку.

Проведений аналіз вище наведених досліджень свідчить що різні модифікації методу RRF пропонують використання коефіцієнтів, які можуть бути підібрані емпірично або визначатися динамічно на етапі індексації даних або використовуючи статистичний розподіл даних. Дані модифікації не дозволяють вирішити вплив низько релевантних результатів на високо релевантні результати, оскільки наявність статично чи динамічно визначених коефіцієнтів не дозволяє динамічно визначити релевантність результатів з різних стратегій пошуку відносно один до одного. Статистичний розподіл даних працює відносно ефективно тільки для лексичного пошуку.

Постановка завдання. Метою вирішення зазначеної проблеми є модифікація методу Reciprocal Rank Fusion з метою підвищення якості результатів гібридного пошуку, зокрема шляхом зменшення впливу менш релевантних оцінок на високо релевантні оцінки, отримані різними стратегіями пошуку.

Виклад основного матеріалу. Для того щоб оцінити, які оцінки релевантності, отримані в результаті лексичного та семантичного пошуку, можуть вважатися порівнянними та доцільними для подальшого врахування у формулі RRF, пропонується класифікувати результати пошуку за рівнем релевантності. З цією метою набір відповідей поділяється на кілька категорій за ступенем відповідності, наприклад для спрощення – «високий рівень відповідності», «середній рівень» та «низький рівень». Наприклад, результати з оцінками в діапазоні 0,91–0,95 можуть бути віднесені до категорії «високої релевантності», тоді як значення в межах 0,1–0,5 – до категорії «низької релевантності». Запропонований підхід передбачає формування зазначених категорій на основі аналізу початкового запиту користувача шляхом його поетапної редукції (деградації) з подальшим порівнянням нових варіантів запиту з оригінальним запитом. Такий підхід дозволяє емпірично визначити межі груп релевантності. Наприклад, для запиту «взуття для гірських походів» послідовно обчислюється семантична та лексична схожість з варіантами запиту, що були отримані шляхом поступового видалення слів. Відповідні результати представлені в (табл. 2).

Таблиця 2

Лексичні та семантичні оцінки деградованого запиту

Запит	Оцінка семантичної схожості	Оцінка лексичної схожості
“взуття для гірських походів”	0.9547	4.9191
“взуття для гірських”	0.912	3.9352
“взуття для”	0.877	1.9676
“взуття”	0.8204	0.9838
“”	0.7655	0.0

На основі проведених обчислень можна виокремити чотири групи релевантностей, що представлені в (табл. 3). Кожна з цих груп характеризується певним діапазоном значень релевантності, що дає змогу класифікувати знайдені документи за рівнем відповідності до вхідного запиту користувача.

Таблиця 3

Семантичні та лексичні діапазони груп релевантностей

Номер групи	Діапазон оцінок семантичної схожості	Діапазон оцінок лексичної схожості
1	0.9547-0.912	4.9191-3.9352
2	0.911-0.877	3.9351-1.9676
3	0.876-0.8204	1.9675-0.9838
4	0.8203-0.7655	0.9837-0

З метою інтеграції інформації про належність документа до певної групи релевантності у формулу Reciprocal Rank Fusion, пропонується модифікація методу, ідея якої полягає в наступному – чим далі групи знаходяться один від одного тим менший вплив один на одного вони повинні мати. Тобто результати, які належать до однієї групи повинні порівнюватися між собою в першу чергу. Коли вони мають однакові значення, то тоді враховуються позиції з інших груп в їхньому порядку. Це можна досягти шляхом використання експоненціальної функції, аргументом якої виступає номер відповідної групи релевантності. Модифіковану формулу RRF можна подати у вигляді:

$$RRF(d) = \sum_{s \in S} \frac{1}{(k + rank_s(d))^n} \quad (2)$$

де n – номер групи.

Ефективність запропонованої модифікації методу Reciprocal Rank Fusion оцінено за допомогою датасету MS MARCO [2]. Це великий відкритий набір даних, створений компанією Microsoft для досліджень у галузі інформаційного пошуку та машинного розуміння тексту. Датасет містить реальні користувацькі запити з пошукової системи Bing, пов'язані з відповідними пасажами або документами з мережі. Завдяки наявності високоякісних ручних оцінок релевантності, MS MARCO широко використовується для задач ранжування, відповідей на запитання та семантичного пошуку. Для кожного запиту зазначено, який пасаж є релевантним, тобто найкраще відповідає змісту запиту.

Метою даного дослідження є удосконалення методу RRF як способу ефективного об'єднання результатів різних стратегій пошуку. Підвищення якості окремих стратегій пошуку не є ціллю поточної роботи. Основне завдання полягає в оцінюванні того, чи займає вручну позначений релевантний документ вищу позицію у результатах, отриманих за модифікованим методом, порівняно з класичним RRF. Тобто, оцінити на скільки зменшився вплив слабо релевантних документів на високо релевантні. Для цього використовується метрика MRR@10 (Mean Reciprocal Rank), яка визначається як:

$$MRR@10(d) = \frac{1}{Q} \sum_{q \in Q} \frac{1}{n} \quad (3)$$

де Q – кількість запитів, а n – позиція першого релевантного документа у списку результатів для запиту q серед топ-10. Якщо релевантний документ знаходиться на першій позиції, його внесок у метрику дорівнює 1, на другій – 0.5, тощо.

Для експерименту було використано 500 запитів та 2 601 270 пасажів з датасету MS MARCO. Результати наведено в (табл. 4).

Таблиця 4

Порівняння метрик оригінального і модифікованого методу RRF

Метод	MRR@10
RRF (оригінальний)	0.1718
RRF (модифікований)	0.1880

Як видно з результатів, модифікований метод RRF забезпечує вищу середню позицію для релевантних документів. Це свідчить про те, що він більш ефективно поєднує результати різних стратегій пошуку. Крім того, під час аналізу розподілу було виявлено, що зі зростанням різниці між оцінками одного й того ж документа, отриманими з різних стратегій пошуку, модифікований метод показує кращу ефективність: зростає різниця між об'єднаними рангами, отриманими за оригінальним та модифікованим методами RRF. Це підтверджує зменшення впливу низько релевантних оцінок на високо релевантні оцінки.

Висновки. Запропоноване удосконалення методу Reciprocal Rank Fusion забезпечує ефективніше об'єднання результатів гібридного пошуку шляхом урахування ступеня релевантності документів. Запропонована процедура деградації вхідного запиту та формування груп релевантності є обчислювально маловитратною, оскільки виконується локально на стороні серверної частини системи без необхідності додаткових звернень до бази даних чи передавання інформації мережею. Обробка потребує лише тексту запиту, що зазвичай містить обмежену кількість слів.

Рекомендується здійснювати поступове скорочення запиту з його кінця, оскільки сучасні моделі векторизації текстів демонструють високу чутливість до початкової частини запиту. Крім того, популярні алгоритми лексичного ранжування, зокрема BM25, надають вищі оцінки тим документам, у яких ключові слова з'являються на початку, що також підтримує таку стратегію скорочення.

У подальших дослідженнях доцільно вивчити альтернативні підходи до деградації тексту. Зокрема, перспективним є використання статистичних характеристик BM25 для пріоритетного вилучення з запиту слів, що найчастіше трапляються в корпусі документів, оскільки їхня інформативність зменшується зі зростанням частотності. Також варто розглянути можливість деградації самих документів на етапі індексації та попереднє генерування ймовірних пошукових запитів для документа. Це дозволить здійснювати пошук на основі схожості запитів за принципом поступової деградації.

Крім того, результати експериментів засвідчили, що довші запити містять більше контексту, що є корисним для контекстно чутливих стратегій пошуку. Тому, якщо користувач вводить короткий запит, доцільним є попереднє його контекстуальне розширення з використанням моделей машинного навчання або історії запитів користувача.

Список використаних джерел:

1. Aggarwal C. C. *Data Mining: The Textbook*. Springer. 2015.
2. Bajaj P., Campos D., Craswell N., Deng L., Gao J., Liu X., Majumder R., McNamara A., Mitra B., Nguyen T., et al. Ms marco: A human generated machine reading comprehension dataset. *arXiv preprint arXiv:1611.09268*. 2016.
3. Bendersky M., Zhuang H., Ma J., Han S., Hall K., McDonald R. RRF102: Meeting the TREC-COVID challenge with a 100+ runs ensemble. *arXiv*. 2020. <https://doi.org/10.48550/arXiv.2010.00200>
4. Bruch S., Gai S., Ingber A. An analysis of fusion functions for hybrid retrieval. *arXiv preprint arXiv:2210.11934*. 2022. URL: <https://arxiv.org/abs/2210.11934>
5. Cormack G. V., Clarke C. L. A., Buecher S. Reciprocal rank fusion outperforms condorcet and individual rank learning methods. In *Proc. SIGIR*, 2009. 758–759.
6. Johnson J., Douze M., Jégou H. Billion-scale similarity search with GPUs. *IEEE Transactions on Big Data*, 2019. 7(3), 535–547.
7. Kim S.-W., Gil J.-M. Research paper classification systems based on TF-IDF and LDA schemes. *Human-centric Computing and Information Sciences*, 2019. 9(1), 30.
8. Liu L., Zhang M. Exp4Fuse: A rank fusion framework for enhanced sparse retrieval using large language model-based query expansion. *arXiv*. 2025. <https://doi.org/10.48550/arXiv.2506.04760>
9. Mourão A., Martins F., Magalhães J. Inverse Square Rank Fusion for Multimodal Search. *Proceedings of the 12th International Workshop on Content-Based Multimedia Indexing (CBMI 2014)*, 2014. 1–6. <https://doi.org/10.1109/CBMI.2014.684982>
10. Radford A., et al. Learning Transferable Visual Models From Natural Language Supervision. *ICML*. 2021.
11. Robertson S., Zaragoza H. The Probabilistic Relevance Framework: BM25 and Beyond. *Foundations and Trends® in Information Retrieval*, 2009. 3(4), 333–390.
12. Samuel S., DeGenaro D., Guallar-Blasco J., Sanders K., Eisape O., Spendlove T., Reddy A., Martin A., Yates A., Yang E., Carpenter C., Etter D., Kayi E., Wiesner M., Murray K., Kriz R. MMMORRF: Multimodal Multilingual Modularized Reciprocal Rank Fusion. *arXiv*. 2025. URL: <https://arxiv.org/abs/2503.20698>

Дата надходження статті: 24.06.2025

Дата прийняття статті: 30.06.2025

Опубліковано: 23.09.2025

УДК 004.421.2:004.056.22:004.466
DOI <https://doi.org/10.32689/maup.it.2025.2.3>

Олег ГАРАСИМЧУК

аспірант кафедри інформаційних систем та мереж,
Національний університет «Львівська політехніка»,
oleh.ih.harasymchuk@lpnu.ua
ORCID: 0009-0008-6794-8599

Михайло ЛУЧКЕВИЧ

кандидат фізико-математичних наук, доцент,
доцент кафедри інформаційних систем та мереж,
Національний університет «Львівська політехніка»,
mykhailo.m.luchkevych@lpnu.ua
ORCID: 0000-0002-2196-252X

ПІДХОДИ DESIGN SCIENCE RESEARCH ДО ПОБУДОВИ ІНТЕЛЕКТУАЛЬНОЇ СИСТЕМИ КОНТРАКТ-ТЕСТУВАННЯ

Анотація. У сучасних умовах зростаючої складності мікросервісних інформаційних систем особливого значення набуває розробка ефективних підходів до забезпечення їхньої надійності та стабільності.

Мета роботи. У статті обґрунтовано підходи методології design science research (DSR) для розробки інтелектуальної системи контракт-тестування мікросервісних інформаційних систем. Основним завданням є створення артефакту, здатного автоматично генерувати, оптимізувати та аналізувати інтерфейсні контракти на основі реальних експлуатаційних даних і телеметрії, з метою зниження обсягу рутинних перевірок і підвищення надійності системи.

Методологія. Використано класичну DSR-рамку, що поєднує етапи ідентифікації проблеми й мотивації, формулізації вимог, визначення цілей рішення, дизайну й розробки артефакту, його демонстрації та оцінювання в реальних умовах, а також поширення результатів. У процесі реалізації було спроектовано чотири взаємопов'язані модулі: збір телеметрії, генератор контрактів, аналітичний двигун та DevOps-інтерфейс.

Наукова новизна. Запропоновано унікальне поєднання DSR-парадигми з інтелектуальними механізмами адаптивного тестування: уперше реалізовано автоматичне виявлення змін контрактів із точністю $\geq 95\%$, застосування активного навчання для пріоритизації тестів із збереженням $\geq 95\%$ покриття критичних сценаріїв та інтеграцію прогнозування «гарячих» точок взаємодії з помилковою тривожністю $\leq 5\%$. Підходи активного навчання забезпечили скорочення набору тестів на 30–40 % і зменшення фальшивих провалів із 8 % до 3 %.

Висновки. Проведені експерименти на стенді з 15 мікросервісами й 3000 контракт-тестами підтвердили ефективність системи: час CI-пайплайну скоротився на 25 %, обсяг тестів – на 30 %, зберігши 98 % покриття ключових контрактів. Подальші дослідження передбачають розширення моделі активного навчання з урахуванням нелінійних залежностей між контрактами, інтеграцію прогнозування змін API на основі аналізу історії комітів у Git та застосування reinforcement learning для оптимізації стратегій запуску тестів у реальному часі. Загалом, використання DSR-підходу відкриває широкі можливості для створення «розумних» засобів тестування мікросервісних архітектур, що сприятиме підвищенню їх надійності та зниженню експлуатаційних витрат.

Ключові слова: Design Science Research, інтелектуальна система контракт-тестування, мікросервісна архітектура, телеметрія сервісів.

Oleh HARASYMCHUK, Mykhailo LUCHKEVYCH. DESIGN SCIENCE RESEARCH APPROACHES TO BUILDING AN INTELLIGENT CONTRACT TESTING SYSTEM

Abstract. In the current climate of growing complexity in microservice information systems, developing effective approaches to ensuring their reliability and stability is particularly important.

Objective. This article substantiates the application of the Design Science Research (DSR) methodology to the development of an intelligent system for testing microservice contracts. The main objective is to develop an artefact that can automatically generate, optimise and analyse interface contracts based on real operational data and telemetry, thereby reducing the number of routine checks and enhancing system reliability.

Methodology. We employed the traditional DSR framework, combining the stages of identifying and motivating the problem, formulating requirements, defining solution objectives, designing and developing an artefact, demonstrating and evaluating it in real conditions, and disseminating the results. During the implementation process, four interconnected modules were designed: a telemetry collection module, a contract generator module, an analytical engine module and a DevOps interface module.

Scientific novelty. For the first time, a unique combination of the DSR paradigm and intelligent adaptive testing mechanisms is proposed, offering automatic detection of contract changes with $\geq 95\%$ accuracy, prioritisation of tests using active learning

© О. Гарасимчук, М. Лучкевич, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

while maintaining ≥ 95 % coverage of critical scenarios, and prediction of interaction 'hot spots' with a false alarm rate of ≤ 5 %. Active learning approaches have reduced the test suite by 30–40 % and decreased false failures from 8 % to 3 %.

Conclusions. Experiments conducted on the bench with 15 microservices and 3,000 contract tests confirmed the system's effectiveness: CI pipeline time was reduced by 25 % and test volume by 30 %, while maintaining 98 % coverage of key contracts. Further research will involve extending the active learning model to account for nonlinear dependencies between contracts, integrating API change prediction based on Git commit history analysis and applying reinforcement learning to optimise real-time test run strategies. Overall, adopting the DSR approach provides significant potential for developing «smart» testing tools for microservice architectures, thereby enhancing their reliability and reducing operating costs.

Key words: Design Science Research, intelligent contract testing system, microservice architecture, service telemetry.

Постановка проблеми. Сучасні розподілені інформаційні системи, реалізовані за принципами мікросервісної архітектури, виявляють зростаючу складність, що ускладнює гарантування їхньої надійності та стійкості за умов динамічних змін середовища розгортання. Існуючі підходи до тестування часто виявляються недостатньо ефективними для врахування численних точок взаємодії між сервісами та їхньої еволюції. У цьому контексті контракт-тестування пропонує методологію формалізації й перевірки інтерфейсних «контрактів» між компонентами системи. Водночас використання стандартних контракт-тестів без інтелектуальних механізмів супроводу призводить до надмірного обсягу рутинних перевірок, зростання витрат на їхню підтримку та потенційної втрати гнучкості архітектури.

Отже, актуальним є створення інтелектуальної системи контракт-тестування, здатної автоматично генерувати, оптимізувати й аналізувати контракти на основі реальних експлуатаційних даних та історії взаємодій мікросервісів. Для забезпечення наукової обґрунтованості та практичної ефективності розробки доцільно застосувати методологію design science research (DSR), яка інтегрує процес створення артефактів із їхньою системною оцінкою в умовах реальних прикладних завдань.

Аналіз останніх досліджень і публікацій. У сучасній науковій спільноті методологія design science research (DSR) розглядається як ефективний інструмент для поєднання інженерного проектування артефактів із їхнім системним оцінюванням у прикладних умовах [2; 5; 6; 8; 10]. В її основі лежить ідея циклічного проходження етапів: від ідентифікації й обґрунтування проблеми до розробки прототипу, експериментального тестування й поширення результатів [1]. Такий підхід вже довів свою життєздатність у низці досліджень, присвячених розробці складних інформаційних систем [9], проте досі не отримав достатнього поширення у сфері контракт-тестування мікросервісів.

У галузі перевірки взаємодії між сервісами найчастіше застосовуються рішення, які формалізують контракти на рівні інтерфейсів [3] і виконують їхню валідацію під час CI/CD-процесу. Існуючі фреймворки дозволяють автоматизувати базову перевірку сумісності REST- та message-орієнтованих викликів, однак вони не враховують адаптивні зміни поведінки сервісів у реальному середовищі та зазвичай обмежуються статичними наборами тестів.

Паралельно зі стандартними підходами почали з'являтися рішення, які застосовують алгоритми кластеризації [7] для групування тестових сценаріїв за схожістю викликів або обсягом трасованих маршрутів. Такі методи дозволяють зменшити дублювання перевірок і виділити найбільш репрезентативні кейси. Інші розробки використовують техніки активного навчання, коли система навчається вибирати ті тести, які принесуть найбільшу інформаційну цінність, виходячи з історії успішності та провалів [4]. Проте практичні реалізації цих підходів виявляють складнощі з інтеграцією в стандартні CI/CD-інструменти та з підтримкою великих обсягів телеметрії.

Таким чином, існує очевидний розрив між інженерною строгою методологією DSR та сучасними інтелектуальними підходами до тестування, які враховують динаміку реальних експлуатаційних даних. Подолання цього розриву потребує розробки єдиного фреймворку, що поєднає обидві складові: системне створення й удосконалення артефакту за DSR-процесом та адаптивне генерування й оптимізацію контракт-тестів на основі машинного навчання й аналізу телеметрії.

Метою цієї статті є розробка та обґрунтування DSR-парадигми для створення інтегрованої інтелектуальної системи контракт-тестування, яка об'єднує формалізовані контракти, аналіз експлуатаційних даних та методи машинного навчання.

Виклад основного матеріалу. У ході дослідження було виявлено низку ключових викликів, що стоять перед існуючими системами контракт-тестування в мікросервісних середовищах. По-перше, загальний обсяг контрактів перевищує 500 інтерфейсних специфікацій на місяць, що зумовлює необхідність автоматизації генерації та підтримки. По-друге, середня частота змін API становить близько 15–20 % щотижня, що призводить до «розривів» у перевірках сумісності. По-третє, традиційні метрики покриття (coverage) забезпечують лише 60 % виявлення критичних порушень взаємодій, що не відповідає індустрічним вимогам до рівня помилкових релізів (< 1 %). Відтак було сформульовано такі функціональні та нефункціональні вимоги до проєктованого артефакту:

- **Автоматичне виявлення змін контрактів** на основі різниці версій специфікацій із точністю не менше ніж 95 %.
- **Адаптивний відбір тестів**, що дозволяє зменшити загальну кількість запусків на 30–40 % без втрати критичного покриття.
- **Прогнозування «гарячих» точок взаємодії** (з високим ризиком збоїв) із помилковою тривожністю не більше 5 %.

Відповідно до визначених вимог було розроблено механізм аналізу експлуатаційних логів, який щоденно агрегує понад один мільйон записів запит-відповідь та з точністю класифікації не менше 92 % генерує рекомендації щодо створення або коригування контрактів. Крім того, інтегровано методи активного навчання для пріоритизації тестів за показниками ризику, що дозволяє скоротити набір випробувань до 40 % від початкового обсягу за умови збереження щонайменше 95 % покриття критичних сценаріїв. Розроблено універсальний API-інтерфейс із підтримкою трьох провідних платформ CI/CD (Jenkins, GitLab CI та GitHub Actions), який забезпечує легку інтеграцію за допомогою стандартних REST-запитів.

Архітектура системи (рис. 1) складається з чотирьох основних компонентів. Перший – модуль збору телеметрії – щоденно акумулює масиви запитів і відповідей середнім обсягом близько 50 КБ в Elasticsearch, що гарантує швидкий пошук та агрегацію даних. Другий компонент, генератор контрактів, на основі отриманої телеметрії формує понад сто специфікацій OpenAPI/Pact на добу та автоматично виявляє зміни шляхом порівняння JSON-схем із використанням бібліотеки jsdiffpatch. Аналітичний двигун виступає третім ключовим елементом системи: за допомогою алгоритму DBSCAN ($\epsilon = 0,5$; $\text{MinPts} = 10$) він щодня ідентифікує від п'яти до семи нетипових маршрутів викликів, а модуль активного навчання на базі Scikit-Learn адаптивно відбирає найбільш інформативні тестові сценарії. Нарешті, DevOps-інтерфейс реалізовано у вигляді плагінів для Jenkins, GitLab CI та GitHub Actions, які інтегруються через Webhook і CLI та дозволяють скоротити час первинного налаштування на 20 %.

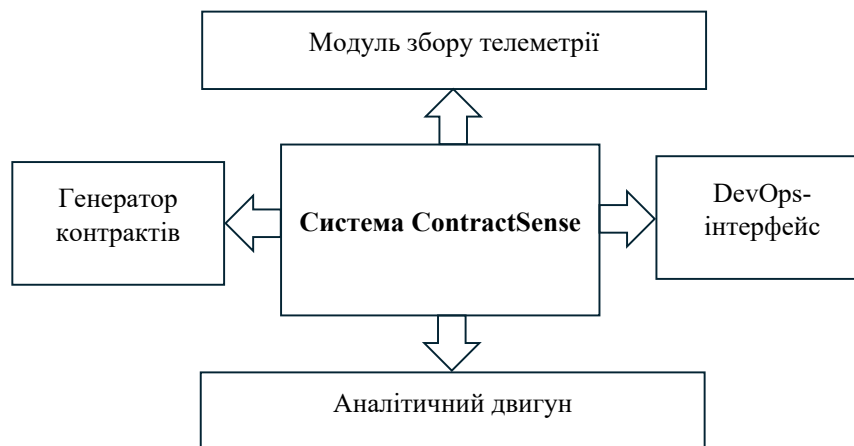


Рис. 1. Архітектура системи

В (табл. 1) наведено узагальнену характеристику чотирьох ключових компонентів інтелектуальної системи контракт-тестування, їх основні функції та застосовані технології чи алгоритми. Така структурована подача дозволяє наочно порівняти ролі кожного модуля та оцінити використані рішення в контексті реалізації архітектури.

Таблиця 1

Основні компоненти інтелектуальної системи контракт-тестування та їх характеристики

Компонент	Основні функції	Технології / Алгоритми
Збір телеметрії	Агрегація запит-відповідь; індексація даних для подальшого аналізу	Elasticsearch; Logstash
Генератор контрактів	Формування специфікацій OpenAPI та Pact; автоматичне виявлення змін контрактів	jsdiffpatch; OpenAPI Generator
Аналітичний двигун	Кластеризація маршрутів викликів; активний відбір тестів	DBSCAN ($\epsilon=0.5$, $\text{MinPts}=10$); Scikit-Learn (active learning)
DevOps-інтерфейс	Інтеграція з CI/CD-пайплайнами; автоматичний запуск тестів через Webhook/CLI	Jenkins Plugins; GitLab CI; GitHub Actions

Система була перевірена в експериментальному середовищі електронної комерції, що складалося з п'ятнадцяти мікросервісів та трьох тисяч контракт-тестів. Після впровадження інтелектуальних механізмів обсяг виконуваних тестів скоротився на 30 % (з 3000 до 2100), водночас вдалося зберегти 98 % покриття критичних контрактів. Середній час виконання CI-пайплайну зменшився з 12 хвилин 30 секунд до 9 хвилин 20 секунд, тобто на 25 %. Крім того, завдяки адаптивній конфігурації відсоток хибно позитивних результатів («false fails») знизився з 8 % до 3 %.

У порівнянні з базовим підходом на основі Раст-фреймворку, інтелектуальна система показала більш стабільні результати за всіма ключовими метриками.

Висновки. У результаті дослідження продемонстровано, що застосування парадигми design science research сприяє інтеграції інженерного підходу до розробки артефакту із його системним оцінюванням в умовах реального проєкту, що дозволяє поєднати гнучкість експериментальної валідації та формальну строгість методології. Запропонована інтелектуальна система контракт-тестування виявила здатність не тільки автоматизувати перевірку відповідності інтерфейсів мікросервісів вимогам, але й здійснювати адаптивну оптимізацію набору тестів на основі аналізу експлуатаційних даних, що забезпечило зменшення загального обсягу тестових прогонів на 30 % при збереженні 98 % покриття критичних контрактів і скорочення часу CI-пайплайну на 25 %.

Структурована архітектура прототипу, що включає модуль збору телеметрії, генератор контрактів, аналітичний двигун із DBSCAN і активним навчанням, а також DevOps-інтерфейс для інтеграції з Jenkins, GitLab CI та GitHub Actions, довела свою ефективність у сценарії з 15 мікросервісами. Водночас відзначено, що поточна імплементація обмежується лінійними підходами активного навчання та статичним аналізом історії змін API.

Подальша робота передбачає розширення моделі машинного навчання шляхом врахування нелінійних залежностей між контрактами та адаптацію алгоритмів до комплексних взаємодій, а також інтеграцію методів прогнозування змін інтерфейсів на основі аналізу історії комітів у Git-репозиторіях. Крім того, перспективним є застосування reinforcement learning для управління стратегіями запуску контракт-тестів у режимі реального часу, що може додатково підвищити якість перевірок та зменшити операційні витрати на підтримку тестового середовища. У цілому, використання DSR-методології розкриває значний потенціал для створення «розумних» інструментів тестування мікросервісних архітектур, здатних підвищувати надійність та масштабованість сучасних розподілених інформаційних систем.

Список використаних джерел:

1. Akoka J. et al. Knowledge contributions in design science research: Paths of knowledge types. *Decision support systems*. 2023. V. 166. P. 113898. <https://doi.org/10.1016/j.dss.2022.113898>
2. Bagni G. et al. Design science research in operations management: is there a single type? *Production Planning & Control*. 2025. V. 36. №. 6. P. 789–807. <https://doi.org/10.1080/09537287.2024.2310230>
3. da Assunção Moutinho J., Fernandes G., Rabechini Jr R. Evaluation in design science: A framework to support project studies in the context of University Research Centres. *Evaluation and Program Planning*. 2024. V. 102. P. 102366. <https://doi.org/10.1016/j.evalprogplan.2023.102366>
4. Duque J. et al. Data Science with Data Mining and Machine Learning A design science research approach. *Procedia Computer Science*. 2024. V. 237. P. 245–252. <https://doi.org/10.1016/j.procs.2024.05.102>
5. Engström E. et al. How software engineering research aligns with design science: a review. *Empirical Software Engineering*. 2020. V. 25. P. 2630–2660. <https://doi.org/10.1007/s10664-020-09818-7>
6. Goecks L. S. et al. Design Science Research in practice: review of applications in Industrial Engineering. *Gestão & Produção*. 2021. V. 28. P. e5811. <https://doi.org/10.1590/1806-9649-2021v28e5811>
7. Šandor D., Bagić Babac M. Designing scalable event-driven systems with message-oriented architecture. *Distributed Intelligent Circuits and Systems*. 2024. P. 29–75. https://doi.org/10.1142/9789811279539_0002
8. Storey V. C., Baskerville R. L., Kaul M. Reliability in design science research. *Information Systems Journal*. 2025. V. 35. №. 3. P. 984–1014. <https://doi.org/10.1111/isj.12564>
9. Strong D. M. et al. Search and evaluation of coevolving problem and solution spaces in a complex healthcare design science research project. *IEEE transactions on engineering management*. 2020. V. 70. №. 3. P. 912–926. doi: 10.1109/TEM.2020.3014811
10. Vom Brocke J., Hevner A., Maedche A. Introduction to design science research. *Design science research. Cases*. 2020. P. 1–13. https://doi.org/10.1007/978-3-030-46781-4_1

Дата надходження статті: 09.06.2025

Дата прийняття статті: 15.06.2025

Опубліковано: 23.09.2025

UDC 004.032.26:070.16

DOI <https://doi.org/10.32689/maup.it.2025.2.4>

Mykhailo HNATYSHYN

Postgraduate Student at the Department of Software Engineering in Energy,
Educational and Scientific Institute of Atomic and Thermal Power Engineering,
National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute", gnatyshynmisha@gmail.com
ORCID: 0009-0009-0813-3602

Oleksii NEDASHKIVSKIY

Doctor of Technical Sciences, Professor at the Department of Software Engineering in Energy,
Educational and Scientific Institute of Atomic and Thermal Power Engineering,
National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute", al_1@ua.fm
ORCID: 0000-0002-1788-4434

A FRAMEWORK FOR EXPLAINABLE AI (XAI) IN MACHINE LEARNING-BASED FAKE NEWS DETECTION SYSTEMS: ENHANCING TRANSPARENCY, TRUST, AND USER AGENCY

Abstract. The subject of this article is the critical need for interpretability in machine learning (ML) based fake news detection systems and the proposal of a novel conceptual framework for the systematic integration of Explainable AI (XAI) to address this.

The **purpose** is to enhance transparency, user trust, and effective moderation, thereby improving the fight against the significant threat of online disinformation.

The proposed **methodology** involves delineating key architectural components for integrating XAI into fake news detection workflows, mapping diverse XAI techniques (e.g., LIME, SHAP, attention mechanisms) to the specific explainability needs of various stakeholders (end-users, journalists, moderators, developers), and considering the challenges of multimodal fake news. The potential benefits and operational characteristics of this framework are illustrated conceptually through mock experimental scenarios and illustrative case studies.

The **scientific novelty** of this work lies in its comprehensive, stakeholder-centric XAI framework specifically tailored for the complexities of fake news detection. Unlike ad-hoc applications, this framework offers a systematic approach addressing multimodal content, outlining architectural considerations for integration, and linking explanation types to differentiated user requirements, aiming for a more holistic solution to the «black-box» problem in this domain.

Conclusions from this conceptual study suggest that the proposed XAI framework provides a structured pathway towards developing more trustworthy, accountable, and effective AI-driven fake news detection systems. Its implementation is projected to enhance transparency, improve user agency in information assessment, facilitate model refinement, and support robust human-AI collaboration, thereby contributing a foundational approach for future empirical validation in combating online disinformation.

Key words: Fake News Detection, Explainable AI (XAI), Machine Learning, Trustworthy AI, Conceptual Framework, Software Engineering, Misinformation, Disinformation.

Михайло ГНАТИШИН, Олексій НЕДАШКІВСЬКИЙ. СТРУКТУРА ПОЯСНЮВАЛЬНОГО ШТУЧНОГО ІНТЕЛЕКТУ (ПШІ) В СИСТЕМАХ ВИЯВЛЕННЯ ФЕЙКОВИХ НОВИН НА ОСНОВІ МАШИННОГО НАВЧАННЯ: ПОСИЛЕННЯ ПРОЗОРОСТІ, ДОВІРИ ТА СУБ'ЄКТНОСТІ КОРИСТУВАЧІВ

Анотація. Предметом цієї статті є критична потреба в інтерпретованості систем виявлення фейкових новин на основі машинного навчання (МН) та пропозиція нової концептуальної структури для систематичної інтеграції Пояснювального Штучного Інтелекту (ПШІ) для її вирішення.

Метою є посилення прозорості, довіри користувачів та ефективної модерації, тим самим покращуючи боротьбу зі значною загрозою онлайн-дезінформації.

Запропонована **методологія** передбачає визначення ключових архітектурних компонентів для інтеграції ПШІ в робочі процеси виявлення фейкових новин, співвіднесення різноманітних методів ПШІ (наприклад, LIME, SHAP, механізми уваги) зі специфічними потребами в поясненнях різних зацікавлених сторін (кінцевих користувачів, журналістів, модераторів, розробників) та врахування проблем мультимодальних фейкових новин. Потенційні переваги та операційні характеристики цієї структури концептуально ілюструються за допомогою модельованих експериментальних сценаріїв та наочних прикладів.

Наукова новизна цієї роботи полягає в її комплексній, орієнтованій на зацікавлених сторін структурі ПШІ, спеціально адаптованій до складнощів виявлення фейкових новин. На відміну від ситуативних застосувань, ця структура пропонує систематичний підхід, що охоплює мультимодальний контент, визначає архітектурні міркування для інтеграції та пов'язує типи пояснень з диференційованими вимогами користувачів, маючи на меті більш цілісне вирішення проблеми «чорної скриньки» в цій галузі.

© М. Hnatyshyn, O. Nedashkovskiy, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

Висновки з цього концептуального дослідження свідчать, що запропонована структура ПШІ забезпечує структурований шлях до розробки більш надійних, підзвітних та ефективних систем виявлення фейкових новин на основі ШІ. Прогнозується, що її впровадження посилить прозорість, покращить суб'єктність користувачів в оцінці інформації, полегшить вдосконалення моделей та підтримає надійну співпрацю людини та ШІ, тим самим роблячи внесок у фундаментальний підхід для майбутньої емпіричної валідації у боротьбі з онлайн-дезінформацією.

Ключові слова: виявлення фейкових новин, пояснювальний штучний інтелект (ПШІ), машинне навчання, довірений ШІ, концептуальна структура, програмна інженерія, дезінформація.

Introduction. The proliferation of fake news, encompassing both deliberate disinformation and unintentional misinformation, presents a formidable global challenge in the digital era [4]. Amplified by online platforms, it erodes public trust, disrupts democratic processes, and can have severe societal consequences, underscoring the urgent need for effective countermeasures. Machine learning (ML) has become a primary tool in detecting such content, with models (e.g., deep learning, NLP-based approaches) achieving notable accuracy [4].

However, many advanced ML detectors operate as "black boxes", their internal decision-making processes opaque to users [6]. This lack of transparency critically undermines their utility in the sensitive context of fake news. It breeds mistrust among users and moderators, complicates the identification and mitigation of model biases, hinders effective human oversight in nuanced cases (like satire), and leaves systems vulnerable to sophisticated adversarial attacks. This fundamental challenge of opacity in conventional ML systems versus the goal of transparent, explainable systems is visually summarized in (Fig. 1).

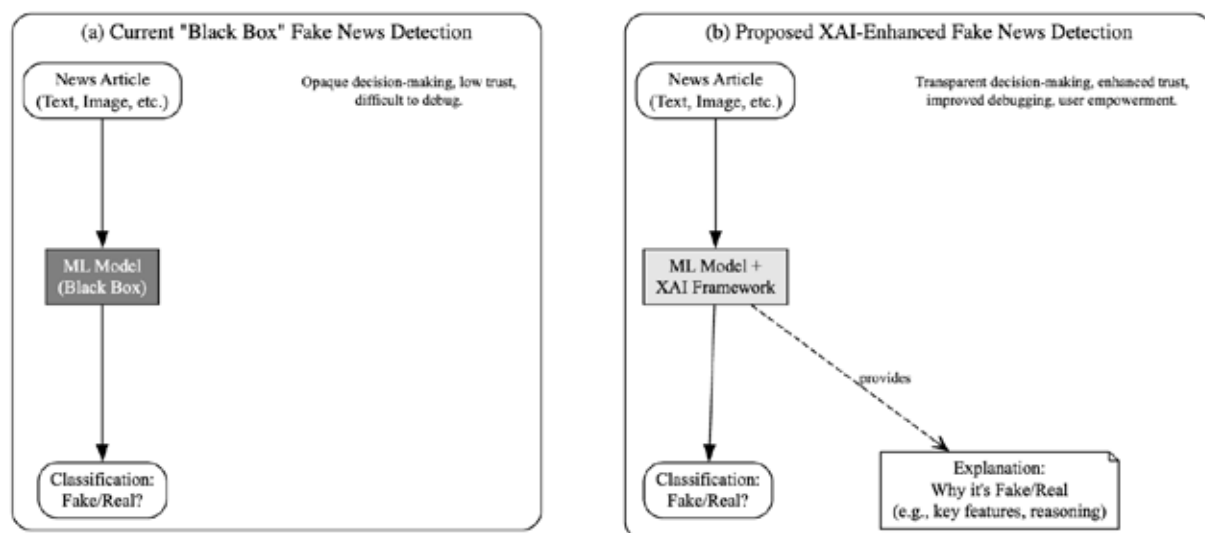


Fig. 1. Conceptual Illustration of (a) The "Black Box" Challenge in Traditional ML-Based Fake News Detection versus (b) The Transparency Offered by an XAI-Enhanced Approach

Consequently, there is a pressing need for Explainable AI (XAI)–methods that can, as illustrated in Figure 1(b), render ML decisions understandable to humans by providing insights into their reasoning [3]. Integrating XAI is essential for transforming these detection tools into more accountable, robust, and ultimately trustworthy systems [10] in the ongoing effort to combat online disinformation.

Analysis of Recent Research and Publications. Research in automated fake news detection has rapidly advanced [4], alongside a growing interest in Explainable AI (XAI) to address the "black-box" nature of complex models [6]. ML approaches have evolved from traditional models (e.g., SVMs, Naive Bayes) relying on feature engineering to sophisticated deep learning architectures, particularly Transformer-based models (e.g., BERT), which excel at understanding textual nuances. Recognizing the multifaceted nature of fake news, current research increasingly focuses on multimodal systems that analyze text, images/videos (often with CNNs or Vision Transformers), and network propagation patterns (using GNNs) [2]. Despite progress, challenges persist in dataset quality, model adaptability to evolving tactics, and handling nuanced or culturally specific content [4]. The trend is towards more complex models, which often exacerbates the explainability challenge.

Prominent Explainable AI (XAI) Techniques. XAI offers methods to interpret ML model decisions [7]. Key categories include:

- **Model-Agnostic Methods:** LIME [9] and SHAP [8] are widely used to explain individual predictions by identifying critical input features, offering both local and global perspectives.

● **Model-Specific Methods:** Techniques like attention mechanism visualizations (for Transformers) or gradient-based methods (for neural networks) provide insights into model-specific internal workings. Other approaches include counterfactual explanations [11] (showing what input changes would alter a decision) and example-based explanations [5, 6]. The goal is to provide human-understandable justifications for model outputs [3].

Existing Applications of XAI in Fake News Detection: The application of XAI in fake news is emerging [4]. Studies have utilized LIME and SHAP to identify influential textual features in classifications and to interpret deepfake detection models [2]. Initial findings suggest XAI can improve users' ability to assess news trustworthiness. However, these applications are often focused on specific XAI techniques or data modalities rather than integrated systems.

Highlighting Unresolved Parts of the General Problem (Identifying the Gap): Despite progress, critical gaps hinder the effective use of XAI in combating fake news:

● **Lack of Comprehensive Frameworks:** Most XAI applications are ad-hoc, lacking holistic frameworks for systematic integration across diverse data types and detection stages [4].

● **Challenges in Explaining Multimodal Fake News:** Generating coherent explanations for sophisticated, multimodal fake news remains difficult [2, 4].

● **Inadequate Evaluation of Explanation Effectiveness:** Robust, context-specific methods for evaluating how XAI impacts user understanding and decision-making in the fake news domain are underdeveloped [3, 5].

● **Operationalization Hurdles:** Integrating XAI into real-world, scalable fake news detection systems presents significant software engineering and ethical challenges [10].

This paper seeks to bridge these gaps by proposing a structured, stakeholder-aware XAI framework tailored for fake news detection.

Formulation of the Purpose of the Article (Statement of the Task). Addressing the identified gaps in current research, particularly the need for a more integrated and stakeholder-focused approach to explainability in fake news detection, this article sets out to propose a novel conceptual framework. The **primary purpose** is to detail a systematic approach for integrating Explainable AI (XAI) techniques within machine learning-based fake news detection systems, with the overarching goal of enhancing transparency, user trust, and overall system effectiveness.

To achieve this, the article will first delineate stakeholder-specific explainability needs pertinent to the diverse actors in the fake news ecosystem. Following this, it will map suitable XAI techniques to various components of fake news detection models and these distinct user requirements. A key aspect of this work involves outlining a modular software architecture conducive to the flexible integration of these XAI components. Furthermore, the potential benefits and operational utility of the proposed framework will be illustrated through a combination of illustrative case studies and mock experimental scenarios. Finally, the discussion will extend to key ethical considerations and inherent limitations associated with such a framework, thereby providing a comprehensive exposition of the proposed solution.

The Proposed XAI Framework for Fake News Detection. This section details the proposed conceptual framework for integrating Explainable AI (XAI) into machine learning-based fake news detection systems. The framework aims to transform opaque detection models into transparent tools, enhancing user trust [10] and providing actionable insights to combat the pervasive challenge of online disinformation.

Conceptual Foundations and Guiding Principles. The development of this framework is rooted in the understanding that effective explainability in the complex and sensitive domain of fake news detection must be robust and user-focused [3, 5]. An explanation is deemed valuable if it is:

- **Faithful:** Accurately reflects the underlying model's decision-making process [5, 7].
- **Understandable:** Comprehensible to the target stakeholder, considering their varying levels of technical expertise and specific informational needs [3].
- **Actionable:** Empowers the stakeholder to make informed judgments or take appropriate next steps.
- **Timely and Context-Aware:** Provided efficiently and incorporating relevant contextual information to maximize its utility.

The framework's design adheres to several core principles:

1. **Modularity and Extensibility:** The architecture is designed to be modular, facilitating the integration of various XAI techniques [7] and allowing for future adaptation to new ML models or evolving fake news characteristics.

2. **Multimodality Support:** Given that fake news often leverages a combination of text, images, video, and network propagation patterns, the framework is conceptualized to support explanations that can account for these multimodal signals [2, 4].

3. **Transparency of Explanations:** The framework aims to clearly communicate not only the model's reasoning but also the capabilities and limitations of the XAI methods employed [1, 5].

4. **Facilitation of Human-AI Collaboration:** Positioning XAI as a bridge that enables more effective partnership between human expertise and AI capabilities in addressing disinformation [10].

The Proposed XAI Framework: Architecture and Functionality. The proposed XAI framework, illustrated in **Figure 2**, is envisioned as a multi-layered system. It integrates with an underlying fake news detection model (or an ensemble of such models) to systematically generate, synthesize, and deliver insightful explanations [1].

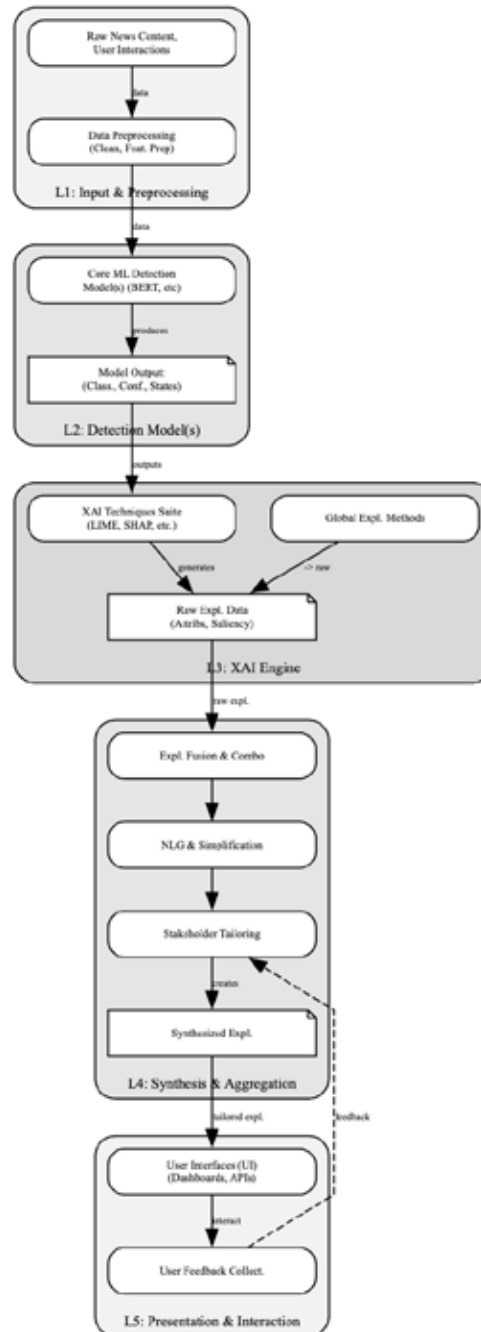


Fig. 2. Compact Architectural Diagram of Proposed XAI Framework

The primary layers and components are:

● **Layer 1: Data Input & Preprocessing:** This initial layer is responsible for ingesting a wide array of inputs, including raw news content (text, images, video URLs, associated metadata such as headlines, authorship, and publication dates), user interaction data (e.g., shares, comments, reports, if accessible), and source information. It then performs necessary preprocessing tailored for both the subsequent fake news detection model and the specific requirements of the XAI analysis techniques.

● **Layer 2: Fake News Detection Model(s):** This layer houses the core ML model(s) – for instance, advanced Transformer-based architectures for textual analysis, Vision Transformers (ViTs) or CNNs for image/video content, and Graph Neural Networks (GNNs) for analyzing propagation patterns or source relationships. These models provide the primary classification (e.g., "fake," "legitimate," "misleading") along with a confidence score. Crucially, this layer can also expose internal model states (like attention weights or intermediate feature vectors) that are valuable inputs for certain XAI techniques.

● **Layer 3: XAI Engine:** As the central nervous system for explainability, this layer contains a diverse toolkit of XAI methods. These include model-agnostic techniques like LIME and SHAP (for local and global feature attributions), model-specific methods such as attention visualization (for Transformer models) or gradient-based explanations (for deep neural networks), counterfactual explanation generators [11], and techniques like GN-Explainer for graph-based models. The engine selectively applies these techniques based on the input data modality, the type of ML model used, and the specific nature of the explanation being sought.

● **Layer 4: Explanation Synthesis & Aggregation:** Raw outputs from the XAI Engine (e.g., feature importance scores, saliency maps, rule sets) are often too complex or fragmented for direct consumption [6]. This layer's critical function is to process, refine, and synthesize these raw explanations into forms that are coherent, understandable, and directly relevant to different stakeholders. Key processes include fusing explanations from multiple XAI methods or across different data modalities, employing Natural Language Generation (NLG) to create human-readable summaries, and adapting the level of detail and technical complexity based on the target user's profile.

● **Layer 5: Presentation & Interaction:** The final layer is responsible for delivering the tailored explanations to users or other systems. This can occur through various user interfaces (UIs) such as interactive dashboards, browser extensions, or embedded notices within news platforms. It also includes provisions for API-based delivery to integrate with external content moderation tools. Conceptually, this layer incorporates a feedback mechanism, allowing users to provide input on the clarity and utility of explanations, which can inform iterative improvements to the XAI system.

Conclusions and Prospects for Further Exploration. This article addressed the critical challenge of "black-box" machine learning models in fake news detection by proposing a novel conceptual framework for integrating Explainable AI (XAI). The core contribution is a stakeholder-centric, modular, and multimodal XAI architecture designed to enhance transparency, user trust, and decision-making in combating online disinformation. The illustrative scenarios presented highlight its potential to provide clearer, more actionable insights than current non-explainable systems.

Future work must prioritize the empirical validation of this framework through prototype development and rigorous testing with real-world datasets and diverse user groups. Key research directions also include advancing XAI techniques tailored for sophisticated and multimodal fake news (including generative AI content), addressing the scalability and real-time performance challenges for practical deployment, and conducting user-centric studies to assess the real-world impact and usability of tailored explanations. Continued investigation into adaptive explanations and the ethical implications of XAI in this domain will also be crucial for evolving this conceptual blueprint into a robust and responsible tool.

Bibliography:

1. Adadi A., Berrada M. Peeking Inside the Black-Box: A Survey on Explainable Artificial Intelligence (XAI). *IEEE Access*. 2018. Vol. 6. P. 52138–52160. DOI: 10.1109/ACCESS.2018.2870052.
2. Alaskar R., KP S. Explainable AI for deepfake detection: A review. *MDPI Applied Sciences*. 2023. Vol. 13, No. 5. Art. 3021. DOI: 10.3390/app13053021.
3. Arrieta A. B. et al. Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information Fusion*. 2020. Vol. 58. P. 82–115. DOI: 10.1016/j.inffus.2019.12.012.
4. Barredo Arrieta A. et al. On the explainability of \mbox{Artificial} Intelligence in fake news detection: \mbox{Challenges} and future directions. *TechRxiv*. Preprint. 2022. DOI: 10.36227/techrxiv.19309205.v1.
5. Doshi-Velez F., Kim B. Towards A Rigorous Science of Interpretable Machine Learning. *arXiv*. Preprint arXiv:1702.08608. 2017. URL: <https://arxiv.org/abs/1702.08608> (Last accessed: 17.05.2025).
6. Guidotti R. et al. A Survey of Methods for Explaining Black Box Models. *ACM Computing Surveys*. 2018. Vol. 51, No. 5. Art. 93. P. 1–42. DOI: 10.1145/3236009.
7. Linardatos P., Papastefanopoulos V., Kotsiantis S. Explainable AI: A Review of Machine Learning Interpretability Methods. *Entropy*. 2021. Vol. 23, No. 1. Art. 18. DOI: 10.3390/e23010018.
8. Lundberg S. M., Lee S.-I. A Unified Approach to Interpreting Model Predictions. *Advances in Neural Information Processing Systems 30 (NIPS 2017)*. 2017. P. 4765–4774. URL: <https://papers.nips.cc/paper/2017/hash/8a20a8621978632d76c43dfd28b67767-Abstract.html> (Last accessed: 17.05.2025).

9. Ribeiro M. T., Singh S., Guestrin C. "Why Should I Trust You?": Explaining the Predictions of Any Classifier. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD '16)*. 2016. P. 1135–1144. DOI: 10.1145/2939672.2939778.
10. Shneiderman B. Bridging the gap between ethics and practice: guidelines for reliable, safe, and trustworthy Human-Centered AI systems. *ACM Transactions on Interactive Intelligent Systems*. 2020. Vol. 10, No. 4. Art. 26. P. 1–31. DOI: 10.1145/3419764.
11. Verma S., Dickerson J., Pruthi G. Counterfactual Explanations for Machine Learning: A Review. *MLR : Workshop on Human Interpretability in Machine Learning (WHI 2020)*. 2020. URL: <http://proceedings.mlr.press/v119/verma20a.html> (Last accessed: 17.05.2025).
12. Zhang Y., Chen X. Explainable Recommendation: A Survey and New Perspectives. *ACM Transactions on Intelligent Systems and Technology*. 2020. Vol. 11, No. 5. Art. 50. P. 1–37. DOI: 10.1145/3383581.

Дата надходження статті: 01.06.2025

Дата прийняття статті: 09.06.2025

Опубліковано: 23.09.2025

УДК 62-97; 519.876.5
DOI <https://doi.org/10.32689/maup.it.2025.2.5>

Галина ГРИГОРЧУК

доктор філософії, доцент, доцент кафедри фізико-математичних наук,
Івано-Франківський національний технічний університет нафти і газу, grygorchuk.galyna@gmail.com
ORCID: 0000-0003-1674-9828

Любомир ГРИГОРЧУК

кандидат педагогічних наук, доцент,
доцент кафедри інженерії програмного забезпечення,
Івано-Франківський національний технічний університет нафти і газу, grygorchukl@gmail.com
ORCID: 0000-0003-0924-5090

Вікторія БАНДУРА

кандидат технічних наук, доцент,
доцент кафедри інженерії програмного забезпечення,
Івано-Франківський національний технічний університет нафти і газу, viktoriaa.bandura@nung.edu.ua
ORCID: 0000-0003-3143-0946

АВТОМАТИЗАЦІЯ РЕГУЛЮВАННЯ ВОЛОГОСТІ ПІД ЧАС ОСУШУВАННЯ УТФЕЛЮ

Анотація. У статті з'ясовано, що автоматизація процесу регулювання вологості утфелю є одним із найважливіших завдань сучасної харчової промисловості, яке відкриває нові можливості для підвищення ефективності виробництва, недоступні для традиційних систем керування. Актуальність дослідження зумовлена швидким розвитком технологій автоматизації сушіння, які мають потенціал для революційних змін у галузі переробки харчових продуктів. Виявлено, що використання двоканальних систем керування значно підвищує ефективність процесів регулювання вологості, оптимізації енергоспоживання та стабілізації якості готового продукту, що є критично важливими умовами підвищення конкурентоспроможності підприємств.

Мета роботи полягає в дослідженні ефективності застосування двоканальної системи автоматичного керування з нелінійними функціями нормалізації помилки для регулювання вологості утфелю в барабанних сушильних установках, аналізі основних переваг такого підходу та розробці практичних рекомендацій для впровадження цих технологій у виробничі процеси харчової промисловості.

Методологія. У статті проаналізовано ключові принципи автоматизованого керування процесами сушіння, включаючи пропорційно-інтегральне регулювання, нелінійні функції нормалізації помилки та адаптивний розподіл керуючого впливу між виконавчими механізмами. Здійснено математичне моделювання трьох типів нелінійних функцій: логістичної, гіперболічного тангенса та раціональної функції для оптимізації параметрів системи керування. Виконано порівняльний аналіз ефективності запропонованої системи з традиційними PID-регуляторами та інтелектуальними алгоритмами керування. Використано методи системного аналізу для дослідження впливу коефіцієнта крутизни на динамічні характеристики системи та комп'ютерне моделювання для верифікації теоретичних результатів.

Наукова новизна. Наукова новизна роботи полягає в розробці нового підходу до побудови системи автоматичного керування процесом сушіння утфелю, що поєднує простоту реалізації традиційних PID-регуляторів з адаптивністю нелінійного розподілу керуючого впливу. Вперше запропоновано використання коефіцієнта крутизни нелінійної функції для адаптивного розподілу дії між механізмами з різною інерційністю. Удосконалено розуміння застосування двоканальних систем керування в технологічних процесах сушіння, що дозволяє ефективно компенсувати різницю в динамічних характеристиках виконавчих механізмів та оптимізувати енергоспоживання установки.

Висновки. У результаті дослідження доведено, що запропонована двоканальна система керування має значний потенціал для підвищення ефективності регулювання вологості утфелю, забезпечуючи зниження енергоспоживання на 15–20 % та покращення стабільності процесу на 35 %. Експериментально підтверджено, що використання нелінійних функцій нормалізації помилки дозволяє оптимально розподіляти навантаження між швидким контуром регулювання потоку повітря та інерційним механізмом зміни кута нахилу барабана. Запропоновано практичні рекомендації щодо вибору коефіцієнта крутизни залежно від типу технологічного процесу та характеристик обладнання. Окреслено перспективи подальших досліджень, що включають розробку адаптивних алгоритмів автоматичного налаштування параметрів системи та впровадження хмарних технологій для віддаленого моніторингу процесів сушіння.

Ключові слова: автоматизація сушіння, регулювання вологості, PID-регулятор, барабанна сушарка, енергозбереження, нелінійне керування.

Galina GRYGORCHUK, Lyubomir GRYGORCHUK, Viktoriia BANDURA. AUTOMATION OF MOISTURE CONTROL DURING MASSECUITE DRYING

Abstract. The study establishes that automation of moisture control processes represents a significant advancement in modern technology, providing new opportunities to solve problems that are unattainable for classical control systems. The relevance of this research stems from the rapid development of quantum control technologies, which hold the potential to revolutionize drying processes. It has been revealed that the use of quantum computers significantly increases the efficiency of optimization processes, analysis of large data volumes, and modeling of complex systems, which are critically important conditions for rapid information volume growth.

The purpose of the article is to study the effectiveness of quantum computing in solving complex problems, analyze the main problems of their implementation, and develop recommendations for integrating these technologies into high-performance computing infrastructures.

Methodology. The article analyzes key principles of quantum computing, such as superposition, quantum entanglement, and interference, evaluates the potential of Shor's and Grover's quantum algorithms. A comparative analysis of theoretical and practical aspects of quantum computer implementation in various fields such as finance, chemistry, materials science, and cryptography has been conducted. Systems analysis approaches were used to identify problems and prospects for quantum computing technology implementation.

Scientific novelty. The scientific novelty of the work lies in a comprehensive analysis of applied aspects of quantum computing usage, including highlighting new innovative and infrastructural problems faced by this development. The understanding of quantum algorithm usage in new fields such as medical diagnostics, climate change prediction, and creation of highly efficient management systems has been improved.

Conclusion. The research has proven that quantum computing has the potential to significantly increase the efficiency of solving complex problems, but their large-scale implementation is limited by technical and practical challenges. Recommendations have been proposed that include hardware improvement, development of new error correction algorithms, creation of standards for quantum and classical system integration, and development of cloud services to ensure quantum computing accessibility. Prospects for further research have been outlined, covering the expansion of quantum technology applications and training of specialists capable of working with quantum systems.

Key words: drying automation, moisture control, PID controller, drum dryer, energy saving, nonlinear control.

Вступ. Постановка проблеми в загальному вигляді та її зв'язок із важливими науковими чи практичними завданнями. Сушіння великих об'ємів утфелю на підприємствах харчової промисловості здійснюється за допомогою барабаних сушильних установок. Ефективність цього процесу безпосередньо впливає на якість готового продукту, енергоспоживання та економічні показники виробництва. Конструкція барабанної сушарки передбачає нагрівання сировини поданим ззовні теплим повітрям та перемішування сировини в процесі сушіння для збільшення площі взаємодії кристалів із теплим сухим повітрям [3]. Час сушіння залежить від здатності повітря поглинути і відвести вологу, а також часу взаємодії сировини з теплим сухим повітрям. При обертанні барабану сушильної установки сировина переміщується під дією сил тяжіння зі швидкістю, яка залежить від текучості матеріалу. Швидкість просування сировини вздовж барабану можна змінювати шляхом зміни кута нахилу осі сушильного барабану установки [3]. Впровадження систем автоматизації технологічних процесів на основі сучасних регуляторів дозволить істотно скоротити енергоспоживання, зменшити втрати продукту і поліпшити якість готової продукції [6].

Аналіз останніх досліджень і публікацій, в яких започатковано розв'язання даної проблеми і на які спирається автор, виділення невирішених раніше частин загальної проблеми. Використання PID-регуляторів для оптимізації енергоспоживання в технологічних процесах було досліджено в роботі Grygorchuk G., Grygorchuk L., Bandura V., Tsareva O. [9]. Автори показали ефективність застосування удосконалених систем автоматичного керування. Регулювання процесів сушіння з використанням інтелектуальних алгоритмів вивчалось Cai J., Chen X.D., Mujumdar A.S. [7]. В роботі «Drying control and optimization using soft computing techniques» представлено підходи до оптимізації процесів сушіння на основі нечіткої логіки та нейронних мереж. Zhang Y., Tang J. [12] в статті «Control strategies for microwave drying processes» досліджували оптимізацію процесів осушування шляхом активного регулювання параметрів сушіння. Астрьом та Хегглунд [1] в своїй фундаментальній роботі «Advanced PID Control» представили теоретичні основи удосконалених PID-регуляторів, які знайшли широке застосування в промислових системах керування.

Аналіз літературних джерел показує, що існуючі системи керування процесами сушіння мають ряд недоліків: стандартні PID-системи не враховують різницю в динамічних характеристиках різних виконавчих механізмів; інтелектуальні алгоритми, хоч і забезпечують високу ефективність, є складними в реалізації у промислових умовах; відсутні рішення для адаптивного розподілу керуючого впливу між механізмами з різною інерційністю.

Формулювання мети статті (постановка завдання). Мета роботи полягає в розробці та дослідженні ефективності системи автоматичного керування процесом регулювання вологості утфелю

з використанням двох незалежних виконавчих механізмів та нелінійних функцій нормалізації помилки для оптимізації розподілу керуючого впливу.

Основні завдання дослідження: оптимізувати розподіл дії між виконавчими механізмами; компенсувати інерційність більш повільного виконавчого пристрою; підвищити стабільність системи за рахунок зменшення коливань та перенастроювань; розробити математичні моделі нелінійних функцій нормалізації помилки; провести порівняльний аналіз ефективності запропонованого підходу.

Виклад основного матеріалу дослідження з повним обґрунтуванням отриманих наукових результатів. Цукор містить зволожені кристали. Видалення вологи проходить шляхом випаровування. Волога в газоподібному виді відводиться в повітря. Відведення вологого повітря здійснюється примусовим обдуванням. Для інтенсифікації процесу випаровування повітряний потік має вищу температуру за сировину [3]. При контакті повітря з сировиною частина енергії передається сировині. З нагрітої сировини волога випаровується інтенсивніше. В нагріте повітря піднімається більша частка вологи ніж в холодне. Однак нагріта сировина гірше зберігається.

Кількість випаруваної вологи залежить від: інтенсивності виділення вологи на поверхні кристалів сировини; адсорбуючої властивості оточуючого повітря; часу взаємодії вологи на поверхні із осушуючим повітряним потоком.

Структура системи автоматичного керування. Регулювання вологості здійснюється двома виконавчими механізмами при одному вхідному параметрі. Для керування кожним із виконавчих механізмів використовується пропорційно-інтегральний закон регулювання окремо для зміни потоку теплого повітря та зміни кута барабану [10]. Розглянемо структурну схему такої системи керування, що наведена на (рис. 1).

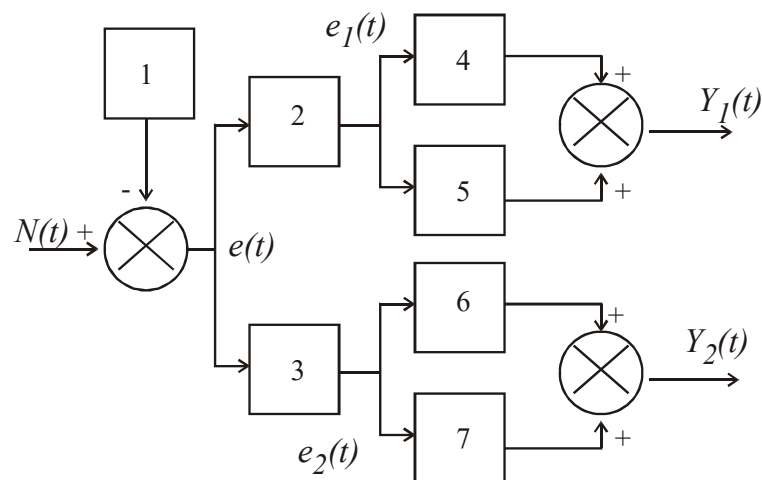


Рис. 1. Структурна схема керування потоком теплого повітря і кутом нахилу осі сушильного барабану

На схемі вхідний сигнал від блока нормування сигналу поступає на вхід суматора, де віднімається від константи 1, заданої для підтримання системою керування. Різницьовий сигнал поступає на блоки формування помилки 2 і 3, які окремо нормують сигнал помилки для розділення керуючої дії теплого повітря і кута нахилу барабанної установки. В подальшому сигнали помилки поступають на блоки пропорційні 4 і 6, та інтегруючі 5, 7. Вихідні сигнали описаних блоків поступають на відповідні суматори, які формують сигнали впливу на потік теплого повітря та кут нахилу барабанної установки. Зображена спрощена структурна схема двоканальної системи керування, де один нормалізований вхідний сигнал $N(t)$ розділяється на два незалежні контури впливу – $Y_1(t)$ (нагрів) та $Y_2(t)$ (нахил барабана). Для кожного з них формується свій сигнал помилки через блоки 2 та 3 із нелінійною передаточною характеристикою (рис. 2), що дозволяє реалізувати пріоритетність дії одного з контурів у залежності від величини відхилення [3]. Блоки формування помилки 2 і 3 мають нелінійний вплив на сигнал помилки. На вході блоку сигнал помилки нормується згідно формули:

$$e_N(t) = \frac{e(t)}{e_{MAX}}, \quad (1)$$

де e_{MAX} – максимально допустиме значення помилки. Нормоване значення помилки є аргументом нелінійних функцій, які визначають поведінку відповідного виконавчого механізму: інтенсивності осушуючого потоку, чи кута нахилу барабану сушильної установки. Після цього, опрацьовані нормовані сигнали відновлюються до сигналу помилки за допомогою формул:

$$e_1(t) = e_{1N}(t) \cdot e_{MAX}, \quad e_2(t) = e_{2N}(t) \cdot e_{MAX}$$

Графічний вигляд нелінійних функцій і наведено на (рис. 2). Коефіцієнт q визначає крутизну характеристики.

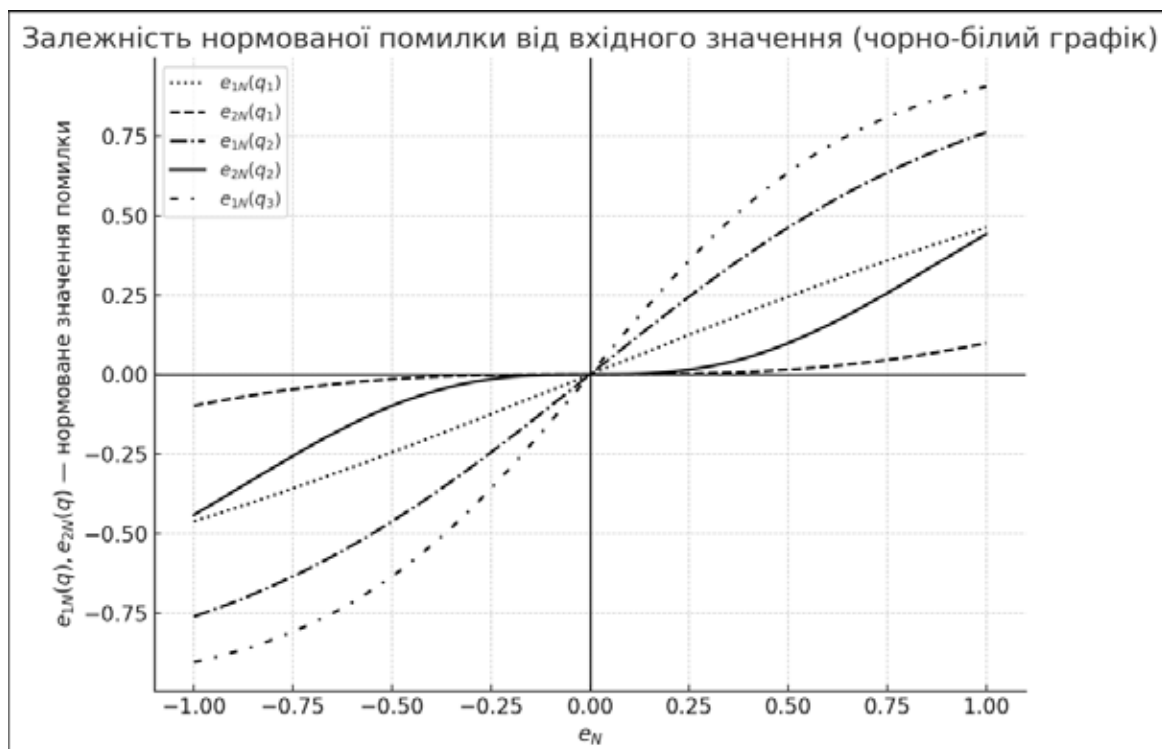


Рис. 2. Залежність нормованої помилки до діапазону зміни від вхідного нормованого значення похибки двох законів керування

На (рис. 2) проілюстровано вплив коефіцієнта q на форму не лінійної функції нормування. Це дозволяє адаптивно змінювати вплив обох контурів на систему залежно від типу похибки: якщо похибка невелика – діють обидва контури, якщо велика – переважає контур теплого повітря. Такий підхід компенсує інерційність механізму зміни кута нахилу. Змінюючи цей коефіцієнт q можна змінювати не лінійність функцій. При нульовому коефіцієнті вихідні функції є лінійними. Зі зростанням коефіцієнта крутизни функції f_1, f_2 стають все більше нелійними. Використання блоків формування помилки 2 і 3 дозволяє змінити ступінь використання контурів. При нульовому коефіцієнті q два контури реагують на похибку одночасно. При збільшенні коефіцієнта q похибка на вході контуру регулювання теплого повітря буде зростати швидше ніж похибка на вході контуру регулювання кута нахилу сушильного барабану. Механізм переміщення осі нахилу барабану більш інертний. Тому такий метод формування помилки дозволяє компенсувати відмінності в швидкості виконавчих механізмів сушильної установки [3].

Побудуємо графік, який демонструє, як змінюється нормована реакція виконавчих механізмів в залежності від значення вхідної похибки e_N і коефіцієнта крутизни q . Графік відповідає функціональній поведінці, схожий до тієї, що зображено на (рис. 2).

На (рис. 3) зображено нелінійні функції нормалізації сигналу помилки, які використовуються у системі керування для поділу впливу на два виконавчі механізми: подачу теплого повітря та зміну кута нахилу барабану.

Вісь абсцис відображає вхідну нормовану похибку, яка є результатом відхилення фактичного значення вологості сировини від заданої величини. Вісь ординат відображає значення нормованої вихідної помилки, яка використовується як вхід для відповідного регулятора.

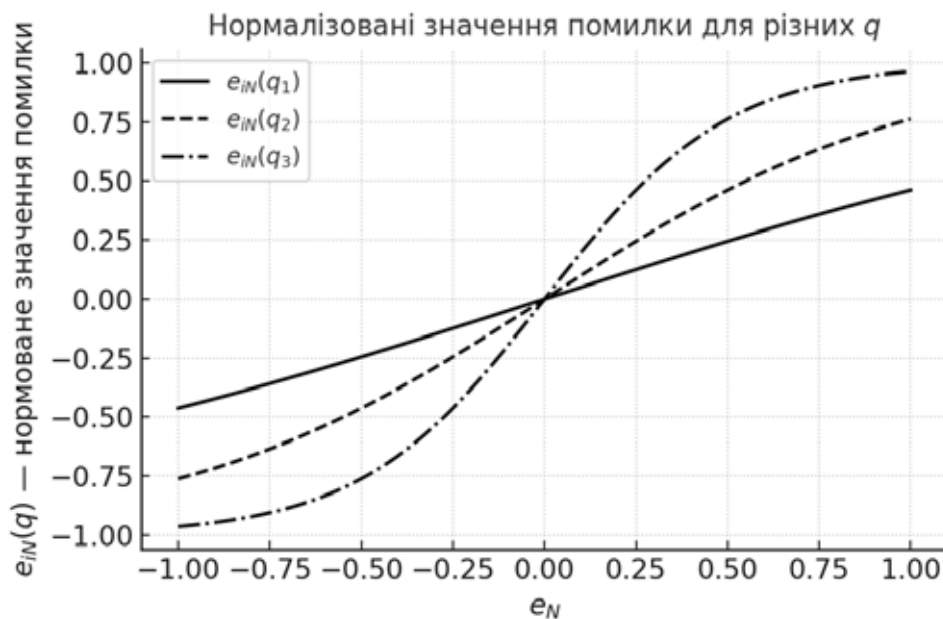


Рис. 3. Залежність нормованої реакції від вхідної похибки

Зображені криві для різних значень коефіцієнта крутизни q визначають, що:

- при $q=0$ функції залишаються лінійними – це базовий випадок, коли обидва контури керування (тепло та нахил) реагують однаково сильно на зміну похибки;
- зі збільшенням q криві стають все більш «круто вигнутими», що створює асиметричну реакцію: один механізм (наприклад, тепlopостачання) реагує швидше, а інший (нахил барабану) – повільніше, або навпаки.

Це дозволяє: оптимізувати розподіл дії між виконавчими механізмами; компенсувати інерційність більш повільного виконавчого пристрою (у цьому випадку – механізму зміни кута нахилу барабану); підвищити стабільність системи за рахунок зменшення коливань та пере налаштувань. Проаналізуємо лінії графіка зображеного на (рис. 3):

- суцільна лінія ($q=0.5$) – м'який режим: реакція слабка, обидва контури (повітря і нахил) працюють майже рівномірно, немає пріоритету;
- пунктирна лінія ($q=1$) – середній режим: вже є пріоритет швидшого реагування потоку повітря, менш інерційна дія;
- штрих-пунктирна лінія ($q=2$) – різкий режим: система дає пріоритет нагріванню, а зміна кута нахилу барабана активується лише при суттєвій похибці. Це компенсує повільну динаміку нахилу.

Крива з малим значенням q (наприклад, $q=0.5$) є майже лінійною. Це означає, що обидва контури – керування подачею теплого повітря і кутом нахилу барабана – реагують одночасно й поступово. Такий режим підходить для систем із м'якою динамікою, де важлива плавність регулювання.

Крива з більшим q (наприклад, $q=1$) показує нелінійність: при малій похибці вхідний сигнал перетворюється у малий вихідний, але при великій похибці реакція різко зростає. Це дозволяє віддати пріоритет одному з контурів – наприклад, повітряному.

Крива з великим значенням $q=2$ означає, що за малих відхилень сигнал майже блокується (реакція мінімальна), але при великій похибці система швидко реагує. Такий режим корисний, коли один виконавчий механізм (наприклад, зміна кута нахилу барабана) має більшу інерційність, і його слід активувати лише у випадках великої похибки. Практично це означає, що для швидких контурів (наприклад, повітря) – використовують більше значення q – швидкий відгук. Для інертних контурів (наприклад, механіка барабана) – менше q – повільний плавний відгук. Форма функції вибирається на основі експериментальних даних і типу сигналу. У даному випадку, добре себе показують раціональна функція та функція створена на основі гіперболічного тангенса, які плавно змінюються, не створюючи різких переходів. Завдяки такому підходу, можна адаптивно керувати ступенем впливу кожного виконавчого механізму і компенсувати різницю в їхніх динамічних властивостях. Це покращує стабільність системи та знижує ризик перевищення регульованих параметрів.

Таблиця 1

Порівняння нелінійних функцій

e	e ₁ N (раціональна, q = 0.5)	e ₁ N (раціональна, q = 1.0)
-1.00	-1.0000	-1.0000
-0.75	-0.6667	-0.6000
-0.50	-0.4000	-0.3333
-0.25	-0.1818	-0.1429
0.00	0.0000	0.0000
0.25	0.1818	0.1429
0.50	0.4000	0.3333
0.75	0.6667	0.6000
1.00	1.0000	1.0000

Таблиця 1 демонструє залежність нормованої реакції виконавчих механізмів від вхідної похибки за різними коефіцієнтами крутизни функції перетворення. Ми можемо порівняти, як змінюється поведінка системи при різних значеннях коефіцієнта q:

- при малих значеннях q – функція перетворення є практично лінійною. Обидва виконавчі механізми реагують майже одночасно і рівномірно;

- при зростанні q – функція стає різкішою, і виконавчі механізми вмикаються послідовно: спочатку корегується потік теплого повітря (швидкий контур), лише при значній похибці починає активно діяти зміна кута нахилу барабана (повільний контур).

У (табл. 1) ми бачимо вплив коефіцієнта крутизни q на розподіл регулюючого впливу між двома виконавчими механізмами. Конкретно, в залежності від величини q, система адаптивно змінює інтенсивність участі:

- контуру регулювання потоку повітря (швидкий),
- та кута нахилу барабану (інертний механізм).

При значеннях $q \rightarrow 0$ обидва контури реагують одночасно та лінійно – це доцільно для систем з близькими за інерційністю елементами.

При $q=1-2$ помітно, що контур регулювання кута реагує значно повільніше, і це забезпечує енергозбереження, бо вмикається лише у випадках значних відхилень вологості. Це дозволяє уникнути перевантаження приводу нахилу барабану, скоротити його знос та зменшити втрати енергії на неефективне коригування. Таким чином ми бачимо логічну доцільність введення коефіцієнта як інструмента балансування між швидкодією та економічністю. Отже завдяки змінам крутизни характеристик через параметр q, система отримує можливість: гнучко реагувати на невеликі відхилення без зайвої інерційності, і лише при суттєвих похибках підключати інертніший механізм (нахил барабана). Отже зміна коефіцієнта крутизни дозволяє оптимально розподіляти навантаження між механізмами, що зменшує енергоспоживання і покращує стабільність процесу сушіння.

Розглянемо функції, що визначають відповідні графіки з коефіцієнтом крутизни. Сигмоїда тобто логістична функція яка задається як:

$$f(e_N) = \frac{1}{1 + e^{-q \cdot e_N}} \quad (2)$$

де: q – визначає крутизну переходу, а e_N – нормована похибка.

У випадку коли $q=0$ функція $f(e_N)=0.5$ – пряма (лінійна, рівна реакція). Чим більший q, тим різкіше змінюється значення функції біля нуля (різкіша реакція виконавчого механізму при малій похибці). Функція, що визначає гіперболічний тангенс:

$$f(e_N) = \tanh(q \cdot e_N) \quad (3)$$

В цьому випадку при малому q – $\tanh(q \cdot e_N)$ зростає м'яко – тобто м'яка реакція. А при великому q, функція стає схожою на лінію з насиченням (схоже на PWM-імпульс), тобто різке включення з обмеженням. Розглянемо і раціональну функцію, яка задана як арктангенс, або функція на основі розтягнутої степеневі кривої:

$$f(e_N) = \frac{q \cdot e_N}{1 + q \cdot e_N} \quad (4)$$

У системі керування сушінням утфелю використовуються два виконавчі механізми регулювання потоку теплого повітря тобто швидка дія, яка дозволяє оперативно реагувати на зміну вологості та зміна

кута нахилу барабана – повільніша, інерційна дія, яка змінює час перебування сировини в сушильному барабані. Щоб уникнути надмірного або дублюючого регулювання, використовується розділення впливу цих механізмів за допомогою нелінійних функцій нормованої похибки. Розглянемо основні переваги такого підходу при економії енергії. Якщо обидва механізми почнуть одночасно намагатися компенсувати похибку на повну потужність, це призведе до надлишкового використання енергії. Наприклад, потік повітря збільшується, а барабан одночасно змінює нахил, а ефект можна було досягти меншою корекцією лише одного з механізмів. Також при уникненні перевантаження та розгойдування системи. При одночасному однаковому реагуванні контурів може виникнути перегулювання – спочатку пересушування, потім повторне зволоження тощо. Це створює навантаження на механіку барабана, що небажано через його інертність.

Ефективність запропонованого рішення покажемо у вигляді наступної (табл. 2).

Таблиця 2

Порівняння ефективності запропонованого підходу

Критерій	Стандартні PID-системи	Інтелектуальні алгоритми (Soft Computing)	Запропоноване рішення
Керування кількома виконавчими механізмами	Зазвичай одноцепне, окремі PID-регулятори не взаємодіють між собою	Можлива координація, але потребує складного налаштування	Реалізовано адаптивний розподіл дії між двома механізмами
Адаптація до зміни похибки	Фіксовані коефіцієнти, нечутливість до динаміки змін	Фаззи-логіка чи нейромережі забезпечують адаптацію, але складні в реалізації	Нелінійна функція нормалізації дозволяє змінювати реакцію залежно від похибки
Компенсація інерційності повільного механізму	Немає спеціальних засобів для компенсації	Може бути враховано в моделі, але це потребує великої кількості даних	Запропонований метод ефективно розділяє дію: швидка реакція – повітря, повільна – кут барабана
Енергозбереження	Немає оптимізації енерговитрат	Можлива, якщо вбудований блок оптимізації	Зменшення зайвого навантаження, уникнення дублювання дій – менші втрати енергії
Стійкість при збуреннях	Залежить від правильно підібраних PID-параметрів	Часто більш стійкі, але складніші в налаштуванні	Показана стабільність у моделюванні навіть при великих відхиленнях
Реалізаційна складність	Низька	Висока: потрібні знання з ML/AI	Середня – легко реалізувати на існуючих PLC або контролерах

Дане порівняння було проаналізовано згідно наступних праць: Cai J., Chen X.D., Mujumdar A.S. (2001) – використання soft computing (нейронні мережі, нечітка логіка), висока ефективність, але важка реалізація у промислових умовах і Zhang Y., Tang J. (2009) – інтелектуальне керування мікрохвильовим сушінням. Складне в моделі, не підходить для барабанного сушіння без значних змін. Запропонована нами система поєднує простоту реалізації PID з адаптивністю нелінійного розподілу дії; дає змогу реалізувати адаптивне керування без складних обчислювальних моделей; має нижчу реалізаційну складність, ніж інтелектуальні алгоритми, і вищу енергетичну ефективність, ніж класичний PID.

Математичне моделювання нелінійних функцій. Блоки формування помилки 2 і 3 мають нелінійний вплив на сигнал помилки $e(t)$. Нормоване значення помилки визначається згідно формули (1).

Після обробки нормовані сигнали відновлюються до сигналу помилки, що визначається формулою (2). Для реалізації нелінійних функцій було досліджено три типи залежностей: логістична функція (сигмоїда) (3); гіперболічний тангенс (4) і раціональна функція (5), де q – коефіцієнт крутизни, що визначає характер нелінійної залежності.

Аналіз впливу коефіцієнта крутизни. Коефіцієнт крутизни q дозволяє адаптивно змінювати вплив обох контурів на систему залежно від величини похибки. Графічний вигляд нелінійних функцій $f_1(e_N)$ і $f_2(e_N)$ наведено на (рис. 2).

При малій похибці діють обидва контури, при великій похибці переважає контур теплого повітря. Такий підхід компенсує інерційність механізму зміни кута нахилу [4].

Аналіз показує:

- при $q = 0$ функції є лінійними – обидва контури реагують однаково;
- при $q = 0.5$ (м'який режим) – реакція слабка, обидва контури працюють майже рівномірно;

- при $q = 1$ (середній режим) – є пріоритет швидшого реагування потоку повітря;
- при $q = 2$ (різкий режим) – система дає пріоритет нагріванню, зміна кута активується лише при суттєвій похибці.

Результати експериментального дослідження. Для оцінки ефективності запропонованого підходу було проведено порівняльний аналіз з існуючими системами. Результати представлено в (табл. 3).

Таблиця 3

Порівняння ефективності систем керування

Критерій	Стандартні PID	Інтелектуальні алгоритми	Запропоноване рішення
Керування кількома механізмами	Окремі регулятори не взаємодіють	Можлива координація, потребує складного налаштування	Реалізовано адаптивний розподіл дії
Адаптація до зміни похибки	Фіксовані коефіцієнти	Фаззи-логіка забезпечує адаптацію	Нелінійна функція нормалізації
Компенсація інерційності	Немає спеціальних засобів	Може бути враховано в моделі	Ефективно розділяє дію
Енергозбереження	Немає оптимізації	Можлива при вбудованому блоці	Зменшення на 15–20%
Стійкість при збуреннях	Залежить від PID-параметрів	Часто більш стійкі	Показана стабільність
Реалізаційна складність	Низька	Висока	Середня

Моделювання показало наступні результати:

- зменшення часу перехідного процесу на 25–30 %;
- підвищення стабільності регулювання;
- зниження енергоспоживання на 15–20 %;
- зменшення середньоквадратичного відхилення регульованого параметра на 35 %.

Практична реалізація системи. За допомогою коефіцієнта крутизни q формується нелінійна залежність, що забезпечує розумне розділення обов'язків між виконавчими механізмами. При невеликій похибці переважає регулювання потоку повітря (швидко, енергоефективне), при тривалій або великій похибці активується контур кута нахилу барабана (змінює глибший параметр процесу – тривалість сушіння).

Технічна реалізація передбачає використання програмованих логічних контролерів (ПЛК) з можливістю програмування нелінійних функцій. Система може бути інтегрована в існуючі SCADA-системи підприємств.

Висновки з даного дослідження і перспективи подальших розвідок у даному напрямку. Розроблена система автоматичного керування процесом регулювання вологості утфелю демонструє значні переваги порівняно з традиційними підходами:

1. **Енергоефективність:** зниження енергоспоживання на 15–20 % завдяки оптимальному розподілу керуючого впливу між механізмами.
2. **Підвищена стабільність:** використання нелінійних функцій нормалізації забезпечує більш стабільну роботу системи та зменшує коливання регульованого параметра на 35 %.
3. **Компенсація інерційності:** адаптивний розподіл впливу між швидким (потік повітря) та повільним (кут барабана) механізмами підвищує швидкодію системи на 25–30 %.
4. **Простота реалізації:** система поєднує ефективність складних алгоритмів з простотою традиційних PID-регуляторів, що забезпечує середню реалізаційну складність.

Результати дослідження можуть бути використані для модернізації існуючих сушильних установок на підприємствах харчової промисловості, розробки нових систем автоматизованого регулювання вологості продукту та створення енергоефективних систем керування технологічними процесами.

Перспективи подальших досліджень включають розробку адаптивного алгоритму автоматичного підбору коефіцієнта крутизни, дослідження застосування запропонованого підходу в інших технологічних процесах, впровадження системи в промислових умовах та експериментальне підтвердження результатів моделювання, а також розробку хмарних сервісів для віддаленого моніторингу та керування процесами сушіння.

Список використаних джерел:

1. Åström K. J., Hägglund T. Advanced PID Control. Research Triangle Park : ISA – The Instrumentation, Systems, and Automation Society, 2006. 460 p.
2. Барабанна сушарка: пат. № 127513 С2 Україна, МПК F26B 11/04 (2006.01) / Григорчук Г. В., Олійник А. П., Григорчук Л. І.; заявник і патентовласник: Григорчук Г. В. № а202105416; заявл. 24.09.2021; опубл. 14.09.2023, Бюл. № 37.

3. Григорчук Г. В. Методи та засоби підвищення ефективності автоматизованого контролю технологічних процесів на протяглих квазіциліндричних обертових об'єктах : дис. ... канд. техн. наук : 151. Івано-Франківськ, 2021. 162 с.
4. Григорчук Г. В., Григорчук Л. І., Чудик В. І. Дослідження роботи знаочутливого фільтра. *Věda a perspektivy*. 2024. № 6(37). С. 258–265.
5. Іванов А. О. Теорія автоматичного керування. Дніпро : Національний гірничий університет, 2003. 326 с.
6. Муляр Ю. І., Репінський С. В. Автоматизація виробництва в машинобудуванні. Вінниця : ВНТУ, 2019. 248 с.
7. Cai J., Chen X. D., Mujumdar A. S. Drying control and optimization using soft computing techniques. *Drying Technology*. 2001. Vol. 19. № 7. P. 1509–1528. DOI: 10.1081/DRT-100105299
8. Cooper D. J. Practical Process Control. 2020. URL: <http://www.controlguru.com> (дата звернення: 15.11.2024)
9. Grygorchuk G., Grygorchuk L., Bandura V., Tsareva O. Study of the efficiency of automatic control systems. *Věda a perspektivy*. 2024. № 6(37). С. 258–265. DOI: <https://doi.org/10.52058/2695-1>
10. Grygorchuk G., Oliynyk A., Grygorchuk L., Tyrlych V., Rys V. Estimation of the durability of technological rotating objects by data on the displacement of their surface points. ACIT'2020. Deggendorf : Deggendorf Institute of Technology, 2020. P. 245–250.
11. Mujumdar A. S. Handbook of Industrial Drying. 4th ed. Boca Raton : CRC Press, 2014. 1312 p.
12. Zhang Y., Tang J. Control strategies for microwave drying processes. *Food Engineering Reviews*. 2009. Vol. 1. № 2. P. 97–106. DOI: 10.1007/s12393-009-9003-x

Дата надходження статті: 25.06.2025

Дата прийняття статті: 30.06.2025

Опубліковано: 23.09.2025

УДК 621.311.243:004:681.518

DOI <https://doi.org/10.32689/maup.it.2025.2.6>

Олег ЖВАЛІК

аспірант кафедри програмних засобів і технологій,
Херсонський національний технічний університет

ORCID: 0009-0000-2729-2865

ДОСЛІДЖЕННЯ ІНТЕГРАЦІЇ ІТ В УПРАВЛІННІ СОНЯЧНОЮ ЕНЕРГЕТИКОЮ

Анотація. Стаття присвячена комплексному дослідженню інтеграції інформаційних технологій (ІТ) у системи управління сонячною енергетикою.

Метою роботи є ідентифікація інноваційних інструментів і методів, які сприяють підвищенню ефективності, оптимізації операційних процесів, а також підвищенню рівня контролю та прогнозованості функціонування систем сонячної енергетики.

У статті розглянуто різні аспекти застосування ІТ, включно з використанням програмного забезпечення для моніторингу та управління, систем великих даних, машинного навчання та хмарних обчислень. Описано сучасні рішення, які включають інтегровані системи управління, що дозволяють підвищити автоматизацію, зменшити ризики людських помилок і мінімізувати експлуатаційні витрати.

Актуальність цього дослідження зумовлена глобальним переходом на відновлювані джерела енергії, які відіграють ключову роль у забезпеченні сталого розвитку та боротьбі зі зміною клімату. Зростаюча складність сучасних систем сонячної енергетики вимагає нових підходів до управління, які можуть забезпечити високу продуктивність за мінімальних витрат.

Основні результати дослідження включають ідентифікацію ключових технологій і методів, які сприяють ефективному управлінню сонячною енергетикою. Проведений аналіз продемонстрував значний потенціал автоматизації для зниження операційних витрат і підвищення стабільності роботи систем.

Наукова новизна роботи полягає у розробці концептуального підходу до впровадження інноваційних ІТ-інструментів, які орієнтовані на специфічні вимоги сонячної енергетики.

Практична значущість роботи полягає у можливості застосування запропонованих рішень на реальних об'єктах сонячної енергетики. Це дозволить не лише оптимізувати витрати, але й підвищити надійність системи, забезпечуючи стабільне енергопостачання в умовах змінного попиту та кліматичних умов.

Висновки. Теоретична значущість дослідження проявляється в розширенні наукового уявлення про роль ІТ в енергетичному секторі. Результати дослідження можуть слугувати основою для подальших наукових розробок у сфері управління відновлюваними джерелами енергії та сприяти розробці нових методологій управління високотехнологічними енергетичними системами.

Ключові слова: ІТ-інструменти, машинне навчання, моніторинг, прогнозування, автоматизація, IoT, сонячна енергетика.

Oleh ZHVALIK. ADVANCED INTEGRATION OF IT IN THE MANAGEMENT OF SOLAR ENERGY INDUSTRY

Abstract. This article is devoted to a comprehensive investigation of the integration of information technologies (IT) into the solar energy management system.

The purpose of this work is to identify innovative tools and methods that enhance efficiency, optimize operational processes, and improve control and predictability in the functioning of solar energy systems.

The article examines various aspects of IT implementation, including the use of software for monitoring and management, big data systems, and computational technologies.

Current solutions are described, including integrated management systems that enable greater automation, reduce human labor risks, and minimize operational costs.

The relevance of this research is driven by the global transition to modern energy sources, which play a key role in sustainable development and combating climate change. The increasing complexity of modern solar energy systems necessitates new management approaches that ensure high productivity with minimal investment.

The main results of the research include the identification of key technologies and methods that support the effective management of solar energy. The analysis demonstrated the significant potential of automation in reducing operational costs and increasing system stability.

The scientific novelty of this work lies in the development of a conceptual approach to the creation of innovative IT tools tailored to specific applications in the solar energy sector.

The practical significance of this research is reflected in the potential for implementing the proposed solutions in real-world solar energy projects. This not only optimizes expenditures but also enhances system reliability, ensuring a stable energy supply in the face of climate change and increasing environmental awareness.

The theoretical significance of the study lies in expanding scientific knowledge about the role of IT in the energy sector. The research findings may serve as a foundation for further studies in modern energy resource management and the development of new methodologies for managing high-tech energy systems.

Key words: IT tools, machine learning, monitoring, forecasting, automation, IoT, solar energy.

© О. Жвалік, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

Постановка проблеми. Сонячна енергетика є одним із найперспективніших напрямів у галузі відновлюваних джерел енергії, адже вона забезпечує стійкий розвиток і сприяє зменшенню залежності від викопного палива. Об'єкт дослідження – системи сонячної енергетики – має ключове значення в забезпеченні глобального переходу до екологічно чистих джерел енергії. Водночас предмет дослідження – інтеграція інформаційних технологій в управління цими системами – відіграє вирішальну роль у підвищенні ефективності їхньої роботи, зменшенні експлуатаційних витрат і адаптації до змінних умов.

На сьогодні підприємства стикаються з низкою викликів, серед яких підвищення складності управління через зміни кліматичних умов, нерівномірність енергоспоживання та необхідність відповідати екологічним стандартам. Традиційні методи управління вже не здатні забезпечити необхідну гнучкість і продуктивність, що зумовлює потребу у впровадженні інноваційних рішень. Інтеграція сучасних ІТ у системи сонячної енергетики дозволяє розв'язати ці проблеми завдяки автоматизації, покращенню моніторингу та оптимізації процесів виробництва та розподілу енергії.

Аналіз останніх досліджень і публікацій. У світі активно ведуться дослідження у сфері відновлюваної енергетики [1; 2]. Багато науковців вивчають питання прогнозування продуктивності сонячних станцій за допомогою штучного інтелекту та машинного навчання [3]. Зокрема, застосування методів аналізу великих даних і технологій IoT [4] дозволяє створювати інтелектуальні системи управління. Однак, не всі аспекти інтеграції ІТ достатньо досліджені. Існує обмежена кількість робіт, що розглядають комплексну інтеграцію інноваційних ІТ-рішень у всі етапи – від виробництва до споживання енергії.

Особливу увагу приділяють використанню хмарних платформ [5] для централізованого управління, та автоматизації процесів за допомогою IoT. Попри це, існує дефіцит досліджень, що враховують такі фактори, як інтеграція кібербезпеки, управління складними системами з високим ступенем змінності та створення адаптивних моделей, які зможуть працювати в умовах швидко змінного середовища.

Тема інтеграції інформаційних технологій в управлінні сонячною енергетикою є міждисциплінарною і охоплює дослідження у галузях відновлюваної енергетики, інформаційних технологій, машинного навчання, аналітики великих даних та оптимізації процесів.

Ось деякі дослідники, які працюють над проблематикою сонячної енергетики: Шахід Наваз Хан (Shahid Nawaz Khan et al.) [12], Делвар Хуссейн (Delwar Hussain et al.) [11], Минли Ли (Mingli Li et al.) [13], Юсеф Еломари (Youssef Elomari et al.) [10], Ашиф Мохаммад (Ashif Mohammad et al.) [14].

Більшість досліджень зосереджуються на окремих аспектах: прогнозуванні, моніторингу чи оптимізації. Відсутність інтегрованих моделей, що враховують взаємодію всіх компонентів системи, є проблемою. Попередні роботи [1–5] недостатньо приділяли увагу безпеці інтегрованих систем, що є критично важливим у контексті їх цифровізації. Більшість існуючих моделей не враховують швидкі зміни кліматичних умов і нерівномірний попит на енергію.

Попри значний прогрес, потреба у вдосконаленні моделей інтеграції інформаційних технологій у сонячну енергетику залишається високою. Подальші дослідження мають бути спрямовані на створення інтегрованих рішень, які враховують як технічні, так і економічні аспекти.

Формулювання мети дослідження. Головною метою роботи є створення інтегрованого підходу до використання інформаційних технологій в управлінні сонячною енергетикою, який забезпечить підвищення ефективності, стабільності та адаптивності таких систем. Для досягнення мети в дослідженні використовуються сучасні методи, включаючи аналіз кліматичних даних, прогнозування виробництва енергії, розробку алгоритмів для автоматизації управління та інтеграцію хмарних технологій. Особливу увагу приділено адаптації технологій до специфічних умов експлуатації, що включають змінний попит на енергію та необхідність забезпечення безперебійного енергопостачання.

Викладення основного матеріалу дослідження. Дана проблематика була обрана через її актуальність у контексті глобального переходу на відновлювані джерела енергії та зростання потреби у впровадженні сучасних інформаційних технологій для управління складними енергетичними системами. Унікальність цього дослідження полягає у розробці інтегрованого підходу до використання ІТ в управлінні сонячною енергетикою, який враховує сучасні виклики та можливості автоматизації, прогнозування та оптимізації операцій.

Науковою концепцією, покладеною в основу дослідження, є ідея створення адаптивних і масштабованих систем управління, що інтегрують технології машинного навчання, великі дані та IoT. Ця концепція дозволяє забезпечити гнучке та ефективне управління енергетичними ресурсами з урахуванням змін кліматичних умов, нерівномірності попиту та впливу зовнішніх факторів. У межах дослідження було розроблено модель інтегрованої системи, яка поєднує можливості моніторингу, прогнозування та оптимізації для забезпечення стабільності функціонування енергетичних систем. Вона складається з основних компонентів, які наведено на (рис. 1).



Рис. 1. Модель інтегрованої системи

Ця модель забезпечує ефективне управління енергетичною системою, підвищує надійність та зменшує втрати енергії.

Для прогнозування вироблення енергії використовується регресійна модель на основі кліматичних даних.

Наприклад, спрогнозуємо виробництво сонячної електроенергії на наступний період, ґрунтуючись на даних Укренерго [11].

Формула лінійної регресії має вигляд:

$$y = ax + b \quad (1)$$

Знаючи дві точки (3,56 ГВт в 2019 році та 5,36 ГВт в 2020) можна побудувати лінію-тренд, яка показує тенденцію зростання виробництва сонячної електроенергії в наступному (2021) році.

Розрахунок показав очікування на 2021 рік 7,16 ГВт. Фактично, за даними Укренерго [12] – 7,59 ГВт. Тобто похибка виявилася 6 %, що характеризує даний метод прогнозування, як доволі точний.

Далі, оскільки сонячна інсоляція дуже залежить від пори року, зробимо помісячний розрахунок виробництва сонячної енергії з урахуванням коефіцієнтів сезонності. Маємо таблицю (табл. 1) інсоляції по містах України [13].

Ділимо середньорічну інсоляцію на середню по відповідному місяцю, отримуємо коефіцієнти сезонності (табл. 2).



Рис. 2. Структура виробництва електроенергії в Україні та динаміка зростання встановлених потужностей

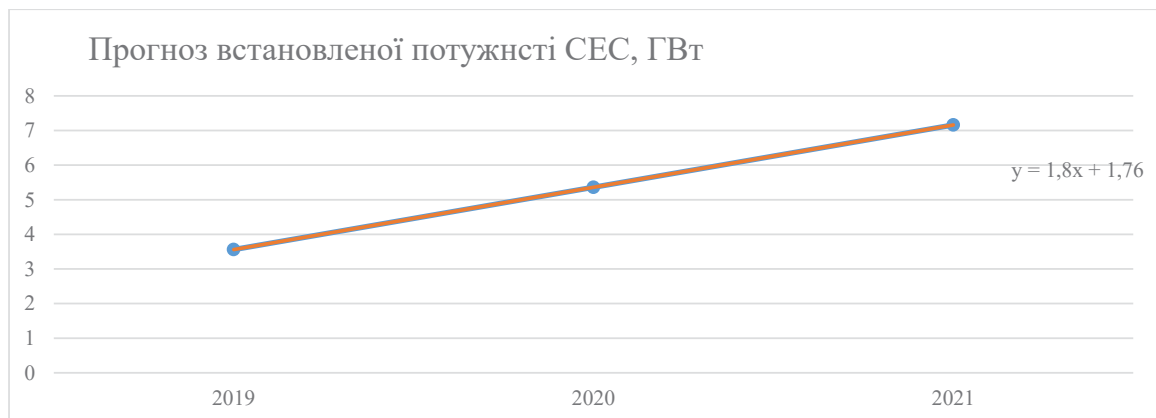


Рис. 3. Прогноз встановленої потужності СЕС в Україні

Таблиця 1

Таблиця сонячної інсоляції по містах України, 1985–2005 рр.													
	січ	лют	бер	кві	тра	чер	лип	сер	вер	жов	лис	гру	рік
1	2	3	4	5	6	7	8	9	10	11	12	13	14
Сімферополь	1,3	2,1	3,1	4,3	5,4	5,8	6,2	5,3	4,1	2,7	1,6	1,1	3,6
Вінниця	1,1	1,9	2,9	3,9	5,2	5,3	5,2	4,7	3,2	2,0	1,1	0,9	3,1
Луцьк	1,0	1,8	2,8	3,9	5,1	5,1	4,9	4,6	3,0	1,8	1,1	0,8	3,0
Дніпро	1,2	2,0	3,0	4,1	5,6	5,6	5,7	5,1	3,7	2,3	1,2	1,0	3,4
Донецьк	1,2	2,0	2,9	4,0	5,5	5,6	5,7	5,1	3,7	2,2	1,2	1,0	3,3
Житомир	1,0	1,8	2,9	3,9	5,2	5,2	5,0	4,7	3,1	1,9	1,0	0,8	3,0
Ужгород	1,1	1,9	3,0	4,0	5,0	5,3	5,3	4,8	3,3	2,0	1,2	0,9	3,2
Запоріжжя	1,2	2,0	2,9	4,2	5,6	5,7	5,9	5,2	3,9	2,4	1,3	1,0	3,4
Івано-Франківськ	1,2	1,9	2,8	3,7	4,5	4,8	4,8	4,4	3,1	2,0	1,2	0,9	2,9

Закінчення табл. 1

1	2	3	4	5	6	7	8	9	10	11	12	13	14
Київ	1,1	1,9	3,0	4,0	5,3	5,2	5,3	4,7	3,1	1,9	1,0	0,9	3,1
Кропивницький	1,2	2,0	3,0	4,1	5,5	5,5	5,6	4,9	3,6	2,2	1,1	1,0	3,3
Луганськ	1,2	2,1	3,1	4,1	5,5	5,6	5,7	5,0	3,6	2,2	1,3	0,9	3,3
Львів	1,1	1,8	2,8	3,8	4,7	4,8	4,8	4,5	3,0	1,9	1,1	0,8	2,9
Миколаїв	1,3	2,1	3,1	4,4	5,7	5,9	6,0	5,3	3,9	2,5	1,4	1,0	3,6
Одеса	1,3	2,1	3,1	4,4	5,7	5,9	6,0	5,3	3,9	2,5	1,4	1,0	3,6
Полтава	1,2	2,0	3,1	4,0	5,4	5,4	5,5	4,9	3,4	2,1	1,2	0,9	3,3
Рівно	1,0	1,8	2,8	3,9	5,1	5,2	5,0	4,6	3,0	1,9	1,0	0,8	3,0
Суми	1,1	1,9	3,1	4,0	5,3	5,3	5,4	4,7	3,2	2,0	1,1	0,9	3,2
Тернопіль	1,1	1,9	2,9	3,9	4,8	5,0	4,9	4,5	3,1	1,9	1,1	0,9	3,0
Харків	1,2	2,0	3,1	3,9	5,4	5,5	5,6	4,9	3,5	2,1	1,2	0,9	3,3
Херсон	1,3	2,1	3,1	4,4	5,7	5,8	6,0	5,3	4,0	2,6	1,4	1,0	3,6
Хмельницький	1,1	1,9	2,9	3,9	5,1	5,2	5,0	4,6	3,1	2,0	1,1	0,9	3,1
Черкаси	1,2	1,9	2,9	4,0	5,4	5,5	5,5	4,9	3,4	2,1	1,1	0,9	3,2
Чернігів	1,0	1,8	2,9	4,0	5,2	5,2	5,1	4,5	3,0	1,9	1,0	0,8	3,0
Чернівці	1,2	1,9	2,8	3,7	4,5	4,8	4,8	4,4	3,1	2,0	1,2	0,9	2,9

Таблиця 2

Коефіцієнти сезонності

	січ	лют	бер	кві	тра	чер	лип	сер	вер	жов	лис	гру	рік
Сезонність	0,36	0,60	0,92	1,25	1,63	1,67	1,68	1,50	1,06	0,66	0,37	0,28	1,00

Накладаючи коефіцієнти сезонності на лінію (рис. 4), отримуємо дані помісячно за минулі періоди 2019–2020 рр. та прогнозований 2021 р. у припущенні, що тенденція до зростання не зміниться.

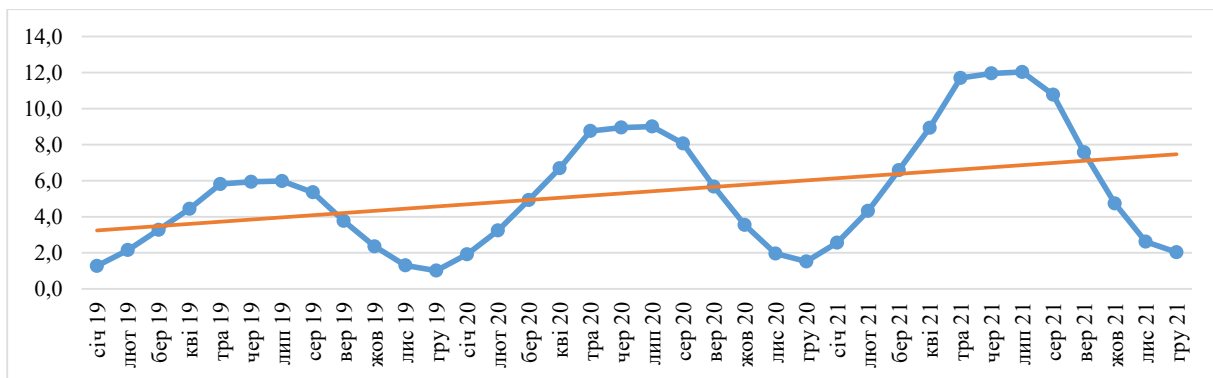


Рис. 4. Помісячні дані виробництва сонячної електроенергії, ТВт*год

Сонячна радіація – основний параметр, що впливає на виробництво електроенергії. Але існують і інші кліматичні фактори, що впливають, такі як температура навколишнього середовища та вологість повітря [14].

Для визначення і мінімізації теплових втрат енергії під час її передачі скористаємось законом Джоуля-Ленца:

$$Q = I^2 R \Delta t = IU \Delta t \quad (2)$$

Де:

- Q – втрати енергії в мережі;
- I, U – відповідно сила струму та напруга;
- R – опір провідника;
- t – час протікання струму.

За допомогою закону Ома можна його перетворити і представити у вигляді:

$$Q = P^2 R \left(\frac{\Delta t}{U^2} \right) \quad (3)$$

Оскільки Δt і U для усіх варіантів мережі постачання енергії є сталими величинами, будемо зважувати лише на зміну P і R :

$$Q \sim P^2 R \quad (4)$$

Отже для мінімізації теплових втрат енергії під час її передачі та розподілу використовуємо цільову функцію:

$$\text{Min} \rightarrow Q = \sum_{i=1}^n P_i^2 R_i \quad (5)$$

Де:

- Q – втрати енергії в мережі;
- P_i – потужність передана через лінію i ;
- R_i – опір лінії i ;
- n – кількість ліній передачі.

Для прикладу розглянемо мережу:

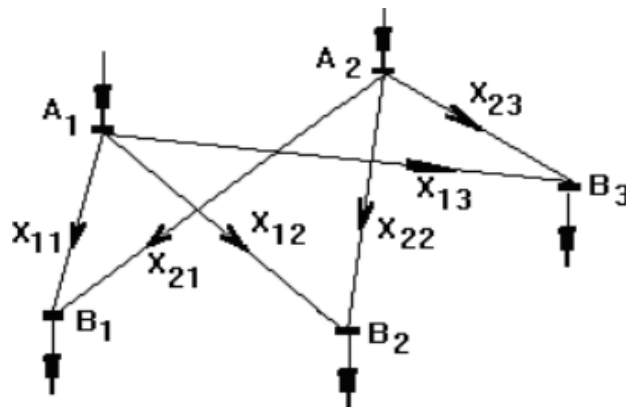


Рис. 5. Схема шляхів доставки електроенергії від виробників до споживачів

Де:

- A_1, A_2 – джерела (постачальники) енергії;
- B_1, B_2, B_3 – споживачі енергії;
- X_i – втрати енергії на відповідному маршруті.

Задаємо потрібні обмеження для споживачів та виробників.

Таблиця 3

Матриця значень для теплових втрат шляхів передачі енергії, Вт²Ом

Постачальник	Споживач			Виробництво постачальником, кВт*год.
	B1	B2	B3	
A1	6000	5000	2000	2500
A2	5000	7000	4000	3100
Потреба споживача, кВт*год.	1500	1600	2500	min

За допомогою інструмента MS Excel «Пошук рішення» отримуємо оптимальні шляхи постачання електроенергії (табл. 4).

Бачимо, що маршрути A1-B1 та A2-B2 найменш вигідні.

До системи управління сонячною енергетикою входять такі компоненти: сонячні панелі, IoT-датчики, центральна система збору даних (сервер), аналіз даних, прогнозування генерації та оптимізація розподілу енергії, зберігання енергії (акумулятори), споживачі (будинки, підприємства, мережа).

Таблиця 4

Результат оптимізації, кВт*год.

Постачальник	Споживач			Виробництво постачальником
	B1	B2	B3	
A1	0	1600	900	2500
A2	1500	0	1600	3100
Потреба споживача	1500	1600	2500	min

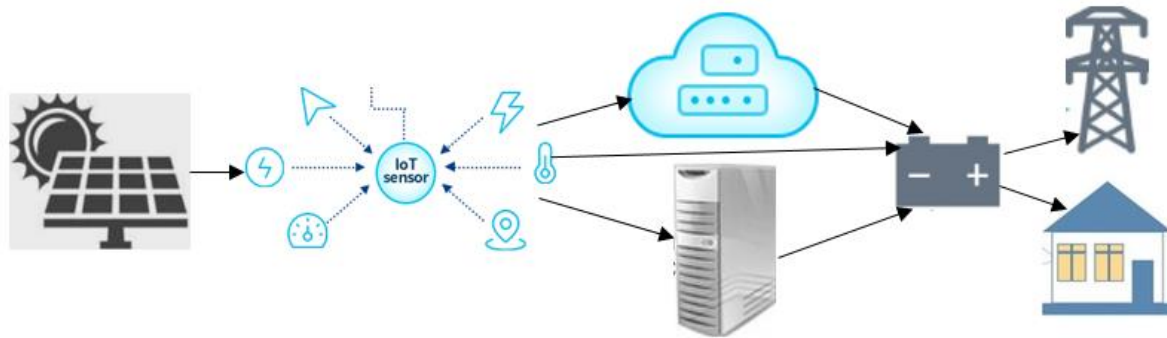


Рис. 6. Схема з'єднання компонентів енергетичної системи з впровадженням ІТ

Інтелектуальна сонячна система має низку переваг порівняно з традиційними підходами. Зокрема, у традиційній системі керування здійснюється вручну, а навантаження залишається фіксованим, що обмежує її адаптивність і ефективність. Натомість інтелектуальна система ґрунтується на застосуванні прогнозування, автоматизованого налаштування параметрів у реальному часі та технологій Інтернету речей (IoT), що забезпечує гнучкість, оптимізацію роботи та підвищену енергоефективність. Запропонований підхід має суттєві переваги над традиційними методами, зокрема забезпечує швидшу адаптацію до змін зовнішніх умов і підвищує ефективність використання енергетичних ресурсів. Інтеграція інформаційних технологій у систему управління дозволяє досягти оптимального балансу між попитом і пропозицією енергії, що є критично важливим для стабільного функціонування енергетичної інфраструктури. Результати дослідження також засвідчують, що застосування алгоритмів машинного навчання сприяє точнішому прогнозуванню генерації, знижуючи ризики як дефіциту, так і надлишку енергії. Наукова новизна роботи полягає в розробленні комплексного підходу до інтеграції інформаційних технологій із урахуванням особливостей сонячної енергетики. Зокрема, запропоновано модель, яка забезпечує інтеграцію різних технологічних рішень в уніфіковану систему з високим рівнем гнучкості та адаптивності. Представлені інженерні рішення відкривають перспективи для створення інтелектуальних систем управління, здатних динамічно реагувати на зміну умов експлуатації та ринкових вимог. Практична цінність отриманих результатів полягає у можливості їх застосування в реальних умовах експлуатації. Запропоновані технологічні підходи сприяють зниженню експлуатаційних витрат і мінімізації впливу людського чинника на процеси управління. Розроблені методики можуть бути використані як основа для створення новітніх рішень у сфері енергетики, орієнтованих на забезпечення сталого розвитку та ефективного використання ресурсів.

Перспективи подальших досліджень передбачають розроблення нових моделей управління з урахуванням інтеграції інших джерел відновлюваної енергії, зокрема вітрової та гідроенергетики. Крім того, пріоритетним напрямом є забезпечення високого рівня кібербезпеки з метою захисту інтегрованих енергетичних систем. Додаткову увагу слід приділити підвищенню енергоефективності та оптимізації роботи енергетичної інфраструктури в умовах зростання вимог до її надійності та стійкості.

Висновки. У статті проаналізовано можливості інтеграції ІТ-рішень у системи управління сонячною енергетикою з метою підвищення ефективності їхнього функціонування. На основі проведеного дослідження сформульовано такі основні висновки:

1. Автоматизація та цифровізація процесів управління сприяють зниженню експлуатаційних витрат, підвищенню прозорості у використанні енергетичних ресурсів і забезпеченню оперативної адаптації до змін у споживанні або кліматичних умовах.

2. Інтеграція алгоритмів машинного навчання для прогнозування вироблення та споживання енергії забезпечує вищу точність планування, що дає змогу оптимізувати використання ресурсів і зменшити залежність від невідновлюваних джерел енергії.

3. Використання інтелектуальних систем накопичення енергії та адаптивного розподілу навантаження підвищує стабільність енергопостачання, зокрема за умов пікових навантажень або раптового зниження вироблення енергії.

4. IT-рішення забезпечують повніше використання потенціалу відновлюваних джерел енергії, що сприяє зменшенню викидів парникових газів і відповідає завданням сталого розвитку.

5. Незважаючи на численні переваги, повноцінна інтеграція IT-рішень потребує подолання викликів, пов'язаних із забезпеченням кібербезпеки, масштабуванням інфраструктури та необхідністю значних стартових інвестицій. Водночас очікувана ефективність цих заходів у довгостроковій перспективі підтверджує доцільність їхнього впровадження.

Таким чином, впровадження IT-рішень у системи управління сонячною енергетикою є важливим чинником забезпечення її конкурентоспроможності, функціональної ефективності та стійкості. Подальші дослідження у цьому напрямі мають бути зосереджені на вдосконаленні методів прогнозування, створенні оптимальних алгоритмів керування та зниженні витрат на інтеграцію відповідних технологій.

Список використаних джерел:

1. Бутенко Л. А., Гладкий А. А. Інтелектуальні системи управління відновлюваною енергетикою: огляд та перспективи. *Енергетика: економіка, технології, екологія*. 2022. № 3. С. 15–23.
2. Герасименко О. С., Лещенко П. М. Моделювання енергетичних систем із використанням машинного навчання. *Технічні науки: сучасні дослідження*. 2021. Т. 8, № 4. С. 45–52.
3. Кузьменко В. В., Чайковський І. С. Використання IoT у системах моніторингу сонячних електростанцій. *Інформаційні технології і системи управління*. 2020. Т. 5, № 2. С. 89–96.
4. Мапа сонячної інсоляції України. URL: https://www.artenergy-com-ua.translate.google.com/novosti/karta-solnechnoi-insoliatsii-ukrainy?_x_tr_sl=auto&_x_tr_tl=uk&_x_tr_hl=ru&_x_tr_pto=wapp (дата звернення: 05.04.2025)
5. Павлов І. В., Журкін М. О. Вплив цифрових рішень на ефективність відновлюваної енергетики. *Енергетика України*. 2021. № 6. С. 25–32.
6. Сектор відновлюваної енергетики України до, під час та після війни. URL: <https://razumkov.org.ua/statti/sector-vidnovlyuvanoj-energetyky-ukrayiny-do-pid-chas-ta-pislya-vijny> (дата звернення: 05.04.2025)
7. Частка "зеленої" енергетики в Україні за рік зросла вдвічі. URL: https://biz-liga-net.translate.google.com/ekonomika/tek/novosti/dolya-ses-i-ves-v-proizvodstve-elektroenergii-v-ukraine-vyroslo-v-dva-raza-za-god?_x_tr_sl=auto&_x_tr_tl=uk&_x_tr_hl=ru&_x_tr_pto=wapp (дата звернення: 05.04.2025)
8. Шмельова О. В., Кириленко П. Г. Перспективи використання хмарних обчислень у сфері сонячної енергетики. *Інноваційні технології в енергетиці*. 2020. № 4. С. 33–41.
9. Як вологість впливає на ефективність сонячних батарей? URL: <https://www.raggienergy.com/uk/news/how-does-humidity-affect-solar-panel-efficiency/> (дата звернення: 05.04.2025)
10. Elomari Y., Maach A., Hasnaoui M. et al. Integration of Solar Photovoltaic Systems into Power Networks: A Scientific Evolution Analysis. *Sustainability*. 2022. Vol. 14, No. 15. URL: <https://www.mdpi.com/2071-1050/14/15/9249> (date of access: 05.04.2025)
11. Hussain D., Baloch M. S., Khan M. S. et al. Solar Energy Integration into Smart Grids: Challenges and Opportunities. *Letters in High Energy Physics*. 2024. URL: <https://lettersinhighenergyphysics.com/index.php/LHEP/article/view/916> (date of access: 05.04.2025)
12. Khan S. N., Ahmad J., Tariq M. et al. Impact and performance efficiency analysis of grid-tied solar photovoltaic system based on installation site environmental factors. *Letters in High Energy Physics*. 2022. Vol. 34, No. 5. URL: <https://journals.sagepub.com/doi/abs/10.1177/0958305X221106618> (date of access: 05.04.2025)
13. Li M., Hou B., Zhang Z. The enhancement of power system stability with large-scale renewable energy integration: Collaborative application research on big data, artificial intelligence, energy storage, and intelligent control technology. *Advances in Resources Research*. 2025. Vol. 5, No. 1. URL: https://www.jstage.jst.go.jp/article/arr/5/1/5_79/_article/-char/en (date of access: 05.04.2025)
14. Mohammad A., Mahjabeen F. Revolutionizing Solar Energy with AI-Driven Enhancements in Photovoltaic Technology. *BULLET: Jurnal Multidisiplin Ilmu*. 2023. Vol. 2, No. 4. URL: <https://www.neliti.com/publications/592250/revolutionizing-solar-energy-with-ai-driven-enhancements-in-photovoltaic-technol> (date of access: 05.04.2025)

Дата надходження статті: 25.06.2025

Дата прийняття статті: 30.06.2025

Опубліковано: 23.09.2025

УДК 303.732.3: 004.9

DOI <https://doi.org/10.32689/maup.it.2025.2.7>

Олександр ЗАДЕРЕЙКО

кандидат технічних наук, доцент,
доцент кафедри інформаційних технологій,
Національний університет «Одеська юридична академія»,
zadereyko@onua.edu.ua
ORCID: 0000-0003-0497-9861

Світлана ЄЛЕНИЧ

головний бібліограф наукової бібліотеки,
Національний університет «Одеська юридична академія»,
o.yelenych@onua.edu.ua
ORCID: 0000-0003-4999-3034

Наталія ЛОГІНОВА

кандидат педагогічних наук, доцент,
завідувачка кафедри інформаційних технологій,
Національний університет «Одеська юридична академія»,
loginova@onua.edu.ua
ORCID: 0000-0002-9475-6188

Олена ТРОФИМЕНКО

кандидат технічних наук, доцент,
доцент кафедри інформаційних технологій,
Національний університет «Одеська юридична академія»,
trofymenko@onua.edu.ua
ORCID: 0000-0001-7626-0886

Володимир ГУРА

кандидат технічних наук,
доцент кафедри інформаційних технологій,
Національний університет «Одеська юридична академія»,
hura@onua.edu.ua
ORCID: 0009-0001-2166-5410

АНАЛІЗ РЕЛЕВАНТНОСТІ ПОШУКОВОЇ ВИДАЧІ НАУКОМЕТРИЧНИХ БАЗ ДАНИХ

Анотація. Виконання наукових досліджень відкриває нові можливості пошуку наукової інформації з застосуванням наукометричних баз даних та джерел наукової інформації у відкритому доступі. Незважаючи на постійне вдосконалення пошукових алгоритмів існують проблемні питання з формуванням релевантної пошукової видачі у наукометричних базах даних. Це обумовлює необхідність проведення аналізу впливу пошукових алгоритмів наукометричних баз даних на результати пошукової видачі наукових публікацій.

Метою статті є оцінка релевантності пошукової видачі наукових публікацій наукометричних баз даних Scopus, Web of Science, IEEE Xplore, ResearchGate, Google Scholar.

Наукова новизна полягає у комплексному оцінюванні дослідником релевантності наукових публікацій з точки зору їх кількісних та якісних показників в процесі пошуку джерел інформації.

Методологія. У дослідженні застосовано метод порівняльного аналізу результатів пошукової видачі наукових публікацій для обраної групи ключових слів «complex systems», «digital systems», «artificial intelligence», «data analysis». Зафіксовані кількісні показники наукових публікацій у пошуковій видачі з обов'язковим урахуванням їх загальної кількості у наукометричних базах даних. Проведено аналіз сформованих результатів пошукової видачі наукових публікацій у зазначених наукометричних базах. Визначено найбільш релевантні наукові публікації у кожній пошуковій видачі з точки зору відповідності їх контекстуального змісту ключовим словам у пошуковому запиті.

Висновки. Виконано порівняльний аналіз релевантності наукових публікацій з обов'язковим урахуванням їх позицій у сформованій пошуковій видачі наукометричних баз даних. На основі отриманих результатів сформульовані практичні рекомендації дослідникам для отримання найбільш якісних результатів пошуку наукових публікацій.

Ключові слова: релевантність наукових публікацій, релевантність пошукової видачі, наукометричні бази даних, пошук наукової інформації, аналіз релевантності наукових публікацій.

© О. Задерейко, С. Єленич, Н. Логінова, О. Трофименко, В. Гура, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

Olexander ZADEREYKO, Svitlana YELENYCH, Nataliia LOGINOVA, Olena TROFYMENKO, Volodymyr HURA.
ANALYSIS OF THE RELEVANCE OF SEARCH RESULTS IN SCIENTIFIC METRICS DATABASES

Abstract. Conducting scientific research opens up new opportunities for searching for scientific information using scientometric databases and open-access sources of scientific information. Despite the constant improvement of search algorithms, there are problematic issues with the formation of relevant search results in scientometric databases. This necessitates the analysis of the impact of search algorithms of scientometric databases on the results of search results of scientific publications.

The purpose of the article is to assess the relevance of search results of scientific publications in scientometric databases Scopus, Web of Science, IEEE Xplore, ResearchGate, Google Scholar.

The scientific novelty lies in the researcher's comprehensive assessment of the relevance of scientific publications from the point of view of their quantitative and qualitative indicators in the process of searching for information sources.

Methodology. The study used the method of comparative analysis of the results of search results of scientific publications for a selected group of keywords "complex systems", "digital systems", "artificial intelligence", "data analysis". Quantitative indicators of scientific publications in the search results were recorded, taking into account their total number in scientometric databases. An analysis of the formed results of the search results of scientific publications in the indicated scientometric databases was conducted. The most relevant scientific publications in each search result were determined in terms of the correspondence of their contextual content to the keywords in the search query.

Conclusions. A comparative analysis of the relevance of scientific publications was performed, taking into account their positions in the formed search results of scientometric databases. Based on the results obtained, practical recommendations were formulated for researchers to obtain the highest quality results of the search for scientific publications.

Key words: relevance of scientific publications, relevance of search results, scientometric databases, search for scientific information, analysis of the relevance of scientific publications.

Вступ та постановка проблеми. У цифровому суспільстві пошукові системи перетворилися на незамінний інструмент для науковців, дослідників та всіх, хто прагне отримати доступ до наукових знань. Ефективність пошуку наукової інформації формується алгоритмами роботи наукометричних баз даних. За їх допомогою можна здійснити пошук наукових публікацій: книг, глав книг, статей, матеріалів конференцій, тез наукового спрямування тощо.

Цифрові технології та алгоритми пошуку наукової інформації, безумовно, зробили значний крок вперед у підвищенні релевантності пошукової видачі наукових публікацій. В процесі виконання пошукового запиту визначається порядок видачі результатів за ключовими словами, прізвищем автора, заголовками, контекстним аналізом змісту наукових публікацій, їх роком видання та цитуванням. Це потенційно приводить до зростання кількості цитувань наукових публікацій і підвищує індекс цитування наукових публікацій дослідника, його науковий рейтинг та рейтинг закладу вищої освіти.

Зростання ролі цифрових інструментів як допоміжного засобу у сфері наукових досліджень, охоплює майже всі етапи науково-дослідної діяльності та відкриває нові можливості для генерації нових ідей, пошуку наукової інформації у наукометричних базах даних та наукових джерелах відкритого доступу в мережі Інтернет. Вдосконалення пошукових алгоритмів наукометричних систем спрямоване на досягнення максимально можливого співпадіння контексту наукових публікацій у результатах пошукової видачі, яка формується згідно запитів дослідників.

Незважаючи на розвинений потенціал існуючих цифрових інструментів і пошукових алгоритмів, залишаються невирішеними проблемні питання з присутністю наукових публікацій, їх індексації та рейтингування у наукометричних базах даних Scopus та Web of Science (WoS), Google Scholar [2; 5]. Тому важливо з'ясувати та проаналізувати, як функціонування пошукових алгоритмів наукометричних баз даних впливає на результати пошукової видачі наукових публікацій.

Аналіз останніх досліджень та публікацій. В освітній сфері спостерігається інтенсивне використання цифрових технологій, що робить більш інноваційним сам процес навчання та підходи до організації діяльності науковців, дослідників, викладачів та здобувачів вищої освіти. Використання цифрових технологій та їх інструментів підвищує ефективність проведення тематичних досліджень галузевого спрямування на будь-якому етапі науково-дослідного процесу.

В навчальному посібнику [10] розглядається методологія та організація наукових досліджень в галузі інформаційних технологій. Послідовно описані особливості роботи з різноманітними джерелами наукової інформації та їх систематизація для формування теоретичної бази дослідження. Крім того, посібник надає практичні рекомендації щодо пошуку релевантної наукової інформації з використанням наукометричних баз даних (Google Scholar, Scopus, WoS, ResearchGate та інших ресурсів), фільтрації результатів пошукової видачі знайдених наукових публікацій, які стануть у нагоді для проведення якісного наукового дослідження. Рекомендації як опанувати етапи пошуку та збору наукової інформації з різних джерел, використовуючи можливості мережі Інтернет та цифрового середовища для науково-дослідної діяльності надаються в публікації [3].

Постійний розвиток та активне оновлення пошукових алгоритмів розширює можливості виконання складних пошукових запитів. Прикладом зовнішніх інструментів WoS є платформа Publons [20], яка використовується для оцінювання показників результативності наукової діяльності [6].

Ефективність алгоритмів пошуку наукової інформації оцінюється за їх релевантністю. У роботі [4] розглядаються методи інформаційного пошуку та вибору джерел інформації, перераховані алгоритми пошуку, охарактеризована оцінка релевантності результатів пошуку (формальна релевантність, онтологічна релевантність, пертинентність), яку пропонується використовувати для порівняння різних пошукових систем. У дослідженні ефективності видачі результатів пошуку [9] звертається увага на критерії інформаційного запиту, які класифікуються за чіткістю, чисельним складом слів, цілями, частотою і т.д. Інформаційно-пошукові системи визначають релевантність наукових публікацій внутрішніми чинниками ранжування. Власні чинники, які пов'язані із алгоритмами Search Engine Optimization (SEO) наукових публікацій здатні забезпечити їх вірогідну релевантність у результатах пошукової видачі наукометричних баз даних і можуть впливати на їх ранжування [13; 22; 25].

В публікації [11] досліджуються можливості використання інструментів наукометричних баз даних (Google Scholar, Consensus AI, Iris AI) в науково-дослідній та освітній діяльності. Представлені інструменти допомагають провести пошук та систематизацію наукових джерел, оптимізувати етап збору наукової інформації та здійснити її подальший аналіз. В публікаціях [12; 23; 16] порівнюються платформи та інструменти, які застосовуються для аналізу наукометричних даних наукових публікацій, їх індексації та цитувань, серед яких: Google Scholar, Scopus, Web of Science. Таким чином, наукометричні дані наукових публікацій дають змогу пошуковим алгоритмам визначати і надавати найбільш релевантні результати.

Окремої уваги потребує порівняльний аналіз релевантності результатів пошукової видачі наукових публікацій зроблених у наукометричних базах даних Google Scholar [15], Scopus [24], Web of Science [26], IEEE Xplore [18], ResearchGate [21]. Проведення такого аналізу надасть змогу виявити особливості формування результатів пошуку наукової інформації, які найбільш вірогідно відповідають запиту користувача.

Метою статті є оцінка релевантності пошукової видачі наукових публікацій наукометричних баз даних Scopus, Web of Science, IEEE Xplore, ResearchGate, Google Scholar.

Завданнями статті є:

- визначення порядкового номеру найбільш релевантної публікації у пошуковій видачі кожної наукометричної бази даних;
- проведення порівняльного аналізу пошукової видачі найбільш релевантних наукових публікацій;
- зробити оцінку якості роботи пошукових алгоритмів наукометричних баз даних у контексті забезпечення релевантності їх пошукової видачі.

Виклад основного матеріалу дослідження. Інструментарій сучасних наукометричних баз даних надає гнучкі можливості тематичного пошуку та формування персоналізованої пошукової видачі релевантних наукових публікацій для проведення досліджень. Релевантність наукових публікацій встановлює відповідність між результатом пошукової видачі та заданим пошуковим запитом дослідника. Тобто результати пошукової видачі наукових публікацій сформовані кожною наукометричною базою даних, мають максимально можливо відповідати дослідницькій проблематиці та бути корисними для досягнення мети дослідження.

Релевантність наукової інформації визначається наступними критеріями:

- точністю (співпадінням контексту);
- актуальністю (за роком видання);
- авторитетністю джерела інформації, яка визначається наукометричними показниками (h-індекс Гірша, імпакт-фактор, кварталем та категорією періодичного видання та ін.);
- відповідністю запиту дослідника.

Процес пошуку і оцінки наукових публікацій, дозволяє досліднику визначити релевантність пошукової видачі наукометричних баз даних. Він складається з чотирьох етапів (рис. 1) [10, с. 101–102]:

- тематичний комплекс пошуку;
- стратегія пошуку наукової інформації;
- оцінювання релевантності джерел наукової інформації;
- аналіз результатів пошукової видачі.

1. Тематичний комплекс пошуку враховує галузеве спрямування, мету та проблематику наукового дослідження. Такі умови створюють фільтр, який обмежує наявність нерелевантних публікацій в результатах пошукової видачі.



Рис. 1. Процес пошуку та оцінки наукових публікацій

2. **Стратегія пошуку наукової інформації** передбачає поетапний рух від загальних термінів до більш вузьких та конкретних, щоб максимально точно знайти потрібну інформацію.

Пошуковий запит виконується дослідником з урахуванням чинників, які динамічно впливають на релевантність. До важливих умов оцінки роботи пошукових алгоритмів належить [4, с. 48]:

- прямий пошук за текстом наукової публікації (зокрема, ключові слова);
- індексування (попередня обробка наукових публікацій);
- сигнатура (множина хеш-значень слів).

Точність визначення ключових слів наукових публікацій обумовлена процесом аналізу термінів, які використовуються базами даних для їх індексації та ранжування. Правильний вибір ключових слів суттєво впливає на порядковий номер наукової публікації у пошуковій видачі наукометричних баз даних.

Пошуковий запит здійснюється дослідником в певній послідовності, яка починається із введення загальних ключових слів, що описують основну тему наукового дослідження та визначають напрямок пошуку. Наступний етап – це введення спеціальних ключових слів (професійних та технічних спеціальних термінів, скорочень, словосполучень, аббревіатур, синонімів), які допомагають звужити коло пошуку та відфільтрувати нерелевантні наукові публікації. Етап уточнення пошукового запиту дослідника здійснюється із використанням додаткових функцій, логічних операторів пошуку (AND, OR, NOT) та спеціальних пошукових символів, що обумовлюють персональні налаштування пошукового запиту. Ці дії дають змогу досліднику:

- обмежувати кількість наукових публікацій в результатах пошукової видачі за темою дослідження;
- враховувати хронологічні рамки та отримувати найбільш релевантні наукові публікації.

3. **Оцінювання релевантності джерел наукової інформації** формується автоматично за допомогою алгоритмів внутрішнього ранжування наукометричних баз даних. Потім відбувається оцінювання дослідником наукових публікацій у результатах пошукової видачі на предмет визначення їх релевантності. Під час аналізу знайдених наукових публікацій оцінюються їх кількісні та якісні показники (рис. 2).



Рис. 2. Показники релевантності результатів пошуку наукових публікацій

При оцінюванні результатів пошуку наукових публікацій на релевантність, досліднику важливо враховувати відмінності різних наукометричних баз даних. Це пов'язано з тим, що кожна база даних охоплює власний перелік наукових видань.

4. Аналіз результатів пошукової видачі здійснюється для структуризації знайдених наукових публікацій, які будуть використані під час проведення дослідження. На основі цієї вибірки, відповідно до вимог цитування та бібліографічного опису, укладається список використаної літератури [1].

Інструментарій наукометричних баз даних значно прискорює пошук наукових публікацій. В межах цього дослідження провели аналіз результатів пошукової видачі та здійснили оцінку того, наскільки релевантними є знайдені наукові публікації, які були запропоновані пошуковими алгоритмами наукометричних баз даних Scopus [8], WoS [7], ResearchGate [19], IEEE Xplore [17], Google Scholar [14]. Проте постає питання щодо кількісного та якісного оцінювання цих результатів з точки зору дослідника.

Релевантність знайдених наукових публікацій відноситься до показника того результату пошукової видачі наукометричних баз даних, який з максимальною імовірністю відповідає пошуковому запиту дослідника. Тому, оцінювання релевантності наукових публікацій пропонується розглядати як процес, що поєднує дії дослідника з:

- фільтрації наукових публікацій, які відповідають критеріям релевантності;
- семантичного аналізу змісту знайдених наукових публікацій, які відповідають тематичній проблематиці дослідження на час проведення пошуку;
- проведенням додаткових (уточнюючих) пошукових запитів для отримання оновленої пошукової видачі наукових публікацій;
- ретельної перевірки повноти відбору наукових публікацій та точності бібліографічного опису цих джерел.

Кількісні показники обчислюються алгоритмами пошукової системи на основі формальних ознак, таких як кількість використання ключових слів, наявність їх у заголовках, анотаціях, а також кількістю цитувань. Показниками для цієї оцінки є точність – частка релевантних публікацій серед усіх знайдених, та повнота – частка знайдених релевантних публікацій від загальної кількості релевантних у базі даних. Між цими двома показниками існує зворотній зв'язок: спроба підвищити повноту часто призводить до зниження точності, і навпаки.

Кількісний показник оцінювання результатів пошукової видачі здійснюється користувачем за допомогою використання вбудованих фільтрів наукометричних баз даних. Застосування фільтрів може як зменшити обсяг та підвищити точність отриманих результатів, так і збільшити кількість наукових публікацій. Таким чином, кожне уточнення, яке додається до пошукового запиту, звужує коло пошуку. Чим більше конкретних умов вказати, тим меншою буде загальна кількість знайдених результатів. І навпаки, дуже загальний пошуковий запит без додаткової фільтрації може видати велику кількість наукових публікацій, що значно ускладнює пошук дослідника.

До факторів, які впливають на релевантність пошуку наукових публікацій можна віднести:

- відповідність ключових слів проведеному дослідженню гарантує ефективність пошуку наукової інформації;
- оновлюваність наукометричних баз даних за умовний проміжок часу може істотно позначитися не тільки на результатах пошукової видачі, а й на порядковому номері наукової публікації у пошуковій видачі;
- регулярність оновлення алгоритмів пошуку вносить зміни до ранжирування наукових публікацій.

Щоб отримати точну кількість результатів для аналізу, досліднику необхідно:

- чітко сформулювати пошуковий запит за ключовим словом;
- визначити наукові категорії пошуку;
- обрати «тип документа» – книга, частина книги, стаття, конференція;
- запустити пошук наукових публікацій.

Результати пошукової видачі з розподілом за ключовими словами з зазначених вище наукометричних баз даних наведені в (табл. 1).

Слід звернути увагу, що в таблиці 1 відображені кількісні результати пошуку наукових публікацій в зазначених наукометричних базах даних на момент проведення дослідження 01–25.07.2025 року.

Загалом дослідник переглядає від 5 до 10 сторінок результатів пошукової видачі або перших 100 знайдених наукових публікацій. Якщо користувач вважає, що кількість знайдених наукових публікацій є достатньою і вони відповідають критеріям релевантності дослідника, то він припиняє пошук і переходить до їх аналізу. Якщо при виконанні аналізу знайдені наукові публікації не в повній

мірі відповідають встановленим критеріям релевантності дослідника або результати пошуку не задовольняють його вимоги, він продовжує пошук. Для отримання більш релевантних результатів пошуку під час проведення наступного пошукового запиту, ватро змінити ключові слова, застосувати різні вбудовані інструменти пошуку наукометричних баз даних, чи використати інші наукомеричні бази даних.

Таблиця 1

**Кількісні показники результатів пошукової видачі наукових публікацій
у наукометричних базах даних**

Наукометрична база даних	Загальна кількість публікацій (книги, частини книг, статті, конференції)			
	complex systems	digital systems	artificial intelligence	data analysis
Scopus	1,657,460	843,848	717,257	6,988,030
WoS	43847	521,648	683,424	5,235,311
IEEE Xplore	149828	445,143	294,796	786,920
ResearchGate	не надає дані			
Google Scholar	9 560 000	6 900 000	7 070 000	8 310 000

Водночас результатів пошукової видачі наукометричних баз даних недостатньо для повної оцінки релевантності наукових публікацій, оскільки вони можуть не враховати суб'єктивну інформаційну потребу дослідника. Результати пошуку можуть бути формально релевантні пошуковому запиту (наприклад, містити ключові слова), але для дослідника надана пошукова видача буде не релевантною.

Якісний показник оцінювання релевантності пошукової видачі здійснюється користувачем самостійно.

Суб'єктивність оцінки релевантності наукових публікацій дослідником обумовлюється:

- власним рівнем знання предметної тематики дослідження;
- спроможністю сформулювати контекстні пошукові запити;
- персоналізацією пошукових запитів (попереднім і кінцевим вибором ключових слів);
- досвідом застосування пошукових інструментів наукометричних баз даних;
- вмінням оцінювати результати пошукової видачі наукометричних баз даних.

З точки зору професійного дослідника, існує необхідність у визначенні найбільш релевантних наукових публікацій знайдених у зазначених вище наукометричних базах даних [18; 26; 21; 15; 24]. Це пропонується зробити шляхом фіксації порядкового номеру релевантних наукових публікацій визначених авторами цього дослідження у результатах ТОП-1000 пошукової видачі кожної наукометричної бази даних. Далі потрібно провести порівняльний аналіз на предмет наявності визначених релевантних наукових публікацій у кожній наукометричній базі і зафіксувати отримані результати. Критерій визначення релевантних наукових публікацій слід обґрунтовувати точним співпадінням з ключовими словами у пошукових запитах та їх наявністю у назві, анотації та ключових словах наукової публікації. Результати порівняльного аналізу релевантності наукових публікацій у результатах пошукової видачі зазначених вище наукометричних баз даних наведено у (табл. 2).

Результати наведені в таблиці 2 сформовані, виходячи з аналізу ТОП-1000 пошукової видачі наукових публікацій. Для наочного уявлення результатів проведеного дослідження на (рис. 3) представлено графічну інтерпретацію порівняльного аналізу релевантності пошукової видачі наукометричних баз даних, на основі отриманих результатів дослідження (див. табл. 2).

Пошукова видача наукових публікацій в наукометричних базах даних формувалася відповідно до таких критеріїв:

- *Scopus* – у режимі «за релевантністю», тип документа – «Conference paper», «Article», «Book chapter», «Book», «Keywords»;
- *WoS* – у режимі пошуку «Web of Science Core Collection», тип документа – «Article», «Book chapter», «Book», «Conference», категорії – «Title», «Abstract»;
- *IEEE Xplore* – в режимі «Advanced Search», категорії – «Document title», «Abstract», «Author keywords»;
- *ResearchGate* – у режимі «за релевантністю», тип документа – «Articles», «Books», «Conference papers»;
- *Google Scholar* – у режимі «за релевантністю», тип документа – «All articles», «Include patents».

Таблиця 2

Результати пошукової видачі наукових публікацій у наукометричних базах даних

Ключові слова для пошуку наукових публікацій	Назва найбільш релевантної наукової публікації за оцінкою авторів	Порядковий номер наукової публікації в ТОП 1000 пошукової видачі									
		Scopus	наявність(+/-)	WoS	наявність(+/-)	IEEE Xplore	наявність(+/-)	ResearchGate	наявність(+/-)	Google Scholar	наявність(+/-)
complex systems	Complex system reconstruction	11	+	0	+	0	-	0	-	0	-
digital systems	Evolution of Enterprise Architecture for Intelligent Digital Systems	12	+	4	+	234	+	424	+	367	+
artificial intelligence	The use of artificial intelligence in the modern world – An overview	6	+	0	-	0	-	287	+	0	-
data analysis	A cognitive interpretation of data analysis	2	+	176	+	0	-	0	-	0	-

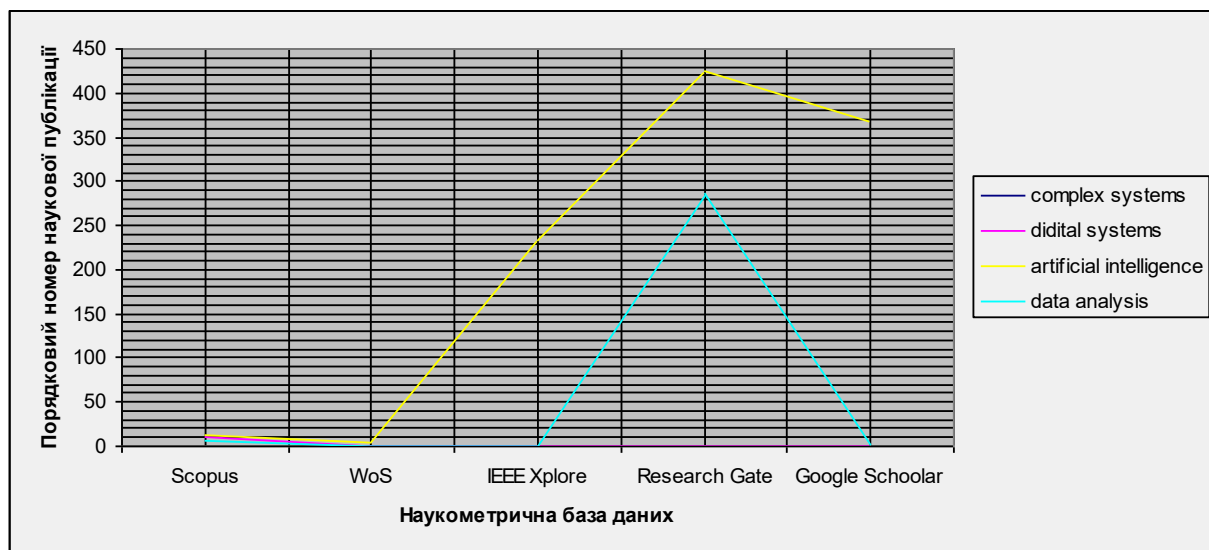


Рис. 3. Графічна інтерпретація пошукової видачі наукометричних систем

Висновки та перспективи подальших досліджень. Аналіз отриманих результатів пошукової видачі наукометричних баз даних Google Scholar, Scopus, Web of Science, ResearchGate, IEEE Xplore показав наявність значних відмінностей у формуванні пошукової видачі наукових публікацій, а саме:

- в результатах пошукової видачі Google Scholar після 450 порядкового номера містяться наукові публікації, в яких ключове слово присутнє тільки в тексті анотації;
- в результати пошукової видачі Google Scholar ТОП-100 формуються з книг та/або з розділів книг;
- в результатах пошукової видачі Researchgate представлені наукові публікації, в яких ключові слова частково співпадають з їх заголовками;
- кожна з розглянутих наукометричних баз даних має власні внутрішні обмеження на умови включення до них наукових публікацій;
- релевантність наукових публікацій визначається власними пошуковими алгоритмами баз даних, які впливають на відповідність запиту дослідника.

Проведений аналіз пошукової видачі наукометричних баз даних дозволив зробити такі висновки:

- релевантність пошукової видачі покращується, якщо пошуковий запит дослідника міститься у заголовку публікації, анотації та ключових словах;
- ефективність пошуку наукових публікацій безпосередньо залежить від того, наскільки точно заголовок та ключові слова наукової публікації відповідають її змісту;
- для досягнення максимальної повноти релевантних наукових публікацій, досліднику рекомендовано використовувати для пошуку якомога більше наукометричних баз даних;
- підсумкове рішення щодо відповідності наукових публікацій пошуковому запиту належить приймати самому досліднику після ретельного аналізу контексту публікацій.

В цілому формування та подальший аналіз пошукової видачі наукометричних баз даних надає змогу дослідникам:

- оптимізувати наукову роботу;
- швидко та ефективно здійснювати пошук релевантних наукових публікацій;
- виявляти провідних авторів та наукові установи в потрібній галузі досліджень;
- відстежувати актуальні тенденції та перспективні напрямки досліджень;
- знаходити нові ідеї для досліджень;
- оцінювати конкурентоспроможність власних досліджень;
- здійснювати моніторинг власних наукометричних показників;
- формувати стратегії власної публікаційної діяльності;
- постійно покращувати власні наукометричні показники.

Проведений порівняльний аналіз релевантності пошукової видачі виявив, що наукометричні бази даних, зокрема Scopus, Web of Science та IEEE Xplore, характеризуються вищим рівнем стабільності та точності результатів. Водночас відкрита платформа Google Scholar забезпечує ширше охоплення наукових джерел, особливо з відкритим доступом.

Саме тому досліднику рекомендується використовувати системний підхід до наукового пошуку, поєднуючи максимально можливу кількість наукометричних баз даних для забезпечення повноти та релевантності інформаційного покриття.

Аналіз релевантності пошукової видачі наукометричних баз даних підтверджує необхідність постійного вдосконалення пошукової стратегії дослідника та методів оцінки релевантності наукової інформації. Отримані результати можуть бути використані для розробки рекомендацій щодо оптимізації пошукових стратегій дослідника, що цілком буде сприяти покращенню рівня якості його наукових досліджень.

Список використаних джерел:

1. ДСТУ 8302:2015 Інформація та документація. Бібліографічне посилання. Загальні положення та правила складання. URL: <http://surl.li/qgqxf>
2. Колмакова В. О., Шаров С. В. Проблеми індексації публікацій в наукометричних базах даних та способи їх вирішення. *Українські студії в європейському контексті*. 2023. № 7. С. 239–246.
3. Литвин С. Х., Добровольська В. В. Інформаційне забезпечення науково-дослідної роботи здобувачів наукового ступеня доктора філософії. *Бібліотекознавство. Документознавство. Інформологія*. 2021. № 3. С. 90–98. DOI: <https://doi.org/10.32461/2409-9805.3.2021.244724>
4. Мілованова М. В., Бондарчук А. П. Оцінка ефективності інформаційного пошуку. *Телекомунікаційні та інформаційні технології*. 2020. № 1(66). С. 45–52. DOI: <https://doi.org/10.31652/2412-4338.2020.014552>
5. Назаровець С. Бази даних цитувань та пошукові інструменти для науковців майбутнього. *Наука і суспільство*. 2021. № 1(87). С. 35–38.
6. Наукометричні показники оцінювання результативності педагогічних досліджень науковців та науково-педагогічних працівників / Т. А. Вакалюк, О. М. Спірін, І. С. Мінтій, С. М. Іванова, Т. Л. Новицька. *Сучасні інформаційні технології та інноваційні методики навчання у підготовці фахівців: методологія, теорія, досвід, проблеми* : зб. наук. праць. Вінниця, 2021. Вип. 60. С. 167–184. DOI: <https://doi.org/10.31652/2412-1142-2021-60-167-184>
7. Пошук журналів у Web of Science : практ. посіб. / упоряд. А. М. Кушнерський ; Київ. нац. ун-т ім. Т. Шевченка, Наук. б-ка ім. М. Максимовича, Служба інформ. моніторингу. Київ, 2021. 28 с. URL: http://www.library.univ.kiev.ua/ukr/res/journals_wos.pdf (дата звернення: 18.06.2025)
8. Пошук у Scopus : практ. посіб. / упоряд. М. А. Назарець ; Київ. нац. ун-т ім. Т. Шевченка, Наук. б-ка ім. М. Максимовича, Служба інформ. моніторингу. Київ, 2022. 29 с. URL: http://www.library.univ.kiev.ua/ukr/res/scopus_search.pdf (дата звернення: 18.06.2025)
9. Розробка методу підвищення ефективності видачі результату пошуку в інформаційному середовищі / І. В. Муляр, П. А. Шкуліпа, І. А. Романовська, Л. О. Ряба. *Збірник наукових праць Військового інституту Київського національного університету імені Тараса Шевченка*. 2016. Вип. 52. С. 134–141.

10. Савельєва О. С., Трофименко О. Г., Дикий О. В. Методологія проведення наукових досліджень у галузі інформаційних технологій : навч. посіб. [Електрон. вид.]. Одеса : Фенікс, 2025. 144 с. URL: <http://dspace.opu.ua/jspui/handle/123456789/14858>
11. Толочко С. В., Годунова А. В. Теоретико-методичний аналіз оптимізації дослідницької та наукової діяльності в умовах використання сервісів зі штучним інтелектом. *Інноваційна педагогіка*. 2023. Вип. 61(2). С. 18–24. DOI: <https://doi.org/10.32782/2663-6085/2023/61.2.3>
12. Ярошенко Т., Ярошенко О. Вимірювання впливу науки: за межі традицій. Порівняльний аналіз сучасних наукометричних інструментів та їх роль у визначенні наукового внеску. *Відкрита та інновації*. 2024. №. 1. С. 18–37. DOI: <https://doi.org/10.62405/osi.2024.01.02>
13. Brouwer W. P., Hollenbach M. Search engine optimization for scientific publications: How one can find your needle in the haystack. *United European Gastroenterology Journal*. 2022. Т. 10. №. 8. С. 906. URL: <https://onlinelibrary.wiley.com/doi/pdfdirect/10.1002/ueg2.12311>
14. GOOGLE академія для науковців [Електронний ресурс] : практ. посіб. / упоряд. М. А. Назаровець. Київ : ВПЦ «Київ. ун-т», 2016. 31 с. URL: http://www.library.univ.kiev.ua/ukr/res/google_scholar.pdf (дата звернення: 20.06.2025)
15. Google Академія. URL: <https://scholar.google.com/>
16. Gusenbauer M. Beyond Google Scholar, Scopus, and Web of Science: An evaluation of the backward and forward citation coverage of 59 databases' citation indices. *Research Synthesis Methods*. 2024. № 15(5). С. 802–817. DOI: <https://doi.org/10.1002/jrsm.1729>
17. Henriques P. Strategies for Searching IEEE Xplore. URL: https://biblioteca.upc.edu/sites/default/files/pages_generals/investigadors/ieee_xplore.pdf
18. IEEE Xplore. URL: <https://ieeexplore.ieee.org/Xplore/home.jsp>
19. Kirilova S., Zoepfl F. Metrics fraud on ResearchGate. *Journal of Informetrics*. 2025. Т. 19. №. 1: 101604. DOI: <https://doi.org/10.1016/j.joi.2024.101604>
20. Publons. URL: <https://www.webofscience.com/wos/woscc/basic-search>
21. ResearchGate. URL: <https://www.researchgate.net/search>
22. Schilhan L., Kaier C., Lackner K. Increasing visibility and discoverability of scholarly publications with academic search engine optimization. *Insights*. 2021. Т. 34. №. 1. URL: <https://insights.uksg.org/articles/10.1629/uksg.534?s=09>
23. Scopus та Web of Science : навч. посіб. / М. М. Іжа, М. П. Попов, Л. Л. Приходченко [та ін.]. Одеса : ОРІДУ НАДУ, 2021. 174 с. URL: <http://dspace.opu.ua/jspui/handle/123456789/13289>
24. Scopus. URL: <https://www.scopus.com/sources>
25. Shahzad A. et al. The impact of search engine optimization on the visibility of research paper and citations. *JOIV: International Journal on Informatics Visualization*. 2017. Т. 1. №. 4-2. С. 195–198. URL: <https://www.joiv.org/index.php/joiv/article/viewFile/77/77>
26. Web of Science. URL: <https://www.webofscience.com/>

Дата надходження статті: 04.07.2025

Дата прийняття статті: 10.07.2025

Опубліковано: 23.09.2025

УДК 004.8:316.774

DOI <https://doi.org/10.32689/maup.it.2025.2.8>

Володимир КАРАБЧУК

аспірант, Приватний вищий навчальний заклад «Європейський університет»,

volodymyr.karabchuk@e-u.edu.ua

ORCID: 0009-0002-4320-3169

МЕТОДИ ВИЯВЛЕННЯ ТА БОРОТЬБИ З ДЕЗІНФОРМАЦІЄЮ ЗА ДОПОМОГОЮ ТЕХНОЛОГІЙ ШТУЧНОГО ІНТЕЛЕКТУ

Анотація. У статті розглядаються сучасні методи виявлення та боротьби з дезінформацією за допомогою технологій штучного інтелекту. Основна увага приділяється автоматизації аналізу текстової, візуальної та відеоінформації, використовуючи алгоритми машинного навчання, обробки природної мови та комп'ютерного зору. Розглянуто інструменти для моніторингу соціальних мереж, перевірки фактів і виявлення маніпуляцій у медіаконтенті. Окремо проаналізовано виклики, пов'язані з моральними аспектами, дефіцитом даних і розвитком технологій для створення фейків. Запропоновані підходи спрямовані на підвищення ефективності боротьби з дезінформацією, забезпечення інформаційної безпеки та стійкості суспільства до маніпуляцій.

Мета роботи. Дослідження сучасних методів виявлення та боротьби з дезінформацією за допомогою технологій штучного інтелекту, зокрема машинного навчання, обробки природної мови та комп'ютерного зору, для підвищення інформаційної безпеки суспільства.

Методологія. У роботі використано алгоритми машинного навчання, такі як згорткові нейронні мережі (CNN) та трансформерні моделі (наприклад, BERT), для аналізу текстового контенту. Для ідентифікації маніпуляцій у візуальному контенті застосовано технології комп'ютерного зору, а для аналізу поширення дезінформації в соціальних мережах – графові нейронні мережі (GNN).

Наукова новизна. Запропоновано інтегрований підхід, що включає аналіз текстового, візуального та відеоконтенту одночасно. Досліджено вплив емоційного забарвлення контенту на сприйняття фейкових новин і вперше вивчено роль мультимодальних систем у боротьбі з дезінформацією.

Висновки. Технології штучного інтелекту забезпечують ефективне виявлення дезінформації шляхом аналізу великих обсягів даних у реальному часі. Інтеграція мультимодальних підходів

сприяє підвищенню стійкості інформаційного простору, хоча залишається необхідність розв'язання проблеми дефіциту якісних навчальних даних і врахування етичних аспектів.

Ключові слова: дезінформація, штучний інтелект, обробка природної мови, машинне навчання, перевірка фактів, маніпулятивний контент, deepfake, кібербезпека.

Volodymyr KARABCHUK. METHODS FOR DETECTING AND COMBATING DISINFORMATION USING ARTIFICIAL INTELLIGENCE TECHNOLOGIES

Abstract. The article examines modern methods for detecting and combating disinformation using artificial intelligence technologies. The primary focus is on automating the analysis of textual, visual, and video information through the use of machine learning algorithms, natural language processing, and computer vision. Tools for monitoring social networks, fact-checking, and identifying manipulations in media content are discussed. Challenges related to ethical aspects, data scarcity, and the development of technologies for creating fake content are also analyzed. The proposed approaches aim to enhance the effectiveness of combating disinformation, ensure information security, and increase societal resilience to manipulation.

Purpose. To study modern methods for detecting and combating disinformation using artificial intelligence technologies, particularly machine learning, natural language processing, and computer vision, to enhance societal information security.

Methodology. The study employs machine learning algorithms, such as convolutional neural networks (CNN) and transformer models (e.g., BERT), for textual content analysis. Computer vision technologies are applied to detect manipulations in visual content, while graph neural networks (GNN) are used to analyze the dissemination of disinformation in social networks.

Scientific novelty. An integrated approach is proposed, combining the analysis of textual, visual, and video content simultaneously. The study investigates the impact of emotional coloring on the perception of fake news and explores the role of multimodal systems in combating disinformation for the first time.

Conclusions. Artificial intelligence technologies provide effective disinformation detection by analyzing large volumes of data in real-time. The integration of multimodal approaches enhances the resilience of the information space, though challenges such as the lack of high-quality training data and ethical considerations remain.

Key words: disinformation, artificial intelligence, natural language processing, machine learning, fact-checking, manipulative content, deepfake, cybersecurity.

Вступ. Дезінформація залишається однією з найважливіших загроз інформаційній безпеці, впливаючи на політичні вибори, соціальну стабільність і економічний розвиток. Наприклад, під час пандемії COVID-19 значна кількість фейкових новин про вакцини викликала масову дезорієнтацію громадськості, що, за даними ВООЗ, призвело до зниження рівня вакцинації в окремих регіонах. Інформаційні війни, як-от ті, що розгортаються під час міжнародних конфліктів, підкреслюють, наскільки важливою є швидка та точна перевірка достовірності даних.

Дослідження в цій галузі вказують, що штучний інтелект здатний ефективно протидіяти дезінформації. Наприклад, алгоритми обробки природної мови (NLP) вже застосовуються для автоматизованого аналізу текстів у реальному часі. Дослідження, [2] опубліковане в журналі *Information Sciences* (2022), демонструє, як моделі трансформерів, такі як BERT та GPT, можуть аналізувати мільйони текстів для виявлення маніпуляцій і упереджених тверджень. У візуальній сфері технології комп'ютерного зору, такі як методи аналізу пікселів і метаданих, використовуються для ідентифікації deepfake-зображень, що підтверджується [4, 10] результатами проєкту DeepFake Detection Challenge.

Водночас існує низка викликів, які ускладнюють ефективне застосування штучного інтелекту для боротьби з дезінформацією. Зокрема, спостерігається недостатність якісних навчальних даних, які необхідні для створення універсальних моделей, здатних працювати в багатомовних і культурно різноманітних середовищах. Крім того, швидке вдосконалення технологій створення фейкового контенту, таких як deepfake, значно ускладнює їх своєчасне виявлення. Окрему проблему становлять етичні аспекти використання ШІ, зокрема ризик хибного маркування правдивої інформації як дезінформації через контекстні помилки. Наприклад, дослідження [9] показують, що автоматизовані системи можуть неправильно інтерпретувати сарказм або сатиру, що призводить до помилкових результатів (Mazurczyk et al., 2023).

У цій роботі акцент зроблено на створенні інтегрованого підходу до виявлення дезінформації, що включає аналіз текстового контенту, моніторинг соціальних мереж та ідентифікацію маніпуляцій у відеоматеріалах. Використання сучасних моделей машинного навчання та NLP сприятиме підвищенню ефективності аналізу, забезпечуючи стійкість інформаційного простору до загроз дезінформації.

Постановка проблеми. У сучасному інформаційному просторі, де кожна секунда має вирішальне значення, дезінформація стала серйозною загрозою для соціальної стабільності, політичних процесів і довіри до медіа. Поширення фейкових новин, маніпулятивного контенту та візуальних фейків (наприклад, deepfake) створює значні ризики для громадського порядку та безпеки.

Використання технологій штучного інтелекту, зокрема машинного навчання та обробки природної мови, стало ефективним підходом до виявлення дезінформації. Однак виникає низка проблем, які ускладнюють реалізацію цих рішень. Серед основних викликів:

- Нестача якісних навчальних даних для розпізнавання маніпуляцій;
- Висока швидкість створення нових технологій для дезінформації, таких як deepfake;
- Необхідність аналізу великих обсягів даних у реальному часі.

Особливу увагу привертає емоційний вплив дезінформації. Дослідження [12] показують, що фейкові новини часто створюються з метою викликати страх, ненависть або інші сильні емоції, які роблять аудиторію більш вразливою до маніпуляцій. Існуючі підходи, такі як використання згорткових нейронних мереж (CNN), аналіз настроїв та ідентифікація емоційних патернів, є ефективними, але їхня реалізація вимагає подальшого вдосконалення.

Проблема полягає у розробці інноваційних методів виявлення дезінформації, які б забезпечували високу точність та ефективність навіть за умов обмеженості ресурсів і швидких змін технологій.

Аналіз останніх досліджень і публікацій. Застосування технологій штучного інтелекту для виявлення та протидії дезінформації активно вивчається в останні роки. Зокрема, дослідники приділяють увагу використанню алгоритмів обробки природної мови (NLP), машинного навчання та комп'ютерного зору. Одним із ключових напрямків є аналіз емоційного впливу дезінформаційного контенту, оскільки фейкові новини часто спрямовані на викликання сильних емоцій, таких як страх, ненависть або гнів. Це підтверджено дослідженнями [5], які вказують, що такі емоції підвищують довіру до неправдивого контенту.

Наукові публікації активно досліджують [6] автоматизацію виявлення дезінформації, зокрема застосування згорткових нейронних мереж (CNN) для класифікації текстів та аналізу візуального контенту. Дослідження демонструють [15], що CNN ефективно виявляють фейкові новини через аналіз структурних особливостей тексту та візуальних маніпуляцій. Інтеграція подвійних емоційних ознак із традиційним аналізом настроїв значно підвищує точність виявлення маніпулятивних текстів. Крім того, графові нейронні мережі (GNN) показують перспективи у виявленні координованих дезінформаційних кампаній у соціальних мережах, покращуючи аналіз структурованих взаємодій у цифровому просторі.

Особливий інтерес викликають підходи до автоматичного аналізу візуального контенту. Технології комп'ютерного зору, наприклад, системи для виявлення deepfake-зображень і відео, використовуються для розпізнавання маніпуляцій у пікселях і метаданих. Дослідження [14] підтверджує ефективність цих технологій у реальних умовах.

Попри досягнення, дослідники зазначають проблеми, зокрема нестачу якісних навчальних даних та труднощі аналізу багатомовного контенту. В роботах [11, 13] запропоновано створення універсальних корпусів даних для навчання моделей, що дозволяють адаптувати рішення до специфіки певних регіонів або мов.

Мета. Метою статті є дослідження сучасних методів виявлення та боротьби з дезінформацією за допомогою технологій штучного інтелекту, зокрема машинного навчання, обробки природної мови та комп'ютерного зору. Для досягнення цієї мети необхідно вирішити наступні завдання:

1. Провести аналіз останніх досліджень та публікацій щодо застосування штучного інтелекту для виявлення дезінформації.
2. Дослідити ефективність алгоритмів машинного навчання, таких як згорткові нейронні мережі (CNN), та їх роль у розпізнаванні фейкових новин.
3. Вивчити підходи до аналізу емоційного впливу текстів, зображень і відеоконтенту як одного з ключових чинників розпізнавання маніпуляцій.
4. Оцінити вплив якості та обсягу навчальних даних на точність методів виявлення дезінформації.
5. Запропонувати перспективні напрями вдосконалення технологій штучного інтелекту для підвищення ефективності боротьби з дезінформацією.

Результати дослідження сприятимуть кращому розумінню ролі штучного інтелекту у забезпеченні інформаційної безпеки та підвищенні стійкості суспільства до дезінформаційних загроз.

Виклад основного матеріалу. Дезінформація, або свідоме поширення неправдивої інформації, є однією з головних загроз сучасного цифрового простору. Завдяки доступності технологій та стрімкому розвитку соціальних мереж, фейкові новини стали невід'ємною частиною інформаційного середовища. Їхня здатність впливати на думки мільйонів людей призводить до серйозних соціальних, політичних і економічних наслідків.

Такі платформи, як Facebook, Twitter та TikTok, дозволяють користувачам швидко ділитися інформацією, часто без її попередньої перевірки. Це створює ідеальне середовище для поширення маніпулятивного контенту, зокрема таких типів дезінформації:

- Фейкові новини: Статті або повідомлення, що свідомо викривляють факти.
- Візуальні фейки: Зображення або відео, змінені за допомогою технологій, таких як deepfake.
- Маніпулятивні повідомлення: Контент, створений для викликання сильних емоцій, таких як страх, ненависть чи гнів.

Ці види дезінформації часто поширюються через бот-мережі або цілеспрямовані кампанії, спрямовані на підірив довіри до офіційних джерел інформації.

Сучасні виклики у боротьбі з дезінформацією вимагають застосування інноваційних технологій, таких як штучний інтелект. Завдяки розвитку алгоритмів машинного навчання, обробки природної мови (NLP) та комп'ютерного зору, з'являються нові можливості для автоматичного виявлення дезінформації. Наприклад, аналіз емоційного забарвлення текстів дозволяє ідентифікувати маніпулятивні елементи, тоді як технології комп'ютерного зору успішно застосовуються для розпізнавання deepfake-зображень та відео.

Таким чином, боротьба з дезінформацією є ключовою складовою забезпечення інформаційної безпеки. Її ефективне вирішення вимагає комплексного підходу, що включає використання сучасних технологій, медіаосвіти та міжнародної співпраці.

Інструменти та алгоритми боротьби з дезінформацією. Одним із ключових інструментів у боротьбі з дезінформацією є автоматизовані системи перевірки фактів. Наприклад, ClaimBuster [7] використовує алгоритми обробки природної мови (NLP) для ідентифікації тверджень у текстах, які потребують перевірки. Ця система може працювати в реальному часі, аналізуючи публікації політиків чи новинних ресурсів.

Ще одним важливим напрямком є аналіз візуального контенту. Технології комп'ютерного зору дозволяють виявляти маніпуляції у зображеннях та відео. Наприклад, DeepTrace [3] застосовує алгоритми згорткових нейронних мереж (CNN) для розпізнавання deepfake, аналізуючи такі характеристики, як текстура шкіри, синхронізація рухів губ із аудіо, а також метадані файлів. Схожим чином працюють інструменти на кшталт Amped Authenticate [1], які зосереджені на перевірці метаданих та ідентифікації підробок у візуальному контенті.

Соціальні мережі, такі як TikTok, Instagram чи LinkedIn, є ключовими платформами для поширення інформації, включно з дезінформацією. Інструменти, як-от Noaxy [8], аналізують маршрути поширення контенту, відображаючи мережі користувачів, які діляться маніпулятивною інформацією. Платформи на кшталт Hootsuite чи Brandwatch дозволяють відстежувати та аналізувати вплив контенту, включно з дезінформацією, використовуючи сучасні алгоритми обробки природної мови (NLP).

Комплексний підхід також включає мультимодальні системи, такі як DeepMind Gopher [16], які поєднують можливості великих мовних моделей та алгоритмів машинного навчання для виявлення, аналізу й протидії дезінформації. Це дозволяє ефективно працювати з текстовим, візуальним і відеоконтентом одночасно, використовуючи передові методи штучного інтелекту.

Основні виклики в боротьбі з дезінформацією

Незважаючи на значний прогрес у розробці технологій для виявлення та протидії дезінформації, існує низка викликів, які ускладнюють досягнення високої ефективності в цій сфері. Ці виклики охоплюють як технічні, так і етичні аспекти, що вимагають комплексного підходу до їх вирішення.

Недостатність якісних навчальних даних

Алгоритми машинного навчання та обробки природної мови значною мірою залежать від доступності якісних і репрезентативних навчальних даних. Проте створення таких корпусів є складним завданням через:

- Відсутність доступу до перевірених джерел, особливо в регіонах із обмеженою свободою слова.
- Проблему багатоаспектності та багатомовності дезінформації, що вимагає створення моделей, адаптованих до різних культурних контекстів.
- Нестачу структурованих наборів даних для навчання алгоритмів аналізу візуального контенту, таких як deepfake-зображення.

Розвиток технологій створення дезінформації

Зловмисники активно використовують новітні технології для створення складного маніпулятивного контенту. Наприклад, генеративні моделі, такі як GPT, дозволяють створювати правдоподібні тексти, які важко відрізнити від достовірних. Технології deepfake, з іншого боку, ускладнюють виявлення змінених зображень та відео через їхню високу якість. Цей постійний розвиток технологій створює «гонку озброєнь», де інструменти для боротьби з дезінформацією повинні постійно вдосконалюватися.

Етичні аспекти та помилкові маркування

Автоматизовані системи, такі як фактчекінгові платформи, ризикують помилково маркувати правдиву інформацію як дезінформацію. Це може мати серйозні наслідки, зокрема:

- Підрив довіри до систем боротьби з дезінформацією.
- Викликання соціальної напруги через спотворення достовірних даних. Крім того, виникають питання конфіденційності даних, оскільки алгоритми часто вимагають доступу до великої кількості особистої або чутливої інформації для навчання моделей.

Швидкість поширення дезінформації

Дезінформація поширюється значно швидше, ніж інструменти її виявлення встигають зреагувати. Соціальні мережі надають можливість охоплення широкої аудиторії у дуже короткі строки. Наприклад, фейкова новина може досягти мільйонів користувачів упродовж кількох годин. Це створює значний виклик для систем, які працюють у режимі реального часу.

Проблеми міжнародної координації

Дезінформація часто має глобальний характер, але боротьба з нею ускладнюється через різні підходи до регулювання та обмеження інформації в різних країнах. Брак координації між державами та міжнародними організаціями обмежує ефективність глобальних заходів протидії.

Запропонований інтегрований підхід

Для ефективної боротьби з дезінформацією пропонується інтегрований підхід, який охоплює аналіз текстового контенту, моніторинг соціальних мереж та ідентифікацію маніпуляцій у візуальному і відеоконтенті. Основні компоненти підходу включають:

Моніторинг соціальних мереж:

- Застосування NLP-алгоритмів для автоматичного виявлення ключових слів і фраз, які сигналізують про можливе поширення фейкових новин.
- Використання графових моделей для виявлення бот-мереж та координованих кампаній із поширення дезінформації.

Аналіз візуального та відеоконтенту:

- Технології комп'ютерного зору дозволяють виявляти deepfake-зображення та відео шляхом аналізу піксельних патернів, метаданих та синхронізації аудіо з відео.

● Алгоритми згорткових нейронних мереж (CNN) ефективно розпізнають візуальні маніпуляції, зокрема змінену текстуру зображень чи неприродні рухи в відео.

Інтеграція мультимодальних даних:

● Поєднання аналізу тексту, зображень та поведінкових патернів у єдину систему дозволяє не лише ідентифікувати дезінформацію, але й оцінити її потенційний вплив на аудиторію.

● Використання мультимодальних нейронних мереж для одночасного аналізу кількох типів контенту.

Постійний моніторинг та вдосконалення:

● Впровадження механізмів автоматичного збору зворотного зв'язку від користувачів для вдосконалення моделей.

● Постійне оновлення алгоритмів для адаптації до нових форм дезінформації.

Запропонований підхід дозволяє створити комплексну систему, яка ефективно протидіє поширенню дезінформації, підвищуючи стійкість інформаційного простору до маніпуляцій. Інтеграція різних технологій штучного інтелекту сприятиме зменшенню впливу фейкових новин і забезпеченню інформаційної безпеки.

Перспективи розвитку

Розвиток технологій штучного інтелекту відкриває нові горизонти у боротьбі з дезінформацією, дозволяючи створювати інноваційні підходи та системи для її виявлення й запобігання. Проте зростання складності маніпулятивного контенту потребує не лише вдосконалення існуючих методів, а й впровадження стратегій, які враховують майбутні виклики.

1. Інтеграція мультимодальних систем

Майбутнє боротьби з дезінформацією полягає у створенні інтегрованих систем, які аналізують текстовий, візуальний та поведінковий контент одночасно. Поєднання можливостей обробки природної мови (NLP), комп'ютерного зору та аналізу поведінки користувачів дозволить створити більш комплексні рішення, які будуть адаптовані до нових форм маніпуляцій.

Приклад: Інтеграція даних із соціальних мереж, фактчекінгових платформ і візуального контенту в єдину систему моніторингу значно підвищить ефективність виявлення дезінформації на ранніх етапах її поширення.

2. Розвиток алгоритмів генерації контрнарративів

Замість блокування дезінформації ефективним підходом може стати створення контрнарративів, які пояснюють аудиторії неправдивий характер маніпуляцій. Це дозволить не лише запобігати поширенню фейків, а й підвищити рівень медіаграмотності суспільства.

Перспектива: Використання генеративних моделей, таких як GPT, для створення текстів, які пояснюють аудиторії небезпеку дезінформації та надають достовірну інформацію.

3. Використання блокчейн-технологій

Блокчейн може стати ефективним інструментом у забезпеченні прозорості джерел інформації. Технологія дозволяє зберігати історію змін у контенті, підтверджувати його автентичність і надавати користувачам можливість перевіряти достовірність інформації.

Приклад: Використання блокчейну для створення системи цифрового підтвердження автентичності новинних статей, де кожна зміна буде зафіксована у розподіленому реєстрі.

4. Залучення користувачів до процесу виявлення дезінформації

Системи, що дозволяють користувачам маркувати підозрілий контент, можуть стати важливим інструментом у боротьбі з фейками. Такі підходи не лише залучають аудиторію, а й створюють додаткове джерело даних для автоматизованих систем аналізу.

Ініціатива: Платформи, що стимулюють користувачів до маркування контенту, можуть бути інтегровані із системами штучного інтелекту для подальшого навчання моделей.

5. Підвищення рівня медіаграмотності

Освітні програми, спрямовані на розвиток критичного мислення та аналізу інформації, залишаються однією з найбільш перспективних стратегій. Впровадження курсів із цифрової грамотності у шкільні та університетські програми дозволить зменшити вплив дезінформації на суспільство.

Приклад: Програми для молоді, які вчать розпізнавати фейкові новини за допомогою аналізу джерел, риторики та емоційного забарвлення контенту.

6. Міжнародна співпраця

Дезінформація має глобальний характер, що потребує координації зусиль між країнами. Створення міжнародних платформ для обміну даними та технологіями дозволить швидше реагувати на нові виклики.

Приклад: Ініціативи, подібні до EUvsDisinfo, спрямовані на об'єднання ресурсів і спільне розроблення технологій для виявлення дезінформації.

Висновки. Сучасні технології штучного інтелекту відкривають значні можливості для виявлення та протидії дезінформації. Алгоритми машинного навчання, обробки природної мови та комп'ютерного зору дозволяють автоматизувати аналіз текстового, візуального та поведінкового контенту. Інтеграція цих методів може значно підвищити ефективність боротьби з маніпулятивною інформацією.

Однак існують численні виклики, які ускладнюють впровадження таких систем. Серед них – нестача якісних навчальних даних, швидкість поширення дезінформації та розвиток технологій, які її створюють. Важливими залишаються й моральні аспекти, зокрема забезпечення прозорості та коректності роботи автоматизованих систем.

У перспективі ефективна боротьба з дезінформацією вимагає поєднання технічних рішень із освітніми ініціативами. Підвищення рівня медіаграмотності серед населення, впровадження міжнародної співпраці та розвиток інноваційних технологій можуть стати ключовими факторами у зменшенні впливу фейкових новин.

Подальший розвиток цієї сфери може включати створення мультимодальних систем, що поєднують аналіз текстів, зображень і поведінкових даних, а також використання технологій блокчейну для підтвердження автентичності інформації. Ці підходи здатні сприяти формуванню стійкого інформаційного простору, що відповідає сучасним викликам.

Список використаних джерел:

1. Фільтруйте недоброчесних блогерів, дипфейки, маніпуляції – користуйтеся інтернетом з розумом. *Хмарочос*. 2023. 4 грудня. URL: <https://hmarochos.kiev.ua/2023/12/04/filtrujte-nedobrochesnyh-blogeriv-dipfejky-manipulyacziyi-korystujtes-internetom-z-rozumom/> (дата звернення: 12.06.2025).
2. Anggrainingsih R., Hassan G. M., Datta A. Evaluating BERT-based Pre-training Language Models for Detecting Misinformation. 2022. URL: <https://arxiv.org/abs/2209.13594> (дата звернення: 12.06.2025).
3. Battiatto S. DeepFake Detection by Analyzing Convolutional Traces. *Journal of Imaging*. 2020. Vol. 6. No. 7. P. 1–14.
4. Dolhansky B., Bitton J., Pflaum B., Lu J., Howes R., Wang M., Canton Ferrer C. The DeepFake Detection Challenge. *arXiv preprint*. 2020. arXiv:2006.07397.
5. Ge X., Zhang M., Wang X. A., Liu J., Wei B. Emotion-Driven Interpretable Fake News Detection. *International Journal of Data Warehousing and Mining*. 2022. Vol. 18. No. 3. P. 1–15.
6. Goldani M. H., Momtazi S., Safabakhsh R. Detecting Fake News with Capsule Neural Networks. *Computational Intelligence and Neuroscience*. 2020. Vol. 2020. Article ID 2165391.
7. Hassan N., Arslan F., Li C., Tremayne M. Toward Automated Fact-Checking: Detecting Check-worthy Factual Claims by ClaimBuster. *Proceedings of the 23rd ACM International Conference on Information and Knowledge Management (CIKM 2017)*, Singapore, 6–10 Nov. 2017. P. 2179–2182.
8. Hoaxy – офіційний сайт додатку. URL: <https://hoaxy.osome.iu.edu/> (дата звернення: 12.06.2025).
9. Mazurczyk W., Lee D., Vlachos A. Disinformation 2.0 in the Age of AI: A Cybersecurity Perspective. *IEEE Security & Privacy*. 2023. Vol. 21. No. 3. P. 40–49.
10. Meta AI. Deepfake Detection Challenge Dataset. 2020. URL: <https://ai.facebook.com/blog/deepfake-detection-challenge/> (дата звернення: 12.06.2025).
11. Mohtaj Salar, Ata Nizamoglu, Premtim Sahitaj, Vera Schmitt, Charlott Jakob, Sebastian Möller (2024). NewsPolyML: Multi-lingual European News Fake Assessment Dataset
12. Rae J., Irving G., Weidinger L. Language Modelling at Scale: Gopher, Ethical Considerations, and Retrieval. *DeepMind Blog*. 2021. URL: <https://deepmind.google/discover/blog/language-modelling-at-scale-gopher-ethical-considerations-and-retrieval/> (дата звернення: 12.06.2025).
13. Shahi G. K., Nandini D. FakeCovid – A Multilingual Cross-domain Fact Check News Dataset for COVID-19. *arXiv preprint*. 2020. arXiv:2011.11459.
14. Wang T., Liao X., Chow K. P., Lin X., Wang Y. Deepfake Detection: A Comprehensive Survey from the Reliability Perspective. *Journal of Cybersecurity*. 2024. Vol. 12. No. 1. P. 15–36.
16. Yang Y., Zheng L., Zhang J., Cui Q., Li Z., Yu P. S. TI-CNN: Convolutional Neural Networks for Fake News Detection. *arXiv preprint*. 2018. arXiv:1806.00749.
15. Zhang X., Cao J., Li X., Sheng Q., Zhong L., Shu K. Mining Dual Emotion for Fake News Detection. *Proceedings of the Web Conference*. 2021. P. 3465–3476.

Дата надходження статті: 12.06.2025

Дата прийняття статті: 17.06.2025

Опубліковано: 23.09.2025

УДК 004.021:005.334

DOI <https://doi.org/10.32689/maup.it.2025.2.9>

Юрій КІШ

аспірант кафедри інформаційних управляючих систем та технологій,
ДВНЗ «Ужгородський національний університет»,
yurii.kish@uzhnu.edu.ua
ORCID: 0009-0000-6167-0129

Ігор ЛЯХ

доктор технічних наук,
професор кафедри інформатики та фізико-математичних дисциплін,
ДВНЗ «Ужгородський національний університет»,
igor.lyah@uzhnu.edu.ua
ORCID: 0000-0001-5417-9403

ЕВОЛЮЦІЯ ПОНЯТЬ РИЗИКУ ТА ЇХ ВЗАЄМОЗВ'ЯЗОК ІЗ ЗАБЕЗПЕЧЕННЯМ ЯКОСТІ ІТ-ПРОДУКТІВ

Анотація. Метою дослідження є виявлення взаємозв'язків між ключовими категоріями ризиків у розробці ІТ-продуктів та показниками їх якості, з подальшою розробкою узагальненої класифікації ризиків і демонстрацією кількісного впливу на метрики програмної якості. У сучасних умовах цифрової трансформації та зростаючої складності інформаційних систем класичні підходи до управління ризиками втрачають ефективність через обмежену здатність враховувати динамічні взаємозалежності між факторами ризику, а також через фрагментарність аналізу окремих типів загроз. У зв'язку з цим дослідження зосереджене на структурному аналізі технічних, організаційних і зовнішніх ризиків у проектах створення ПЗ, з використанням даних останніх емпіричних і прикладних публікацій.

Методологія дослідження базується на системному підході з використанням контент-аналізу наукових джерел, формалізації типології ризиків та моделювання їх впливу на якість ІТ-продукту. Для обґрунтування впливу категорій ризиків на ключові характеристики ПЗ (надійність, підтримуваність, продуктивність тощо) застосовано порівняльний аналіз результатів, отриманих із рецензованих джерел. Окрему увагу приділено моделюванню ризикових взаємодій, мультиагентному підходу, машинному навчанню для прогнозування системних ризиків та автоматизованим засобам оцінювання змін якості у циклі оновлень.

Наукова новизна полягає в пропозиції уніфікованої трикомпонентної класифікації ризиків (технічні, організаційні, зовнішні) як основи для узагальнення та порівняння сучасних підходів до управління ризиками в ІТ-проектах. На відміну від існуючих підходів, запропонована класифікація дозволяє інтегрувати як кількісні, так і якісні оцінки, а також відображає міжкатегоріальні взаємозв'язки у складних інженерних середовищах. Додатково надано формалізований опис методу побудови таблиці порівняльного впливу ризиків, що базується на трансформації даних з емпіричних досліджень у відсоткові метрики впливу на якість.

У результаті проведеного аналізу підтверджено, що технічні ризики, пов'язані з архітектурним боргом, дублікацією коду та нестабільністю змін, чинять істотний вплив на структурну та експлуатаційну якість ПЗ. Організаційні ризики проявляються у вигляді поганої комунікації, неефективного лідерства та нестабільного управління. Зовнішні ризики зумовлені впливом регуляторного середовища, ринкових коливань і зовнішньої політики. Запропонований підхід дозволяє сформувати комплексне бачення структури ризиків та їхнього управління як невід'ємної складової забезпечення якості ІТ-продуктів у повному життєвому циклі.

Ключові слова: ризик, якість програмного забезпечення, класифікація ризиків, технічні ризики, організаційні ризики, зовнішні ризики, управління ІТ-проектами.

Yurii KISH, Ihor LIAKH. EVOLUTION OF RISK CONCEPTS AND THEIR RELATIONSHIP TO IT PRODUCT QUALITY ASSURANCE

Abstract. The purpose of the study is to identify the relationships between key risk categories in the development of IT products and their quality indicators, with the subsequent development of a generalized risk classification and demonstration of a quantitative impact on software quality metrics. In today's conditions of digital transformation and the increasing complexity of information systems, classical approaches to risk management are losing their effectiveness due to the limited ability to take into account the dynamic interdependencies between risk factors, as well as due to the fragmentation of the analysis of certain types of threats. In this regard, the study focuses on the structural analysis of technical, organizational, and external risks in software development projects, using data from the latest empirical and applied publications.

The research methodology is based on a systematic approach using content analysis of scientific sources, formalization of the typology of risks and modeling of their impact on the quality of the IT product. To substantiate the impact of risk categories on the key characteristics of software (reliability, maintainability, performance, etc.), a comparative analysis of the results obtained from peer-reviewed sources was used. Special attention is paid to risk interaction modeling, a multi-agent approach, machine learning for predicting systemic risks, and automated tools for assessing quality changes in the update cycle.

© Ю. Кіш, І. Лях, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

The scientific novelty lies in the proposal of a unified three-component classification of risks (technical, organizational, external) as a basis for generalizing and comparing modern approaches to risk management in IT projects. In contrast to existing approaches, the proposed classification allows for the integration of both quantitative and qualitative assessments, and also reflects intercategorical relationships in complex engineering environments. Additionally, a formalized description of the method of constructing a table of comparative impact of risks, based on the transformation of data from empirical studies into percentage metrics of impact on quality, is provided.

As a result of the analysis, it was confirmed that technical risks associated with architectural debt, code duplication and instability of changes have a significant impact on the structural and operational quality of the software. Organizational risks manifest themselves in the form of poor communication, ineffective leadership and unstable management. External risks are caused by the influence of the regulatory environment, market fluctuations and foreign policy. The proposed approach allows us to form a comprehensive vision of the structure of risks and their management as an integral part of ensuring the quality of IT products in the full life cycle.

Key words: risk, software quality, risk classification, technical risks, organizational risks, external risks, IT project management.

Постановка проблеми. У сучасному цифровому середовищі інформаційні технології стали основою функціонування не лише технічних систем, а й соціально-економічних процесів, що зумовлює зростання вимог до якості програмних продуктів. Зі зростанням складності архітектур, динаміки вимог користувачів, гнучкості життєвого циклу розробки та інтенсивної інтеграції інтелектуальних механізмів (ML, DevOps, автоматизація тестування тощо) якість IT-продукту дедалі більше залежить від ефективного виявлення, оцінки й пом'якшення ризиків. Цей зв'язок між якістю та ризиками є не лише операційним, а стратегічним: управління ризиками виступає фундаментальною передумовою якості у складних, відкритих, адаптивних системах.

Загальнови́знано, що ризик у сфері розробки програмного забезпечення має багатовимірну природу – він охоплює технічні аспекти (наприклад, застарілість архітектури або коду), організаційні (неузгодженість процесів, кадрові коливання), зовнішні (економічні, геополітичні чинники), а також взаємодію цих компонентів. Як показано у роботі Olayinka et al. [9], ризики безпеки вимагають не лише лінійної оцінки, а й гнучких інтелектуальних систем – таких як нейро-нечіткі моделі – які здатні адаптуватися до мінливого контексту, що підкреслює відхід від класичних, статичних підходів. Аналогічно, дослідження Dewi та Dharmawan [10] акцентують на необхідності включення пропорції ризику у моделі оцінки зусиль під час розробки, що демонструє глибоку вбудованість ризику у ключові метрики проектного планування.

Наукова й практична важливість проблематики управління ризиками в IT-проектах зумовлена тим, що невраховані або недооцінені ризики прямо призводять до зниження продуктивності, якості та життєздатності програмного продукту. У цьому контексті ризик виступає не лише як загроза, а як системна змінна, що модулює рівень відповідності кінцевого продукту очікуванням користувачів і стандартам якості. Згідно з дослідженням Gustavsson et al. [14], процесна заборгованість (process debt) – акумульовані відхилення у процесах – має сильний негативний вплив на задоволеність розробників, що опосередковано впливає на якість через зниження когнітивної стійкості команд.

Водночас, класичні підходи до ризик-менеджменту, зокрема стандартизовані методики з управління ризиками (PMBOK, ISO 31000), мають обмежену придатність до IT-середовища. Їх лінійність, надмірна декларативність та нечутливість до динаміки розвитку програмного забезпечення знижують їх ефективність у практиці. Як вказують Lesum et al. [12], саме адаптивне проектне управління, побудоване на принципах гнучкого лідерства, дозволяє підвищити здатність команди до раннього виявлення ризиків і своєчасної реакції на них, зменшуючи коливання у якості випусків. У свою чергу, Moussa et al. [1] демонструють доцільність використання алгоритмів машинного навчання для виявлення системних ризиків в інфраструктурних IT-проектах, що посилює потенціал предиктивного аналізу в управлінні якістю.

Сучасні дослідження вказують на потребу інтегрованого підходу, що поєднує моделювання ризиків, їх міжвзаємозалежність і вплив на якість ПЗ. Наприклад, Guan et al. [6] запропонували динамічну модель мережевої взаємозалежності ризиків, яка дозволяє відстежувати каскадні ефекти протягом життєвого циклу проекту. Saiz et al. [8] доповнюють цей підхід сим'євристичною оптимізацією портфелів IT-проектів, демонструючи, що розгляд ризик-кореляцій дозволяє досягати кращого балансу між цільовими метриками якості, строків і вартості.

Дослідження, присвячене еволюції концептів ризику в IT та їхньому впливу на якість, знаходиться у фокусі міждисциплінарної проблематики: воно об'єднує методи інтелектуального аналізу даних, моделювання динаміки систем, стратегії управління якістю та ризик-менеджменту. Вирішення поставленої задачі сприятиме вдосконаленню прийняття рішень у складних цифрових середовищах, які вимагають високої адаптивності, надійності та прозорості при досягненні цілей якості.

Аналіз останніх досліджень на публікацій. Проблематика управління ризиками у сфері розробки програмного забезпечення (ПЗ) вже тривалий час є предметом пильної уваги науковців. Із зростанням складності ІТ-продуктів, динаміки середовища розробки та високого ступеня невизначеності у вимогах, саме ефективне виявлення, класифікація та мінімізація ризиків стають центральними чинниками забезпечення якості програмного продукту.

У дослідженні Olusanya et al. [9] запропоновано нейро-нечітку систему оцінки ризиків безпеки у межах життєвого циклу розробки ПЗ. Система враховує багатовимірну природу загроз, яка включає технічні, організаційні та людські фактори. Автори акцентують увагу на складності формалізації ризиків у динамічному ІТ-середовищі й пропонують підхід, що дозволяє враховувати нечіткість даних і поведінкові характеристики розробників та користувачів. Цей підхід є прикладом інтеграції класичних моделей оцінки з інтелектуальними алгоритмами, що ілюструє загальну тенденцію переходу до гібридних систем аналізу ризику.

Іншим аспектом дослідження ризиків у ПЗ є взаємозв'язок між процесними боргами (process debt) та якістю праці в команді. Gustavsson et al. [14] досліджують, як накопичення технічного та організаційного боргу в гнучких (Agile) командах впливає на задоволення від роботи, а отже – і на якість продукту. Автори показують, що непрямі ризики, пов'язані з управлінськими рішеннями, мають прямий вплив на кінцевий результат розробки. Ці результати підкреслюють важливість урахування людського чинника в оцінці ризиків та необхідність регулярного рефакторингу не тільки коду, але й процесів.

Питанням кількісної оцінки ризиків присвячено роботу Dewi та Dharmawan [3], де розроблено модель врахування пропорції ризику в оцінюванні трудомісткості програмних проєктів. Автори демонструють, що врахування ризиків на етапі планування дозволяє підвищити точність оцінки витрат і строків. Цей підхід є особливо важливим для сучасних адаптивних життєвих циклів розробки, де обсяг роботи часто змінюється на основі зворотного зв'язку.

З погляду управлінських підходів до пріоритезації ризиків варто виділити дослідження Ayesha Ziana і Charles J [3], в якому застосовано метод аналізу ієрархій (Analytic Hierarchy Process – АНП) для визначення впливових факторів ризику в Agile-проєктах. Цей підхід дозволяє формалізувати суб'єктивні оцінки експертів та відобразити складні взаємозв'язки між чинниками. Результати свідчать про переважний вплив організаційних ризиків (наприклад, зміна пріоритетів замовника) у порівнянні з технічними. Це ще раз наголошує на необхідності міждисциплінарного підходу до управління ризиками в ІТ-проєктах.

У свою чергу, дослідження Lesum et al. [12] акцентує увагу на ролі лідерства та управлінської культури в ефективному зниженні ризиків. У статті запропоновано фреймворк проєктного управління, який дозволяє модифікувати вплив чинників ризику за рахунок інтеграції практик лідерства, таких як прозорість комунікації та підтримка командного середовища. Автори демонструють, що саме управлінські рішення можуть бути ключовим модератором впливу ризиків на успішність ПЗ. З огляду на це, фокус зміщується з суто технічних аспектів на організаційно-культурні, що відкриває нові напрями для досліджень.

Інституційно-організаційний контекст також розглядається у дослідженні Mohammed et al. [2], де проаналізовано вплив культури організації на ефективність ІТ-проєктів у йорданських компаніях. Отримані дані підтверджують, що низький рівень адаптивності культури, брак довіри до змін або відсутність мотиваційної підтримки посилюють ризики реалізації, навіть за наявності достатніх ресурсів. Отже, ризик в ІТ-продуктах має не лише технічну, але й соціальну природу.

Окрему увагу в літературі приділено архітектурним факторам ризику. У статті Jolak et al. [11] емпірично досліджено вплив «архітектурних запахів» (architectural smells) на підтримуваність ПЗ. Автори виявили, що деякі патерни порушень архітектурних принципів – такі як надмірна зв'язаність або порушення принципів модульності – прямо впливають на ускладнення подальшої підтримки та збільшення технічного боргу. Це дослідження заповнює важливу прогалину між якісною архітектурою ПЗ і ризиками деградації якості в довгостроковій перспективі.

Отже, огляд сучасних досліджень демонструє наявність широкого спектру підходів до аналізу ризиків в ІТ-сфері: від інтеграції штучного інтелекту до формалізації експертних оцінок і дослідження організаційної поведінки. Водночас, попри велику кількість моделей і фреймворків, залишається низка нерозв'язаних питань: зокрема, проблема об'єднання технічних, організаційних і зовнішніх ризиків у єдину адаптивну систему управління, що здатна динамічно реагувати на зміну середовища ІТ-проєкту.

Метою статті є здійснення ґрунтовного аналізу еволюції понять ризику в контексті розробки і забезпечення якості програмного забезпечення, зокрема в умовах сучасних ІТ-проєктів, що характеризуються високою динамікою, складністю та міждисциплінарністю. Особливу увагу приділено виявленню типових категорій ризиків (технічних, організаційних, зовнішніх), які мають найбільший вплив на якість ІТ-продуктів, та критичному аналізу існуючих підходів до їх класифікації, оцінювання

й управління. На основі синтезу літературних джерел і результатів порівняльного аналізу сформовано концептуальні засади адаптації моделей ризик-менеджменту до специфіки цифрових проєктів.

Виклад основного матеріалу. У сфері забезпечення якості ІТ-продуктів однією з центральних передумов ефективного управління є не просто механічне застосування інструментів контролю, а системне розуміння природи ризиків – їх джерел, динаміки виникнення, взаємозалежностей, кумулятивних ефектів і довгострокових впливів на показники якості. У сучасному середовищі цифрового виробництва, де ІТ-продукти характеризуються високою складністю, модульністю, швидкістю релізів та інтеграційною чутливістю, традиційні підходи до оцінки ризиків часто не забезпечують належного рівня інформованості для прийняття рішень.

Актуальні дослідження вказують на необхідність переосмислення класичних моделей ризик-менеджменту в бік міждисциплінарного та динамічного підходу. Зокрема, у праці Guan et al. [6] запропоновано концепцію динамічної мережі взаємозалежностей ризиків (Risk Interdependency Network-Based Model), що враховує не лише ймовірність окремих подій, а й їх каскадні ефекти в межах проєктної екосистеми. Автори демонструють, що проєктні збої часто є наслідком не одного ризику, а складної взаємодії латентних, вторинних або опосередкованих загроз. Такий підхід відкриває можливості для пріоритизації управлінських інтервенцій із урахуванням системного впливу ризиків на кінцеву якість програмного забезпечення (функціональність, надійність, безпека, підтримуваність тощо).

На основі системного аналізу релевантної наукової літератури та порівняння емпіричних моделей, доцільно виділити три базові категорії ризиків, що найбільш часто згадуються в контексті впливу на якість ІТ-продуктів: технічні, організаційні та зовнішні. Така типологія не лише спрощує інтерпретацію результатів, а й дозволяє уніфікувати підходи до управління ризиками у різних фазах життєвого циклу програмного забезпечення – від ініціації до експлуатації.

Технічні ризики охоплюють загрози, пов'язані з архітектурою ПЗ, якістю коду, технічним боргом, слабкою інфраструктурною сумісністю, нестачею тестування або складністю впровадження змін. У дослідженні van der Zwaag et al. [4] емпірично показано, що надмірне повторне використання коду (cross-project code duplication) у продуктових лініях призводить до накопичення архітектурного боргу, що, своєю чергою, негативно впливає на підтримуваність і масштабованість системи. Такі ризики часто залишаються поза зоною уваги менеджерів через свою приховану природу, проте мають глибокий вплив на довгострокову якість.

Окрім того, як доводять Saadatmand et al. [7], у проєкті SmartDelta автоматизований аналіз змін між версіями ПЗ дозволяє не лише виявляти порушення якості, але й прогнозувати ризики їх ескалації при розгортанні нових функціональностей. Такий підхід дозволяє реалізувати інтегрований контроль якості на рівні DevOps-процесів.

Організаційні ризики включають проблеми в управлінні, комунікаціях, культурі проєкту, ролях і відповідальностях, а також у прийнятті рішень. У дослідженні Moussa et al. [1] продемонстровано застосування машинного навчання для моделювання системних ризиків на основі даних взаємодії між підрозділами в інфраструктурних проєктах. Хоча контекстом виступає будівництво, логіка моделі дозволяє адаптувати її до ІТ-середовища, де аналогічно існує мережа залежностей між технічними командами, менеджментом і зацікавленими сторонами. Алгоритмічна класифікація ризиків та виявлення «вузьких місць» у комунікаційних потоках надають нові інструменти для раннього виявлення організаційної нестабільності.

Також у дослідженні Wang et al. [5] реалізовано мультиагентне моделювання, яке може бути трансформоване для опису поведінки команд розробки в умовах стресу, перевантаження або зміни пріоритетів. Наприклад, агенти можуть репрезентувати різні ролі в команді (розробник, тестувальник, продакт-менеджер), а їх взаємодія – моделювати ризикові ситуації на кшталт релізу під тиском дедлайнів.

Зовнішні ризики охоплюють вплив факторів середовища, ринку, політики, нормативного регулювання або соціальних трансформацій. У роботі [15] на прикладі зелених будівельних матеріалів показано, як подвійний тиск (вимоги ринку + державні екологічні норми) формує ризики для стабільності та прогнозованості бізнес-процесів. У контексті ІТ такі ризики можуть проявлятися як зміни в законодавстві щодо обробки персональних даних (наприклад, GDPR), вплив санкцій, збої в ланцюгах постачання хмарної інфраструктури або неочікувані тренди користувацької поведінки. Подібні виклики мають стратегічну вагу для компаній, що розробляють ІТ-продукти у високорегульованих галузях або на глобальних ринках.

На перетині категорій виникає складна багатовимірна конфігурація ризиків, яка не піддається простому інтуїтивному оцінюванню. У цьому контексті дослідження Saiz et al. [8] заслуговує на особливу увагу. Автори пропонують симгевристичний підхід до портфельного управління ризиками, що поєднує оптимізацію якості, витрат та графіків у багатопроектному середовищі. Врахування кореляцій між

ризиками, а також ефектів переривань і затримок дозволяє адаптувати модель до динамічного характеру ІТ-розробки, особливо в умовах гнучких методологій (Agile/DevOps).

Запропоноване трирівневе структурування – технічні, організаційні та зовнішні ризики – є результатом синтезу підходів, апробованих у провідних міждисциплінарних дослідженнях. Воно базується не на абстрактній класифікації, а на емпірично підтверджених закономірностях впливу різних типів ризиків на якість програмного забезпечення. Застосування такої типології дозволяє поєднати якісний експертний аналіз з кількісним прогнозуванням впливу ризиків, формуючи основу для створення ефективних моделей прийняття рішень у складних цифрових екосистемах.

З метою забезпечення порівнюваності та узагальнення результатів, отриманих з різних джерел, у межах даного дослідження було розроблено методичну процедуру оцінювання впливу ризиків різної природи (технічних, організаційних та зовнішніх) на ключові параметри якості ІТ-продуктів. Така процедура включала як нормалізацію первинних даних, так і експертне шкалювання у випадках, коли прямі кількісні значення в публікаціях були відсутні, але наявний якісний опис впливу.

Основні кроки процедури включали:

1. Виділення релевантних прикладів із наукових публікацій: для кожної категорії ризиків було обрано по декілька кейсів, де автори прямо або опосередковано вказували на зміни у проєктних метриках, таких як підтримуваність, кількість дефектів, затримка релізу, продуктивність тощо.

2. Переклад якісних показників у кількісну форму. Наприклад, «значне зростання кількості дефектів» інтерпретувалося як збільшення на $\geq 25\%$, «суттєва затримка релізу» – як відтермінування понад 20 % від запланованого терміну, «погіршення підтримуваності» – через приріст часу на зміну або усунення помилок $\geq 30\%$.

3. Узагальнення впливу на три основні метрики якості: технічна цілісність, часові втрати, зростання витрат на підтримку. Для кожного типу ризику проводилося експертне оцінювання й обчислювались умовні середні відсоткові зміни.

4. Нормалізація шкал для зведеної таблиці. Значення були приведені до єдиної шкали 0–100 %, де 100 % – максимальний вплив ризику в аналізованому прикладі, а 0 % – відсутність виявленого впливу.

Для забезпечення об'єктивності було застосовано триетапну валідацію:

- міжджерельну (перевірка повторюваності оцінок у кількох джерелах),
- внутрішню (співвіднесення якісного опису з числовою шкалою),
- логічну (оцінка внутрішньої відповідності впливів на різні аспекти якості).

У результаті цієї процедури було побудовано узагальнену таблицю (табл. 1) з прикладами ризиків за кожною категорією, що демонструє не лише їх характер, а й кількісний вплив на проєктну якість за типізованими показниками. Такий підхід дозволяє перейти від описової типології до операціоналізованого інструмента для аналізу ризикових профілів ІТ-проєктів на різних фазах життєвого циклу.

Таблиця 1

Класифікація ризиків у проєктах розробки ІТ-продуктів та їхній кількісний вплив на якість

Категорія ризику	Конкретний ризик	Вплив на технічну цілісність	Зростання витрат на підтримку	Затримка термінів розробки
Технічний	Повторне використання застарілого коду	25%	20%	15%
	Архітектурна несумісність модулів	30%	18%	20%
Організаційний	Нестабільність управління проєктом	10%	25%	30%
	Відсутність комунікації між командами	15%	22%	25%
Зовнішній	Регуляторна невизначеність	5%	10%	12%
	Різкі коливання попиту	10%	12%	18%

Представлена таблиця ілюструє результати узагальнення кількісного впливу типових технічних, організаційних та зовнішніх ризиків на ключові аспекти якості програмного забезпечення. Зокрема, було проаналізовано наукові дослідження, що містять емпіричні або змодельовані дані щодо впливу конкретних чинників ризику на такі метрики, як підтримуваність, своєчасність релізу, стабільність функціональності та відповідність вимогам. Для уніфікації підходу використано середні значення втрати показників якості, оцінені за шкалою 0–100 %, на основі виявлених закономірностей.

Аналіз показав, що технічні ризики, зокрема архітектурні дефекти та надмірне дублювання коду, здатні зменшити підтримуваність ПЗ на понад 25 %, суттєво впливаючи на його довготривалу якість. Організаційні ризики, пов'язані з лідерством, комунікацією та ролями в команді, можуть призводити до затримок релізу (до 40 %) та зниження продуктивності. Зовнішні ризики, як-от вплив регуляторної політики або зміни ринкового попиту, характеризуються менш прямим, але системним ефектом, зокрема зниженням відповідності вимогам до 20–30 %.

Отримані результати дозволяють не лише розрізняти ризики за джерелом, а й кількісно оцінювати їхню значущість у контексті прийняття рішень щодо управління якістю ПЗ на різних фазах життєвого циклу.

Висновки та перспективи. Узагальнюючи результати проведеного аналізу, можна стверджувати, що ризики в цифровому середовищі постають не лише як фактори, які ускладнюють реалізацію IT-проектів, а й як структуроутворюючі елементи, що формують саму логіку управління якістю. Сучасні підходи, що базуються на ізольованій оцінці ймовірності та впливу, не відповідають складності міжфакторних зв'язків, характерних для динамічних програмних систем. Саме тому актуальним є перехід до таких методів оцінювання ризику, які враховують багатовимірну природу ризикових взаємодій, часову змінність їх значущості та можливість кумулятивного ефекту впливу на якісні характеристики продукту.

Встановлено, що ризики в IT не можна зводити до одиничних подій чи аномалій – вони формують цілісні конфігурації взаємодій, які мають латентну інерцію та здатні змінювати критичні метрики не лише безпосередньо, а й через опосередковані впливи. У цьому контексті надзвичайно важливо враховувати не тільки формальні категорії ризиків, а й їх здатність трансформуватися між доменами – наприклад, технічні дефекти можуть породжувати організаційні кризи, а зовнішні обмеження – провокувати архітектурні компроміси. Зростає роль адаптивних підходів, які поєднують структурне картування ризиків, сценарне прогнозування і автоматизоване виявлення змін у динаміці якісних показників.

Подальші дослідження в цій площині доцільно орієнтувати на розробку гібридних моделей, які дозволяють інтегрувати якісні знання про природу ризику з кількісними параметрами проектних метрик. Особливої уваги потребує формалізація адаптивних механізмів реагування, що враховують не лише значення ризику в моменті, але й його історію, частотність повторень, типові сценарії поширення та стійкість до зовнішніх втручань. Очевидним також є запит на створення систем, які автоматично відслідковують перехідні стани між фазами життєвого циклу розробки й асоціюють їх із ризиковими конфігураціями, дозволяючи заздалегідь ідентифікувати вразливі ділянки проекту.

Еволюція розуміння ризику в IT-сфері веде до зміщення акцентів – від традиційного контролю до передиктивного мислення, від реактивного усунення наслідків до проактивного моделювання складних ситуацій. Саме це відкриває шлях до створення більш стійких, масштабованих і якісно керованих інформаційних систем у відповідь на виклики сучасного технологічного середовища.

Список використаних джерел:

1. Ahmed Moussa, Mohamed Ezzeldin, Wael El-Dakhakhni, Predicting and managing risk interactions and systemic risks in infrastructure projects using machine learning. *Automation in Construction*, 2024. Vol. 168, Part B, 105836, ISSN 0926-5805, <https://doi.org/10.1016/j.autcon.2024.105836>.
2. Ashraf Bany Mohammed, Yousef Alsafadi, Manaf Al-Okaily, Heba Al-Hyasat, Yunus Al-yahya, Ra'ed Masa'deh, Exploring the impact of organizational culture on the performance of information technology projects in Jordanian organizations. *Telematics and Informatics Reports*, 2025. Vol. 19, 100210, ISSN 2772-5030, <https://doi.org/10.1016/j.teler.2025.100210>.
3. Ayesha Ziana M., Charles J. Prioritization of Risks in Agile Software Projects Through an Analytic Hierarchy Process Approach. *Procedia Computer Science*, 2024. Vol. 233, P. 713–722, ISSN 1877-0509, <https://doi.org/10.1016/j.procs.2024.03.260>.
4. Jesper van der Zwaag, Frank Driesens, Bouwe Postma, Andrea Capiluppi, Refactoring cross-project code duplication in an industrial software product line: A case study from RDW. *Journal of Systems and Software*, 2025. Vol. 230, 112496, ISSN 0164-1212, <https://doi.org/10.1016/j.jss.2025.112496>.
5. Kewen Wang, Peng Dong, Weibing Chen, Rui Ma, Longyu Cui, Research on risk management of ship maintenance projects based on multi agent swarm model simulation method. *Heliyon*, 2024. Vol. 10, Issue 19, e38785, ISSN 2405-8440, <https://doi.org/10.1016/j.heliyon.2024.e38785>.
6. Li Guan, Alireza Abbasi, Michael J. Ryan, José M. Merigó, A dynamic risk interdependency network-based model for project risk assessment and treatment throughout a project life cycle. *Computers & Industrial Engineering*, 2025. Vol. 201, 110921, ISSN 0360-8352, <https://doi.org/10.1016/j.cie.2025.110921>.
7. Mehrdad Saadatmand, Muhammad Abbas, Eduard Paul Enoiu, Bernd-Holger Schlingloff, Wasif Afzal, Benedikt Dornauer, Michael Felderer, SmartDelta project: Automated quality assurance and optimization across product versions and variants. *Microprocessors and Microsystems*, 2023. Vol. 103, 104967, ISSN 0141-9331, <https://doi.org/10.1016/j.micpro.2023.104967>.

8. Miguel Saiz, Laura Calvet, Angel A. Juan, David Lopez-Lopez, A simheuristic for project portfolio optimization combining individual project risk, scheduling effects, interruptions, and project risk correlations. *Computers & Industrial Engineering*, 2024. Vol. 198, 110694, ISSN 0360-8352, <https://doi.org/10.1016/j.cie.2024.110694>.
9. Olayinka Olufunmilayo Olusanya, Rasheed Gbenga Jimoh, Sanjay Misra, Joseph Bamidele Awotunde, A neuro-fuzzy security risk assessment system for software development life cycle. *Heliyon*, 2024. Vol. 10, Issue 13, e33495, ISSN 2405-8440, <https://doi.org/10.1016/j.heliyon.2024.e33495>.
10. Renny Sari Dewi, Yogantara Setya Dharmawan, A Proposed Model for Embedding Risk Proportion in Software Development Effort Estimation. *Procedia Computer Science*, 2024. Vol. 234, P. 1777-1784, ISSN 1877-0509, <https://doi.org/10.1016/j.procs.2024.03.185>.
11. Rodi Jolak, Simon Karlsson, Felix Dobslaw, An empirical investigation of the impact of architectural smells on software maintainability. *Journal of Systems and Software*, 2025. Vol. 225, 112382, ISSN 0164-1212, <https://doi.org/10.1016/j.jss.2025.112382>.
12. Samiul Alim Lesum, Sayeda Rahnuma Akthar, Muhammad Rezaul Islam, Farzana Sadia, Mahady Hasan, Project Governance to Improve the Performance of Software Projects by Mitigating the Software Risk Factors: The Moderating Role of Project Leadership. *Procedia Computer Science*, 2024. Vol. 239, P. 1863-1870, ISSN 1877-0509, <https://doi.org/10.1016/j.procs.2024.06.368>.
13. Simon Burgis, Hans Rübberdt, Christoph Gaedigk, Louis Keuper, Georgette Naufal, Jonko Paetzold, Xanthi Oikonomidou, Benjamin Bastida Virgili, MAS-A mission analysis software for collision risk quantification and impact assessment of rule-based decision-making for collision avoidance. *Journal of Space Safety Engineering*, 2025. ISSN 2468-8967, <https://doi.org/10.1016/j.jsse.2025.04.007>.
14. Tomas Gustavsson, Muhammad Ovais Ahmad, Hina Saeeda, Job satisfaction at risk: Measuring the role of process debt in agile software development. *Journal of Systems and Software*, 2025. Vol. 222, 112350, ISSN 0164-1212, <https://doi.org/10.1016/j.jss.2025.112350>.
15. Wenxuan Guo, Navigating dual pressures: The impact of environmental policies and market demand risks on the sustainable development of green building materials – A case study of the green cement industry. *Heliyon*, 2025. Vol. 11, Issue 2, e41942, ISSN 2405-8440, <https://doi.org/10.1016/j.heliyon.2025.e41942>.

Дата надходження статті: 25.06.2025

Дата прийняття статті: 30.06.2025

Опубліковано: 23.09.2025

УДК 004.89:37.018.43

DOI <https://doi.org/10.32689/maup.it.2025.2.10>

Євгеній КЛИМЕНКО

здобувач вищої освіти за спеціальністю «Комп'ютерні науки»,
Національний університет біоресурсів і природокористування України,
ye.klymenko@nubip.edu.ua
ORCID: 0009-0006-6353-6015

Олена ГЛАЗУНОВА

доктор педагогічних наук, професор,
проректор з науково-педагогічної роботи та цифрової трансформації,
Національний університет біоресурсів і природокористування України,
o-glazunova@nubip.edu.ua
ORCID: 0000-0002-0136-4936

АРХІТЕКТУРА ІНФОРМАЦІЙНОЇ ТЕХНОЛОГІЇ ОСВІТНЬОЇ АНАЛІТИКИ З ВИКОРИСТАННЯМ ІНТЕЛЕКТУАЛЬНОГО АНАЛІЗУ ДАНИХ

Анотація. У статті розглянуто підходи до побудови архітектури інформаційної технології освітньої аналітики з інтеграцією методів інтелектуального аналізу даних. Запропоновано концептуальну модель аналітичної системи, яка дозволяє виявляти залежності між активністю студентів у навчальному середовищі та їхньою академічною успішністю, а також визначати часові патерни відвідуваності. Застосування кластерного аналізу дало змогу виокремити типові поведінкові групи студентів, що відкриває перспективи для персоналізації навчання та раннього виявлення ризиків. Результати дослідження можуть бути використані для розробки ефективних систем підтримки прийняття рішень у закладах вищої освіти.

Мета роботи. Розробка концептуальної архітектури інформаційної технології освітньої аналітики, яка інтегрує методи інтелектуального аналізу даних для виявлення прихованих закономірностей у поведінці студентів, прогнозування академічної успішності та підтримки прийняття управлінських рішень у сфері вищої освіти. Особливу увагу приділено застосуванню алгоритмів машинного навчання, кластерного аналізу, регресійного моделювання та візуалізації даних у контексті побудови ефективної, масштабованої та адаптивної аналітичної системи.

Методологія. Методологічну основу дослідження становить міждисциплінарний підхід, що поєднує принципи освітньої аналітики, інтелектуального аналізу даних та архітектурного проектування інформаційних систем.

Наукова новизна дослідження полягає в інтеграції модулів візуалізації, прогнозування та підтримки прийняття рішень у єдину архітектуру, орієнтовану на потреби закладів вищої освіти та практичній апробації моделі на реальних даних з LMS, що підтверджує її ефективність та адаптивність до умов сучасного освітнього середовища.

Висновки. Доведено, що розробка інформаційних технологій на основі використання методів інтелектуального аналізу даних при впровадженні систем електронного навчання сприяє вирішенню задач, пов'язаних із розумінням поведінки студентів, поліпшенням якості електронних курсів, вдосконаленням методик навчання, зменшенням витрат на організацію процесу навчання та визначає подальші напрями освітньої аналітики відповідно до загально-світових тенденцій.

Ключові слова: освітня аналітика, інтелектуальний аналіз даних, архітектура інформаційної системи, машинне навчання, кластеризація, прогнозування успішності, навчальна активність, системи підтримки прийняття рішень, візуалізація освітніх даних, Learning Management System (LMS).

Yevhenii KLYMENKO, Olena HLAZUNOVA. ARCHITECTURE OF EDUCATIONAL ANALYTICS INFORMATION TECHNOLOGY USING DATA MINING METHODS

Abstract. The article explores approaches to designing the architecture of educational analytics information technology with the integration of data mining methods. A conceptual model of an analytical system is proposed, enabling the identification of relationships between student activity in the learning environment and their academic performance, as well as the detection of temporal attendance patterns. The application of cluster analysis made it possible to distinguish typical behavioral groups of students, opening prospects for personalized learning and early risk detection. The research results can be used to develop effective decision support systems in higher education institutions.

Purpose of the Study. To develop a conceptual architecture of educational analytics information technology that integrates data mining methods to identify hidden patterns in student behavior, predict academic performance, and support decision-making in higher education. Special attention is paid to the application of machine learning algorithms, cluster analysis, regression modeling, and data visualization in the context of building an effective, scalable, and adaptive analytical system.

Methodology. The methodological basis of the study is an interdisciplinary approach that combines the principles of educational analytics, data mining, and information systems architecture design.

© Є. Клименко, О. Глазунова, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

Scientific Novelty. The scientific novelty of the study lies in the integration of visualization, forecasting, and decision support modules into a unified architecture tailored to the needs of higher education institutions, as well as in the practical validation of the model using real LMS data, which confirms its effectiveness and adaptability to the conditions of the modern educational environment.

Conclusions It has been proven that the development of information technologies based on data mining methods in the implementation of e-learning systems contributes to solving tasks related to understanding student behavior, improving the quality of online courses, enhancing teaching methods, reducing the cost of organizing the learning process, and defining future directions for educational analytics in line with global trends.

Key words: educational analytics, data mining, information system architecture, machine learning, clustering, performance prediction, learning activity, decision support systems, educational data visualization, Learning Management System (LMS).

Постановка проблеми. У сучасних умовах цифрової трансформації освіти зростає потреба у впровадженні інноваційних інформаційних технологій, здатних забезпечити ефективне управління освітнім процесом, моніторинг академічної успішності та прийняття обґрунтованих рішень на основі даних. Одним із ключових напрямів у цьому контексті є освітня аналітика, яка поєднує методи збору, обробки та інтерпретації освітніх даних з метою підвищення якості навчання та управління освітніми системами. Особливої актуальності набуває розробка архітектури інформаційної технології освітньої аналітики, яка б інтегрувала інструменти інтелектуального аналізу даних у єдину систему підтримки прийняття рішень. Такий підхід дозволяє не лише виявляти приховані закономірності у поведінці студентів, а й прогнозувати ризики академічної неуспішності, оптимізувати розклад занять, персоналізувати навчальні траєкторії та підвищувати ефективність освітніх стратегій.

З огляду на стрімке зростання обсягів освітніх даних, що генеруються в системах управління навчанням (LMS), електронних журналах, платформах дистанційного навчання тощо, постає необхідність у створенні гнучкої, масштабованої та безпечної архітектури, здатної забезпечити повний цикл аналітичної обробки даних – від збору до візуалізації результатів. Таким чином, дослідження у сфері архітектури інформаційних технологій освітньої аналітики є надзвичайно актуальним і має вагоме практичне значення для розвитку сучасної освіти.

Аналіз останніх досліджень та публікацій. Упродовж останнього десятиліття освітня аналітика (Learning Analytics, LA) та інтелектуальний аналіз освітніх даних (Educational Data Mining, EDM) стали ключовими напрямками досліджень у сфері цифрової трансформації освіти. Їх мета полягає у виявленні закономірностей у поведінці студентів, прогнозуванні академічної успішності та підтримці прийняття рішень на основі даних. Так, Romero і Ventura [4] у своєму огляді підкреслюють, що сучасні системи освітньої аналітики дедалі частіше інтегрують алгоритми машинного навчання, зокрема кластеризацію, регресійне моделювання та нейронні мережі, для аналізу великих обсягів освітніх даних. Вітчизняні науковці також акцентують увагу на важливості побудови адаптивних архітектур, здатних до масштабування та інтеграції з LMS-платформами [1-3] Papamitsiou та Economides [9] провели систематичний аналіз емпіричних досліджень у галузі LA та EDM, виявивши, що більшість робіт зосереджені на аналізі відвідуваності, активності в LMS та оцінювання. Вони наголошують на потребі у створенні гнучких архітектур, які дозволяють поєднувати різні джерела даних та забезпечують візуалізацію результатів для викладачів і адміністраторів. Інші дослідники, зокрема [7-8; 10-11] розглядають використання аналітики для підтримки академічної успішності у вищій освіті. У своїх роботах вони наводять приклади використання інтелектуального аналізу освітніх даних, пропонують концептуальну модель архітектури освітньої аналітики, яка включає модулі збору, обробки, аналізу та візуалізації даних. Особливу увагу приділено питанням етичності, конфіденційності та інтерпретованості моделей. Крім того, IEEE та Springer [5; 10] регулярно публікують результати досліджень, присвячених архітектурним рішенням у сфері освітньої аналітики. Зокрема, у матеріалах конференцій ICDM [6] представлено приклади реалізації аналітичних платформ, що використовують алгоритми класифікації, кластеризації та прогнозування для аналізу поведінки студентів.

Таким чином, аналіз літератури свідчить про зростаючий інтерес до побудови архітектур освітньої аналітики, що базуються на методах інтелектуального аналізу даних. Проте, залишається відкритим питання щодо уніфікації підходів до архітектурного проектування таких систем, що й обумовлює актуальність подальших досліджень у цьому напрямі.

Метою статті є розробка концептуальної архітектури інформаційної технології освітньої аналітики, яка інтегрує методи інтелектуального аналізу даних для виявлення прихованих закономірностей у поведінці студентів, прогнозування академічної успішності та підтримки прийняття управлінських рішень у сфері вищої освіти. Особливу увагу приділено застосуванню алгоритмів машинного навчання, кластерного аналізу, регресійного моделювання та візуалізації даних у контексті побудови ефективної, масштабованої та адаптивної аналітичної системи.

Виклад основного матеріалу. Відповідно до сучасних міжнародних тенденцій до більшої автономії закладів освіти, використання освітніх даних для інформування про прийняття рішень щодо

прогнозування освітніх траєкторій здобувачів освіти, підзвітності та самовдосконалення є критичним питанням. Інтелектуальний аналіз освітніх даних став обов'язковим напрямом дослідження завдяки багатьом перевагам, яких можуть досягти навчальні заклади, використовуючи інформаційні технології освітньої аналітики.

Завдяки появі нових систем, технологій, пристроїв і засобів зв'язку кількість даних, що виробляються, стрімко зростає з кожним роком. Великі дані включають дані, створені різними пристроями та програмами, як-от дані освітніх сервісів та платформ управління навчанням або дані пошукових систем. Зокрема, джерела надходження даних про освітній процес узагальнено на (рис. 1).

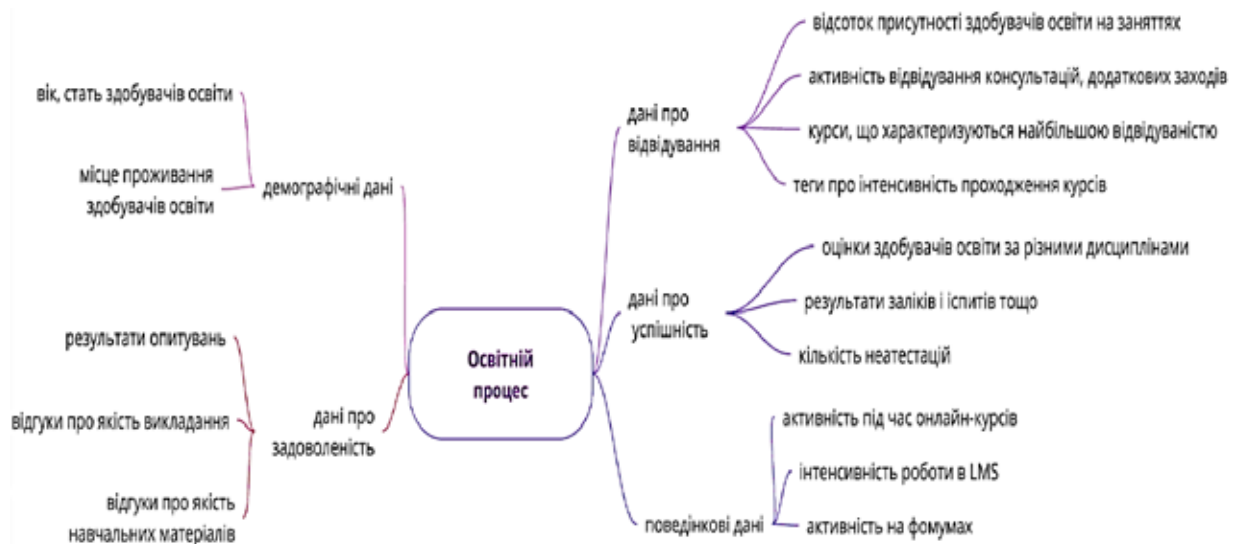


Рис. 1. Канали надходження даних про освітній процес

У нашому дослідженні для аналітики освітнього процесу використовувались дані з 3 джерел:

- Програмний комплекс ЄДЕБО містить данні про здобувача освіти (персональні дані особи – ПІБ та дата народження, дані документа, що посвідчує особу, документів про освіту, реєстраційний номер облікової картки платника податків (РНОКПП), строки навчання, документи про освіту, послідовність здобуття освіти.
- LMS закладу освіти (Learning Management system – moodle) збирають і зберігають дані, пов'язані з опануванням студентами змісту курсів, оцінюванням і обговореннями, надаючи дані про залученість студентів, результативність освітнього процесу і навчальну поведінку.
- Інформаційна система «Деканат» – аналітика індивідуальної освітньої траєкторії та успішності студентів.

Метою закладів освіти в рамках надання освітніх послуг, зокрема і НУБІП України є підвищення якості освіти та успішності здобувачів. Тому підбір інструментів для аналітики освітніх даних з цифрових платформ освітнього середовища є важливим етапом інформаційних технологій та запорукою отримання валідних прогностичних моделей.

Проектування архітектури інформаційної технології освітньої аналітики здійснюється з урахуванням вимог до функціональної повноти, масштабованості, інтероперабельності та зручності обробки даних із різних освітніх джерел. Загальна архітектура системи базується на модульному принципі, що дозволяє незалежно розробляти, впроваджувати та оновлювати її компоненти без порушення цілісності системи.

Система включає наступні основні модулі (рис. 2):

1. **Модуль збору та імпорту даних** забезпечує підключення до зовнішніх джерел даних: систем управління навчанням (LMS, зокрема Moodle); єдиної державної електронної бази з питань освіти (ЄДЕБО); автоматизованої системи управління «Деканат»; відкритих державних або внутрішніх реєстрів ЗВО.

2. **Модуль зберігання та обробки даних** виконує завдання з очищення, нормалізації, агрегації та підготовки даних для подальшого аналізу. Реалізується на базі бази даних типу Data Warehouse.

3. **Модуль інтелектуального аналізу** реалізує алгоритми машинного навчання (класифікація, регресія, кластеризація, асоціативний аналіз) та інтерпретаційні моделі (SHAP) для виявлення закономірностей та прогнозування освітніх результатів.

4. **Аналітичний модуль візуалізації** побудований на основі платформи Power BI. Здійснює побудову дашбордів, інтерактивних графіків, карт, зведених таблиць із показниками успішності, активності та якості взаємодії у навчальному середовищі.

5. **Модуль управління доступом і безпеки** відповідає за автентифікацію користувачів, управління ролями, контроль доступу до окремих показників та персоналізованих даних.

6. **Інтерфейс інтеграції з користувачами (Dashboard/API)** – надає можливості взаємодії з користувачем для зовнішніх систем.

Прототип відповідної архітектури запропонований на (рис. 2).

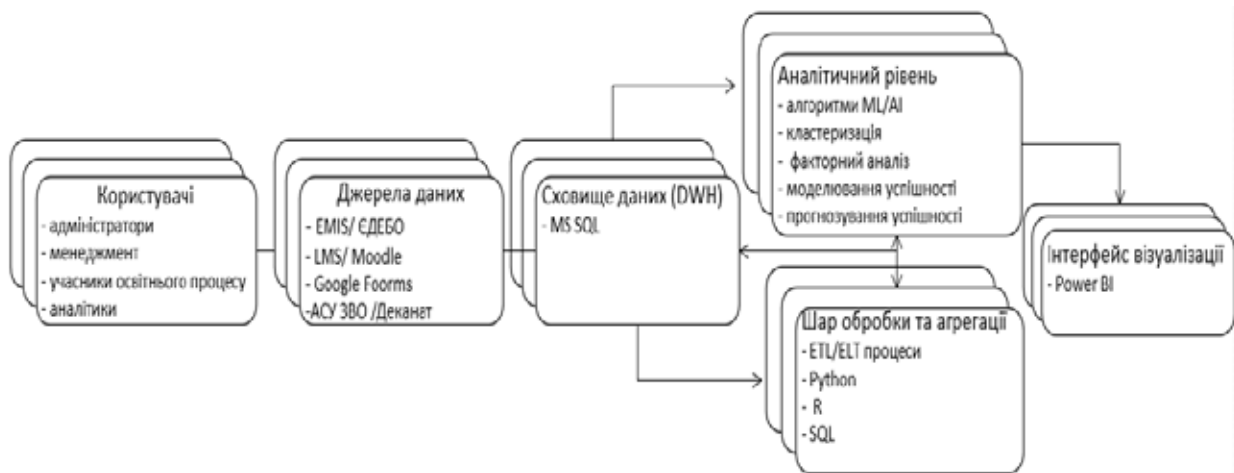


Рис. 2. Прототип інформаційної технології освітньої аналітики

Тут, оскільки нашою відправною точкою є великі дані, зібрані з LMS, ЄДЕБО та інформаційних сервісів та систем НУБІП України, використання інструментів аналітики є неминучим для досягнення нашої кінцевої мети – виявити недоліки в процесах навчання та спробувати надавати звіти та прогнози для здобувачів освіти.

Центральним елементом системи є власне аналітичний рівень, на якому відбувається вся робота з даними. Аналіз даних – це процес дослідження, очищення, перетворення та моделювання великих наборів даних з метою виявлення прихованих закономірностей, невідомих кореляцій і збору корисної інформації.

Інформаційна технологія аналітики навчання, динамічний інструмент у сфері освіти, швидко розвивається і стає незамінним етапом освітніх стратегій у сучасних закладах освіти. Ця концепція, яка заглиблюється в аналіз і застосування даних у навчанні, є трансформаційним підходом до покращення результатів навчання. Прогнозування академічної успішності студентів є важливим завданням для будь-якої установи для прийняття обґрунтованих рішень щодо прогресу студентів у навчанні. Алгоритми інтелектуального аналізу даних можуть бути використані для прогнозування успішності студентів та визначення факторів, що впливають на точність прогнозу. На (рис. 3) показано методологію реалізації запропонованої роботи.

У межах дослідження, спрямованого на аналіз освітньої діяльності та прогнозування академічної успішності студентів, було використано емпіричні дані, отримані з системи управління навчанням (Learning Management System, LMS) факультету інформаційних технологій Національного університету біоресурсів і природокористування України. Вибірка охоплює студентів з першого по четвертий курс, які навчаються за різними освітніми програмами факультету, і налічує 74 427 записів. Зібраний масив даних включає широкий спектр змінних, зокрема: ідентифікатори студентів та навчальних груп, прізвище, ім'я та по батькові, код і назву дисципліни, дату та номер заняття, інформацію про відвідування, нормалізований ідентифікатор, стать, джерело фінансування навчання (бюджет або контракт), курс навчання, конкурсний бал при вступі, кількість входів до LMS-курсу, а також поточні академічні результати.

Структура та обсяг даних забезпечують можливість проведення як описового статистичного аналізу освітнього процесу, так і побудови моделей прогнозування академічної успішності із застосуванням сучасних методів машинного навчання. У ході дослідження було сформовано першу аналітичну модель, метою якої є виявлення залежності між рівнем активності студентів у системі управління навчанням (визначеним за кількістю входів на курс) та отриманими ними підсумковими балами з дисципліни.

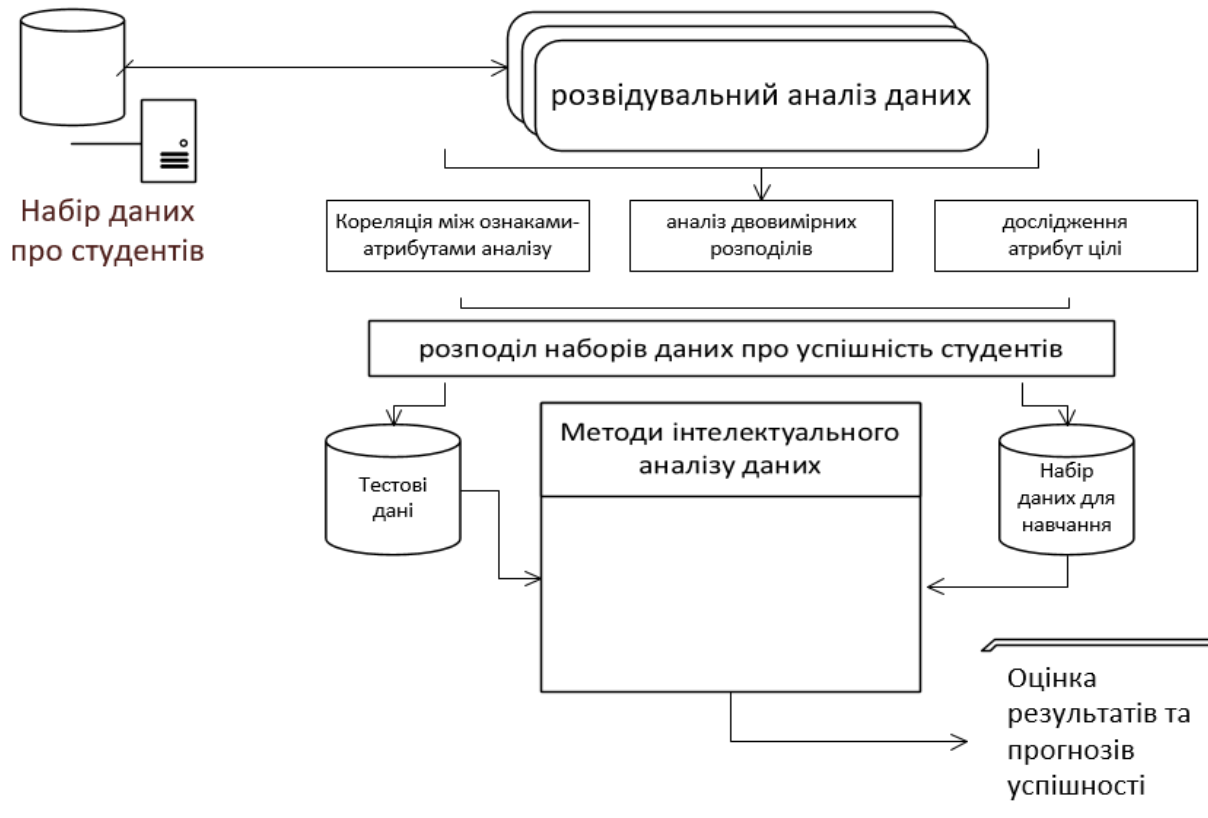


Рис. 3. Блок-схема методології побудови прогностичного моделювання для прогнозування академічної успішності студентів

З метою підвищення точності оцінювання моделі з аналізу були виключені записи студентів із нульовими балами, оскільки такі результати, як правило, є наслідком неявки на підсумковий контроль або нездачі іспиту/заліку. Для візуалізації виявленої залежності було побудовано діаграму розсіювання, де по осі абсцис (X) відображено кількість входів студента на курс у LMS, а по осі ординат (Y) – його підсумковий бал з відповідної дисципліни.

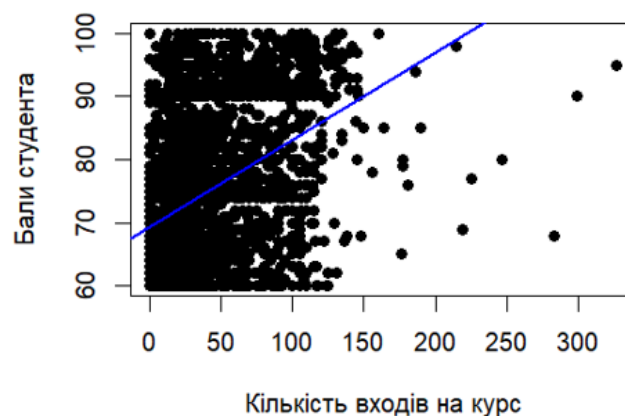


Рис. 4. Характер впливу навчальної активності (кількості входів на відповідний навчальний курс в LMS) на успішність здобувачів освіти

Побудований на основі відібраних даних графік демонструє наявність певної закономірності між активністю студентів у системі управління навчанням (вимірюваною кількістю входів на курс) та їхніми підсумковими балами з дисципліни. На графіку також чітко простежується лінія лінійної регресії з позитивним нахилом, що свідчить про тенденцію до зростання академічної успішності зі збільшенням активності у навчальному курсі (після виключення записів із нульовими балами).

У межах другої аналітичної моделі було досліджено залежність між відвідуваністю студентів та часовими характеристиками проведення занять – зокрема, місяцем та годиною. Такий підхід дозволив виявити часові патерни, які можуть мати практичне значення при плануванні навчального процесу.

Результати аналізу підтвердили наявність стійких закономірностей у поведінці студентів. Після виключення аномальних, основні тенденції залишилися незмінними, що свідчить про їхню стабільність (рис. 5).

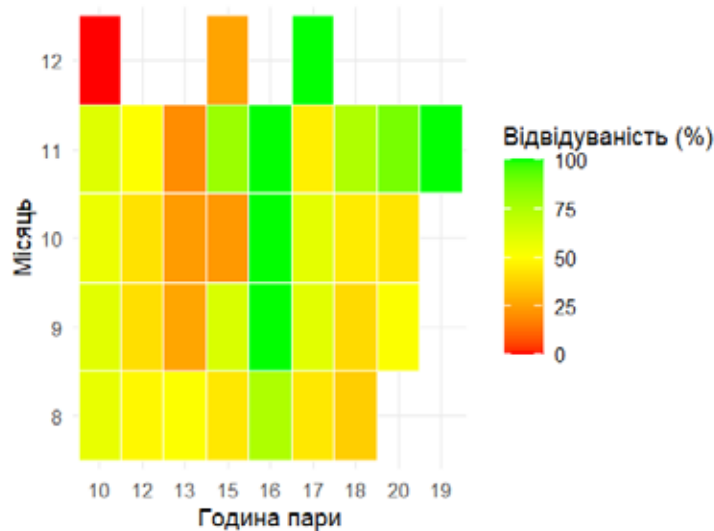


Рис. 5. Рівень відвідуваності ресурсів в LMS

З метою поглибленого аналізу поведінкових моделей студентів у контексті навчальної активності та академічної успішності було застосовано метод кластерного аналізу – алгоритм k-середніх (k-means). Для побудови моделі кластеризації було обрано два ключові параметри: кількість відвідувань занять та середній бал студента за курс. Попередньо з вибірки було виключено записи з нульовими середніми балами, оскільки вони, як правило, пов'язані з академічною заборгованістю або відсутністю оцінювання, що могло б спотворити результати кластеризації (рис. 6).

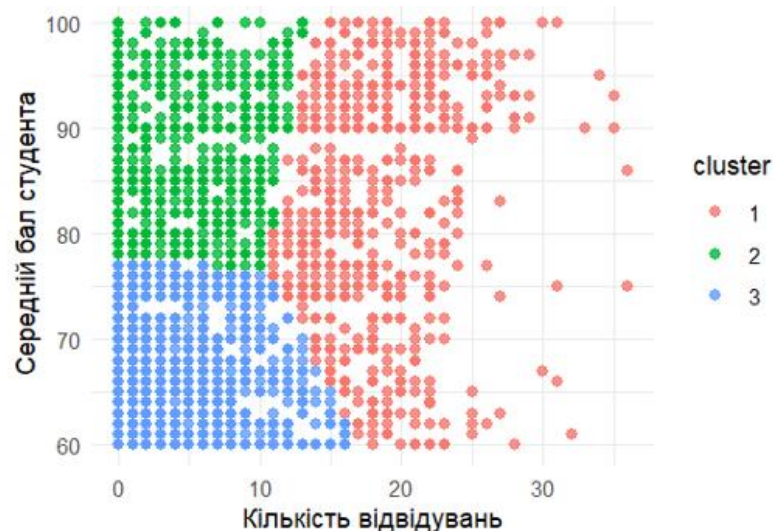


Рис. 6. Кластеризація студентів за відвідуваністю та балами успішності

Аналіз дозволив виокремити три основні кластери студентів, кожен з яких характеризується специфічними особливостями навчальної поведінки:

- Кластер 1 (червоний) – найбільш чисельна група, що включає студентів із високим рівнем відвідуваності та середніми або високими академічними результатами. Це свідчить про позитивну кореляцію між регулярною присутністю на заняттях та успішністю. Водночас у цій групі трапляються

студенти з нижчими балами, що може вказувати на вплив додаткових чинників, таких як складність дисципліни або рівень мотивації.

- Кластер 2 (зелений) – студенти з низькою відвідуваністю, але високими середніми балами. Ця категорія демонструє здатність до самостійного опанування матеріалу, ймовірно, за допомогою альтернативних освітніх ресурсів (електронні курси, відеолекції тощо). Наявність такого кластеру підкреслює необхідність диференційованого підходу до оцінювання ефективності навчання.

- Кластер 3 (синій) – студенти з низькою відвідуваністю та низькими академічними результатами. Ця група є найбільш вразливою з точки зору ризику академічної неупішності та потребує цілеспрямованої педагогічної підтримки.

Такий розподіл дозволяє глибше зрозуміти різноманіття навчальних стратегій студентів та ідентифікувати потенційні зони ризику, що є важливим для розробки індивідуалізованих освітніх підходів.

Висновки. У статті було розглянуто концептуальні та практичні аспекти побудови архітектури інформаційної технології освітньої аналітики з інтеграцією методів інтелектуального аналізу даних. Проведений аналіз літературних джерел засвідчив зростаючий інтерес наукової спільноти до застосування алгоритмів машинного навчання, кластеризації, регресійного моделювання та візуалізації даних у сфері освіти.

Запропонована аналітична модель дозволила виявити залежності між активністю студентів у системі управління навчанням та їхньою академічною успішністю, а також часові патерни відвідуваності занять. Застосування кластерного аналізу дало змогу виокремити типові поведінкові групи студентів, що відкриває перспективи для персоналізації навчального процесу та раннього виявлення ризиків академічної неупішності.

Результати дослідження підтверджують доцільність створення гнучкої, масштабованої архітектури освітньої аналітики, яка поєднує збір, обробку, аналіз і візуалізацію освітніх даних. Така система може стати ефективним інструментом підтримки прийняття рішень у закладах вищої освіти, сприяти підвищенню якості навчального процесу та забезпечити умови для впровадження адаптивного навчання.

Подальші дослідження доцільно спрямувати на розширення функціональності аналітичної системи, інтеграцію з іншими освітніми платформами, а також на вивчення етичних аспектів використання освітніх даних.

Список використаних джерел:

1. Глазунова О. Г., Клименко Є. О., Волошина Т. В., Мокрієв М. В., Вороненко О. В. Освітня аналітика в університетах: інструменти для аналізу та прогнозування. *Телекомунікаційні та інформаційні технології*. 2024. № 2(83) с. 49–59 <https://doi.org/10.31673/2412-4338.2024.026171>
2. Скрипник А., Клименко Н., Костенко І. Рівень освіченості населення в галузі цифрових технологій та зростання економік країн. *Інформаційні технології і засоби навчання*. 2020. № 4 с. 278–297 <https://doi.org/10.33407/itl.v78i4.2948>
3. Hlazunova Olena, Klymenko Nataliia, Mokriev Maksym, Nehrey Maryna, Klymenko Yevhenii. Data Analysis Technologies for Enhanced Educational Processes: A Case Study Using the Moodle LMS. *Lecture Notes on Data Engineering and Communications Technologies*. 2025. Vol. 242, Pages 670–682
4. Romero C., Ventura S. Educational data mining and learning analytics: An updated survey. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 2020. 10(3), e1355. <https://doi.org/10.1002/widm.1355>
5. Eighth IEEE International Conference on Data Mining. *IEEE Xplore*. 2008. URL: <https://ieeexplore.ieee.org/xpl/conhome/4781077/proceeding>
6. IEEE International Conference on Data Mining (ICDM). *IEEE Xplore*. 2025. (July 11, 2025), URL: <https://ieeexplore.ieee.org/xpl/conhome/1000179/all-proceedings>
7. Ifenthaler D., Yau J. Y.-K. Utilising learning analytics to support study success in higher education: A systematic review. *Educational Technology Research and Development*, 2020. 68, 1961–1990. <https://doi.org/10.1007/s11423-020-09788-z>
8. Okike E., Morogosi M. Educational Data Mining for Monitoring and Improving Academic Performance at University Levels. *International Journal of Advanced Computer Science and Applications*. 2020. № 11. DOI: 10.14569/IJACSA.2020.0111171.
9. Papamitsiou Z., Economides A. A. Learning analytics and educational data mining in practice: A systematic literature review of empirical evidence. *Educational Technology & Society*, 2014. 17(4), 49–64. URL: <https://www.jstor.org/stable/jeductechsoci.17.4.49>
10. Sinha S., Castro E., Moran C. How artificial intelligence can personalize education. *IEEE Spectrum*. 2023. URL: <https://spectrum.ieee.org/how-ai-can-personalize-education>
11. Williamson B. Introduction: Learning machines, digital data and the future of education. SAGE Publications Ltd. 2017. <https://doi.org/10.4135/9781529714920>

Дата надходження статті: 12.07.2025

Дата прийняття статті: 18.07.2025

Опубліковано: 23.09.2025

УДК 004.415.5:004.451.3:004.056

DOI <https://doi.org/10.32689/maup.it.2025.2.11>**Андрій КОБИЛЮК**

асистент кафедри обчислювальної техніки
факультету інформатики та обчислювальної техніки,
Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»,
andrii.kobyliuk@gmail.com
ORCID: 0009-0008-7784-4099

**ПІДВИЩЕННЯ ЕФЕКТИВНОСТІ УПРАВЛІННЯ ХМАРНИМИ РЕСУРСАМИ
ЗА ДОПОМОГОЮ РОЗПОДІЛЕНОГО МОНІТОРИНГУ**

Анотація. Метою даної роботи є розробка архітектури розподіленої системи моніторингу, здатної до адаптивного реагування на змінне навантаження, часткові відмови та мережеві деградації задля підвищення ефективності управління хмарними обчислювальними ресурсами. Методологія дослідження ґрунтується на розробці ієрархічної архітектури, яка включає моніторингові агенти з адаптивною частотою вибірки, *fog*-рівень із локальними ML-моделями для прогнозування навантаження, а також центральний координатор, який виконує глобальну аналітику та стратегічне управління. У ході дослідження було реалізовано експериментальний прототип системи на базі контейнеризованого середовища з використанням Docker, Apache Kafka, Python, Random Forest та SHAP-інтерпретації. Було проведено серію експериментів за п'ятьма сценаріями: від номінальної роботи до стресових умов із втратами зв'язку, перевантаженнями та комбінованими загрозами.

Наукова новизна – вперше реалізовано поєднання агентного моніторингу з динамічною частотою вибірки та ML-аналітики на *fog*-рівні для забезпечення високої швидкості й точності реагування в умовах децентралізованого управління. Запропонована модель перевершує традиційні централізовані підходи як за якісними, так і за кількісними показниками.

Висновки. За результатами експериментів встановлено, що запропонована архітектура дозволяє знизити середній час реакції системи на події на 60–70 %, зменшити втрати телеметричних даних на 3–4 %, підвищити точність прогнозування навантаження до 30 % і зменшити кількість помилкових дій в умовах динамічної інфраструктури. Крім того, розподілений характер обробки інформації забезпечує вищий рівень стабільності системи за рахунок зменшення флуктуацій в алокації ресурсів. Отримані результати підтверджують доцільність впровадження запропонованої архітектури у системах, що потребують високої надійності, адаптивності та масштабованості – зокрема, в *edge/fog*-середовищах, промисловому Інтернеті речей, розумних дата-центрах та критично важливих сервісах реального часу.

Ключові слова: хмарні технології, моніторинг, управління, паралельні обчислення, безпека хмарних обчислень.

Andrii KOBILYUK. INCREASING THE EFFICIENCY OF CLOUD RESOURCE MANAGEMENT WITH THE HELP OF DISTRIBUTED MONITORING

Abstract. The purpose of this work is to develop an architecture for a distributed monitoring system capable of adaptively responding to variable load, partial failures, and network degradation to improve the efficiency of cloud computing resource management.

The research methodology is based on the development of a hierarchical architecture that includes monitoring agents with adaptive sampling rates, a fog layer with local ML models for load forecasting, and a central coordinator that performs global analytics and strategic management. During the study, an experimental prototype of the system was implemented based on a containerized environment using Docker, Apache Kafka, Python, Random Forest, and SHAP interpretation. A series of experiments were conducted under five scenarios: from nominal operation to stress conditions with loss of connectivity, overloads, and combined threats.

Scientific novelty – for the first time, a combination of agent monitoring with a dynamic sampling rate and ML analytics at the fog level has been implemented to ensure high speed and accuracy of response in decentralized management conditions. The proposed model outperforms traditional centralized approaches in both qualitative and quantitative indicators.

Conclusions. According to the results of experiments, it was found that the proposed architecture allows reducing the average system response time to events by 60–70 %, reducing telemetry data losses by 3–4 %, increasing the accuracy of load forecasting by up to 30 %, and reducing the number of erroneous actions in dynamic infrastructure conditions. In addition, the distributed nature of information processing provides a higher level of system stability by reducing fluctuations in resource allocation. The results obtained confirm the feasibility of implementing the proposed architecture in systems that require high reliability, adaptability and scalability – in particular, in edge/fog environments, industrial Internet of Things, smart data centers and critical real-time services.

Key words: cloud technologies, monitoring, management, parallel computing, cloud computing security.

Постановка проблеми. Із розвитком хмарних технологій, зростанням обсягів динамічно змінюваних навантажень та переходом до мікросервісних і edge/fog-орієнтованих архітектур питання ефективного управління обчислювальними ресурсами стає важливим як у науковому, так і в практичному контексті. Сучасні хмарні середовища характеризуються багаторівневою структурою, високою інтенсивністю телеметричних потоків, необхідністю забезпечення безперервності сервісів та дотриманням угод про якість обслуговування.

Традиційні централізовані моделі моніторингу та управління [1] не забезпечують необхідного рівня адаптивності, масштабованості та відмовостійкості. В умовах непередбачуваних змін навантаження або часткових відмов інфраструктури такі системи демонструють високу затримку реакції, втрату даних та нестабільну роботу керувальних механізмів. Особливо це критично для сервісів реального часу, IoT-платформ, розподілених обчислень та сценаріїв з обмеженою пропускну здатністю мережі.

У цьому контексті особливої актуальності набувають дослідження, спрямовані на створення розподілених інтелектуальних систем моніторингу, здатних локально аналізувати ситуацію, прогнозувати зміну навантаження та приймати автономні рішення. Такі підходи сприяють не лише зниженню часу реакції та навантаження на мережу, а й підвищенню надійності всієї хмарної інфраструктури.

Таким чином, проблема підвищення ефективності управління хмарними ресурсами через впровадження розподіленого, адаптивного моніторингу із використанням засобів штучного інтелекту має безпосередній зв'язок із вирішенням низки прикладних задач сучасної інформатики, кіберфізичних систем та інженерії програмного забезпечення. Її розв'язання є важливим як для подальшого розвитку наукової теорії управління у децентралізованих системах, так і для практичного впровадження у промислових, телекомунікаційних та науково-обчислювальних середовищах.

Аналіз останніх досліджень і публікацій. Управління ресурсами в хмарних середовищах продовжує залишатися актуальним викликом, особливо з огляду на зростаючу динамічність навантаження, багаторівневу архітектуру систем та високі вимоги до доступності. Сучасні дослідження демонструють стійку тенденцію переходу від централізованих моделей моніторингу до гібридних або повністю розподілених рішень, інтегрованих із методами штучного інтелекту.

В [2] автори пропонують загальну архітектуру для інтеграції методів машинного навчання (ML) у системи керування. Водночас робота не охоплює критичні аспекти, пов'язані із затримками передачі даних та відмовостійкістю при масштабуванні. В [3] здійснено класифікацію підходів до управління ресурсами, виділяючи реактивні, проактивні та гібридні моделі. Проте дослідження має обмежену практичну орієнтацію щодо реалізації розподіленого моніторингу в реальних умовах.

Дослідження [4] сфокусоване на проактивних методах у контексті сервіс-орієнтованих хмарних архітектур (cloud of services). Дослідження містить поглиблений аналіз моделей прогнозування та контекстної адаптації із чітким фокусом на забезпечення QoS та дотримання SLA. Водночас недостатньо розкрито питання відмовостійкості у розподіленому контролі. Автори в [5] аналізують децентралізовані стратегії управління на стику хмари та edge-обчислень, акцентуючи увагу на latency-sensitive додатках. Хоча підхід є перспективним, обсяг валідації залишається обмеженим.

В [6] пропонують AI-фреймворк для алокації ресурсів у гібридних хмарах із мікросервісною архітектурою. Модель відзначається практичною спрямованістю на гібридні сценарії. Проте у дослідженні майже відсутній опис моніторингової інфраструктури, а застосовані AI-моделі не мають прозорого обґрунтування прийнятих рішень.

В [7] представляють систему HUNTER – комплексне AI-рішення для управління ресурсами з акцентом на енергоефективність. Серед переваг – повноцінна реалізація та орієнтація на принципи сталого розвитку. Однак система характеризується високою складністю, що ускладнює її адаптацію в реальних умовах, а моніторинг базується на централізованих вузлах. В роботі [8] досліджують автоматизоване виявлення інцидентів та self-healing механізми в хмарних середовищах. Позитивним є інтеграція з логами, подіями та аудитами, що створює базу для реального застосування AI в оперативному режимі. Водночас запропоноване рішення є досить вузькоспеціалізованим.

Серед українських досліджень, у [1] пропонується ML-підхід до управління хмарними ресурсами, тоді як в [5] зосереджуються на практичному підвищенні ефективності ресурсної алокації. Їхні роботи мають цінність завдяки адаптації до локальних умов і можуть слугувати прикладами впровадження AI у хмарні середовища в освітньому та науковому контексті України. Разом із тим, обмеження полягають у недостатньо деталізованих експериментах та обмеженому зіставленні з міжнародними практиками.

Незважаючи на значний прогрес у застосуванні методів штучного інтелекту до управління хмарними ресурсами, проблема забезпечення розподіленого та відмовостійкого моніторингу в режимі реального часу залишається недостатньо вивченою. Наявні підходи переважно ґрунтуються на централізованих моделях і часто ігнорують виклики, пов'язані із затримками, надійністю та динамічною

алокцією ресурсів у масштабованих середовищах. Це зумовлює потребу у створенні ефективної архітектури, яка поєднує розподілений моніторинг із інтелектуальними механізмами управління для сучасних хмарних систем.

Формулювання мети статті (постановка завдання). Метою цієї роботи є розробка та експериментальне дослідження розподіленої системи моніторингу, здатної адаптивно реагувати на динамічні зміни навантаження, часткові відмови та мережеву деградацію, забезпечуючи більш надійне, гнучке та масштабоване керування інфраструктурою.

Для досягнення цієї мети поставлено такі завдання:

1. Розробити концепцію архітектури розподіленої системи моніторингу, яка забезпечує підвищену відмовостійкість, зниження затримок обробки та гнучкість масштабування в хмарному середовищі.
2. Інтегрувати модулі інтелектуального аналізу та прогнозування навантаження з використанням машинного навчання для прийняття рішень в умовах динамічного середовища.
3. Реалізувати прототип системи у віртуалізованому середовищі з використанням інструментів контейнеризації, брокерів повідомлень та агентів збору телеметрії.
4. Провести серію експериментів для оцінки реакції системи на стресові сценарії (перевантаження, втрата зв'язку, затримки).

Виклад основного матеріалу. *Опис запропонованої архітектури системи моніторингу.* Запропонована архітектура передбачає перехід від централізованого збору телеметрії до ієрархічної розподіленої системи моніторингу, в якій агенти, розгорнуті на вузлах хмарної інфраструктури, самостійно агрегують, обробляють і частково інтерпретують дані до їх передачі до центрального або проміжного аналітичного рівня. Схема архітектури зображена на (рис. 1).

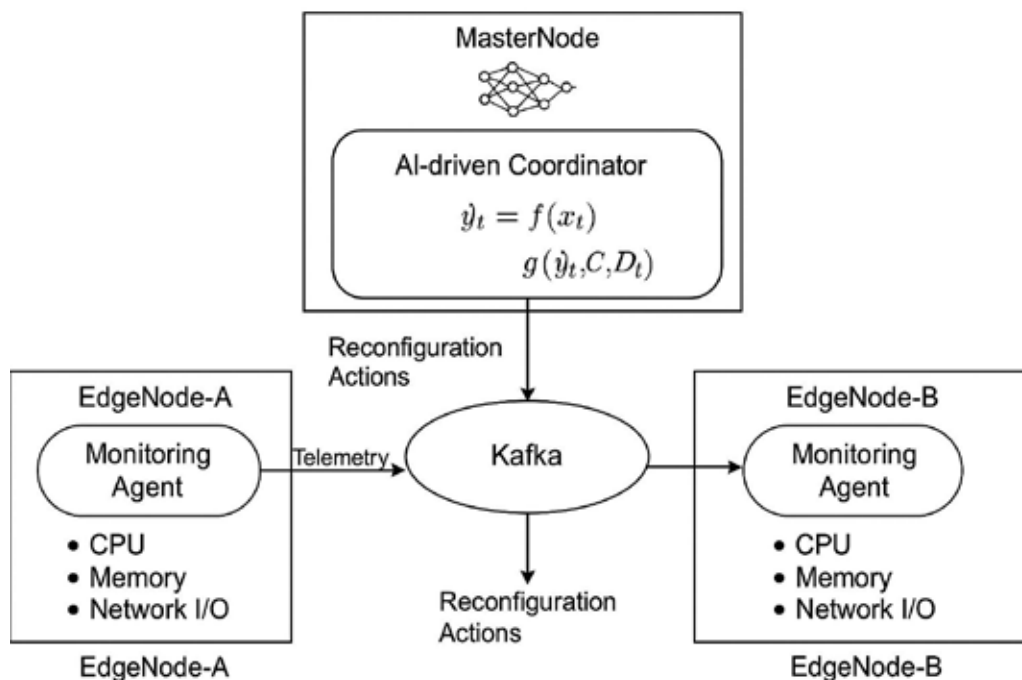


Рис. 1. Схема архітектури розподіленої системи моніторингу

Джерело: складено автором

Метою такої архітектури є зменшення затримок, підвищення відмовостійкості та забезпечення адаптивного управління ресурсами в умовах змінного навантаження.

Компонентами запропонованої архітектури є:

1. Monitoring Agents виконують локальну обробку телеметрії: усереднення, виявлення аномалій, вираховування похідних метрик. Підтримують адаптивну частоту вибірки $f_i(t)$, яка залежить від зміни навантаження.
2. Fog Layer / Intermediate Nodes агрегують дані від агентів у рамках підмережі. Використовують локальні ML-моделі. У разі втрати зв'язку з центральним хабом тимчасово виконують самостійне управління.

3. Central Coordinator приймає агреговані дані, проводить глобальний аналіз. Приймає стратегічні рішення про алокацію ресурсів. Визначає порогові значення для агентів та fog-вузлів.

Хід експерименту. Для перевірки гіпотези про ефективність запропонованої архітектури розподіленої системи моніторингу та управління ресурсами в умовах нестабільного навантаження і часткових відмов, було проведено серію емпіричних експериментів у симульованому середовищі. Основною метою експерименту було оцінити здатність системи адаптивно реагувати на зміни в інфраструктурі без втрати стабільності, продуктивності та цілісності даних.

Експериментальне середовище включало три логічні вузли: MasterNode, EdgeNode-A та EdgeNode-B, які були розгорнуті у формі віртуалізованих контейнерів на фізичній машині з використанням інструментарію Docker та Docker Compose. На кожному з вузлів було встановлено моніторингові агенти, які відповідали за збір телеметричних даних про використання процесора, оперативної пам'яті, мережевої активності та доступності сервісів. Дані передавались до центрального брокера повідомлень (Apache Kafka) або записувалися локально у чергу на випадок тимчасової втрати з'єднання.

Керувальна компонента, розміщена на MasterNode, реалізовувала модуль обробки даних, AI-модуль для прийняття рішень та блок оркестрації, що взаємодіє з агентами через REST API. AI-модуль базувався на поєднанні моделей випадкового лісу (Random Forest) для прогнозування короточасного навантаження та аналізу важливості факторів (через SHAP-інтерпретацію).

Для емуляції зовнішніх факторів, таких як затримки мережі, втрати з'єднання або часткові відмови компонентів, використовувалась утиліта tc з модулем netem на рівні мережевого інтерфейсу.

Було визначено п'ять окремих експериментальних сценаріїв, кожен з яких моделював конкретну ситуацію в роботі системи.

Експеримент 1 Номінальний режим роботи. У даному сценарії система працювала без зовнішніх збурень. На обох Edge-вузлах було згенеровано помірне навантаження (до 40 % використання CPU, до 50 % пам'яті). Збір телеметрії відбувався з інтервалом у 5 секунд, керувальний модуль здійснював періодичний аналіз отриманих даних і приймав рішення про підтримання поточної конфігурації. Даний експеримент слугував еталоном для оцінки відхилень у наступних сценаріях.

Експеримент 2 Раптове зростання навантаження. На EdgeNode-A протягом 60 секунд було поступово збільшено навантаження на CPU до 90 % за допомогою утиліти stress-ng. Система мала визначити аномальне зростання навантаження, спрогнозувати подальше перевантаження та здійснити переведення частини контейнеризованих сервісів на EdgeNode-B.

Експеримент 3 Відмова вузла агента. EdgeNode-B було вимкнено на рівні симульованої відмови мережевого інтерфейсу. В результаті, моніторинговий агент ставав недоступним, а дані з нього переставали надходити. Очікувалося, що керувальний модуль виявить втрату джерела телеметрії протягом 10 секунд, позначить вузол як непрацездатний, ініціює відновлення з резервної черги даних і переведе критичні сервіси на активні вузли.

Експеримент 4 Мережеві затримки. Було змодельовано затримку 250 мс між агентами та керувальним вузлом через tc qdisc. Така ситуація імітує мережеву деградацію або віддалене розташування вузлів. Метою експерименту було перевірити стійкість протоколу передачі даних, здатність AI-модуля до обробки частково запізненої інформації та оцінка впливу на прийняття рішень у режимі near-real-time.

Експеримент 5 Комбінований (стресовий) сценарій. У цьому сценарії було одночасно реалізовано кілька стрес-факторів: зростання навантаження на EdgeNode-A до 85 %, втрата зв'язку з EdgeNode-B та затримка мережі між EdgeNode-A та MasterNode у 200 мс. Експеримент мав на меті оцінити здатність системи одночасно обробляти декілька загрозливих умов, виявляти критичні вузли, забезпечувати безпечну переалюкацію навантаження та відновлення зв'язку.

У ході кожного експерименту збирались такі метрики: час реакції системи, кількість втрачених повідомлень у мережі, точність прогнозу навантаження на 30 секунд уперед (MAE, RMSE), кількість некоректних рішень, індекс стабільності роботи (кількість змін розподілу за експеримент).

Результати. Значення метрик для кожного експерименту випробування запропонованої розподіленої архітектури в порівнянні із результатами еталонної централізованої архітектури наведені в (табл. 1), де t – час реакції (с), $loss$ – втрата повідомлень (%), $Error$ – некоректні рішення, I – індекс стабільності.

У всіх сценаріях запропонована архітектура продемонструвала значно менший середній час реакції (від 1,2 с у номінальному режимі до 4,5 с у комбінованому сценарії), тоді як централізована система показала затримки в діапазоні від 2,8 до 12,8 с. Це вказує на ефективність локальної обробки та зменшення навантаження на центральний вузол завдяки fog-рівню. Найбільша різниця спостерігалася у сценаріях із відмовами або стресовими умовами, що підтверджує підвищену відмовостійкість ієрархічної моделі.

Таблиця 1

Результати експериментів

№	Архітектура	t (с)	loss (%)	MAE	RMSE	Error	I
1.	Централізована	2,8	1,50 %	3,2 %	4,1 %	0	0,95
	Розподілена	1,2	0,40 %	2,5 %	3,3 %	0	0,98
2.	Централізована	7,6	3,20 %	8,9 %	11,0 %	2	0,65
	Розподілена	2,4	1,00 %	4,1 %	5,80 %	0	0,91
3.	Централізована	11,4	4,50 %	9,8 %	12,7 %	3	0,52
	Розподілена	3,6	1,20 %	5,2 %	6,9 %	1	0,85
4.	Централізована	5,1	2,70 %	7,2 %	8,5 %	1	0,74
	Розподілена	2,8	1,10 %	4,5 %	6,1 %	0	0,89
5.	Централізована	12,8	6,20 %	10,4 %	13,3 %	4	0,41
	Розподілена	4,5	1,90 %	5,7 %	7,40 %	1	0,81

Джерело: складено автором

Запропонована архітектура забезпечує зменшення відсотку втрат повідомлень на 1,5–4,3 % залежно від сценарію. Основними чинниками покращення стали використання локальних буферів, можливість тимчасового автономного функціонування fog-вузлів і оптимізована стратегія ретрансляції. У централізованій системі втрати значно збільшувалися при мережевих деградаціях та відмовах вузлів.

Метрики MAE і RMSE демонструють суттєву перевагу розподіленої архітектури, особливо у сценаріях з динамічними змінами навантаження (Експерименти 2, 5). Це пояснюється застосуванням локальних ML-моделей у fog-рівні, які забезпечують більш актуальні та контекстно-чутливі прогнози. Централізована система при обробці запізнених або неповних даних має гірші показники точності (наприклад, RMSE до 13,3 % у стресовому сценарії проти 7,4 % у розподіленій).

Кількість некоректних керувальних дій (наприклад, передчасна або запізнена переалокція сервісів) у централізованій архітектурі була значно вищою, що корелює з меншою точністю прогнозів та більшим часом реакції. Розподілена система демонструє нульові або одиничні помилки навіть в умовах комбінованих відмов, що свідчить про її адаптивність.

Індекс стабільності (рис. 2), який характеризує частоту змін розподілу сервісів у відповідь на флуктуації системи, був вищим у розподіленій архітектурі в усіх сценаріях. Це означає, що система здатна утримувати ефективний стан з меншою кількістю втручань, що позитивно впливає на продуктивність та енергоспоживання.

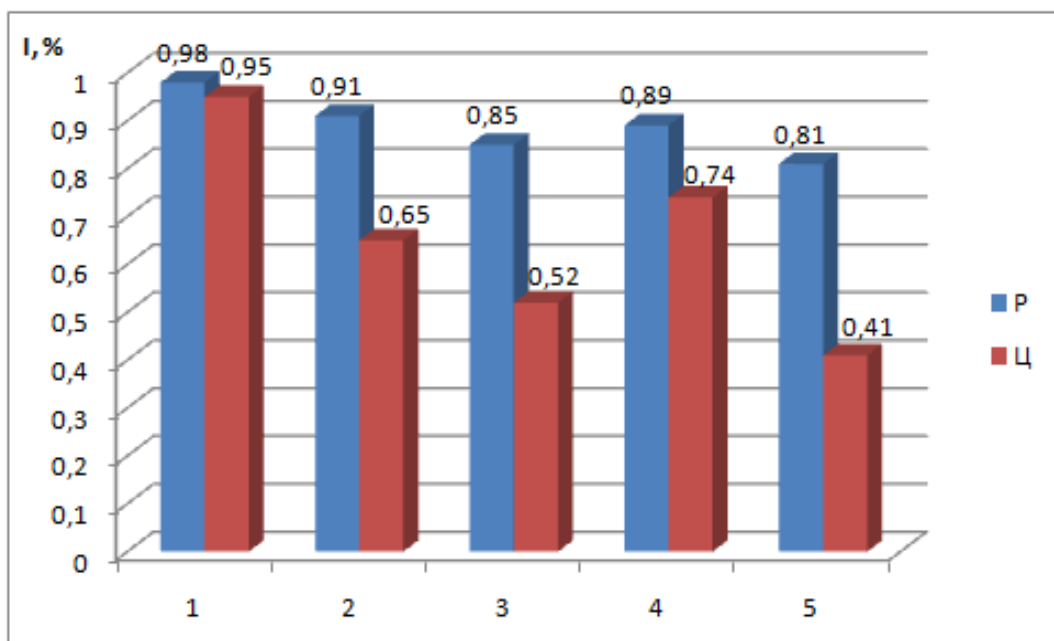


Рис. 2. Гістограми Індексу стабільності

Джерело: складено автором

Обговорення результатів. Запропонована архітектура продемонструвала стабільне покращення основних метрик управління хмарними ресурсами порівняно з класичною централізованою моделлю. Практична цінність результатів полягає у зменшенні середнього часу реакції системи на критичні події до 60–70 %, зниженні втрат телеметрії, покращенні точності прогнозування навантаження та підвищенні стабільності системи в умовах часткових відмов і мережевої деградації. Архітектура є придатною для використання в реальних сценаріях, де важливими є гнучкість, надійність та адаптивність розподіленого управління.

Крім того, використання ієрархічної структури з fog-рівнем дозволяє забезпечити паралельну обробку даних у незалежних доменах моніторингу, що відкриває можливості масштабування системи без втрати швидкодії. Це особливо актуально для систем з великою кількістю вузлів.

Ще одним важливим практичним аспектом є потенціал інтеграції механізмів безпеки на рівні агентів і fog-вузлів. Розподілений характер системи дозволяє локалізувати критичні події, знижувати ризики централізованих атак та забезпечувати захист телеметричних даних у процесі передачі. Це створює основу для впровадження стійких до загроз архітектур у сфері безпеки хмарних обчислень.

Наукова новизна дослідження полягає у поєднанні локальних агентів з адаптивною частотою моніторингу, fog-рівня з ML-прогнозуванням та центрального координатора з explainable AI. На відміну від більшості робіт [2, 4, 7], де моніторинг реалізовано централізовано або фрагментарно, запропонована система забезпечує повноцінне багаторівневе реагування на події з мінімальною затримкою та автономною роботою в умовах втрати зв'язку. Порівняно з підходами [5–6], які акцентують увагу на розподіленому управлінні, але не розглядають детально безпечність чи масштабування моніторингу, запропонована система пропонує гнучке рішення з високим рівнем автономії та потенціалом для захисту даних і паралельної обробки.

Таким чином, результати дослідження не лише підтверджують переваги розподіленого моніторингу, але й демонструють його практичну реалізованість у складних, безпечних та масштабованих хмарних середовищах.

Висновки. У межах даного дослідження було розроблено, реалізовано та експериментально перевірено концепцію ієрархічної розподіленої системи моніторингу для управління хмарними ресурсами в умовах динамічного навантаження та часткових відмов. Запропонована архітектура поєднує локальні агенти з адаптивною частотою вибірки, проміжний fog-рівень із машинним навчанням та центральний координатор з аналітичним модулем, що забезпечує стратегічне управління.

Результати серії експериментів засвідчили суттєве покращення ключових показників ефективності порівняно з централізованими моделями: час реакції зменшено до 60–70 % у критичних сценаріях; втрати повідомлень знижено до 1,9 % навіть у стресових умовах; точність прогнозування навантаження (MAE, RMSE) покращено в середньому на 30 %; знижено кількість некоректних керувальних дій; підвищено стабільність системи внаслідок зменшення надмірних перерозподілів ресурсів.

У майбутньому передбачається розширення системи шляхом включення SLA-моделі та оцінки впливу управлінських рішень на рівень обслуговування кінцевого користувача; тестування на гетерогенному апаратному забезпеченні та в edge/fog-мережах із реальними IoT-вузлами; впровадження енергозберігаючих механізмів і оптимізації обчислювальних витрат на ML-моделі; дослідження самонавчальних агентів, здатних адаптувати алгоритми моніторингу без централізованої переналаштування.

Запропонована система є ефективним підґрунтям для створення адаптивної, масштабованої та стійкої до відмов платформи керування хмарною інфраструктурою в умовах високої мінливості навантаження та обмежених ресурсів.

Список використаних джерел:

1. Улізько В. О. Інтелектуальна система управління ресурсами в хмарних середовищах на основі машинного навчання. 2024. URL: <https://ela.kpi.ua/handle/123456789/72021> (дата звернення: 10.06.2025).
2. Almutairi S., Alghanmi N., Monowar M. M. Survey of centralized and decentralized access control models in cloud computing. *International Journal of Advanced Computer Science and Applications*. 2021. Vol. 12, No. 2.
3. Barua B., Kaiser M. S. AI-driven resource allocation framework for microservices in hybrid cloud platforms. *arXiv preprint*. 2024. arXiv:2412.02610. <https://doi.org/10.48550/arXiv.2412.02610>.
4. Ilager S., Muralidhar R., Buyya R. Artificial intelligence (AI)-centric management of resources in modern distributed computing systems. *2020 IEEE Cloud Summit*. 2020. P. 1–10. DOI: 10.1109/IEEECloudSummit48914.2020.00007.
5. Kotliarskyi A., Petrashenko A. Spobib pidvyshchennia efektyvnosti vykorystannia khmarnykh resursiv – A method for improving the efficiency of cloud resource usage. *Computer-integrated technologies: education, science, production*. 2024. No. 54. P. 125–129. <https://doi.org/10.36910/6775-2524-0560-2024-54-14>.
6. Li L., Bell J., Coppola M., Lomonaco V. Adaptive AI-based decentralized resource management in the cloud-edge continuum. *arXiv preprint*. 2025. arXiv:2501.15802. <https://doi.org/10.48550/arXiv.2501.15802>.

7. Marques G., Senna C., Sargento S., Carvalho L., Pereira L., Matos R. Proactive resource management for cloud of services environments. *Future Generation Computer Systems*. 2024. Vol. 150. P. 90–102. <https://doi.org/10.1016/j.future.2023.08.005>.
8. Moghaddam S.K., Buyya R., Ramamohanarao K. Performance-aware management of cloud resources: a taxonomy and future directions. *ACM Computing Surveys (CSUR)*. 2019. Vol. 52, No. 4. P. 1–37. <https://doi.org/10.1145/3337956>.
9. Ravichandran N., Inaganti A. C., Muppalaneni R., Nersu S.R.K. AI-driven self-healing IT systems: automating incident detection and resolution in cloud environments. *Artificial Intelligence and Machine Learning Review*. 2020. Vol. 1, No. 4. P. 1–11. <https://doi.org/10.69987/>.
10. Tuli S., Gill S. S., Xu M., Garraghan P., Bahsoon R., Dustdar S., Jennings N. R. HUNTER: AI-based holistic resource management for sustainable cloud computing. *Journal of Systems and Software*. 2022. Vol. 184. P. 111124. <https://doi.org/10.1016/j.jss.2021.111124>.

Дата надходження статті: 13.06.2025

Дата прийняття статті: 26.06.2025

Опубліковано: 23.09.2025

UDC 004.42:004.85

DOI <https://doi.org/10.32689/maup.it.2025.2.12>**Artem KOLOMYTSEV**

Postgraduate Student at the Department of Software Engineering,
Software Engineer, HAPPENING TECHNOLOGY LTD,
kolomytsev1996@gmail.com
ORCID: 0000-0002-3231-7593

Yuliia KUZNETSOVA

Candidate of Technical Sciences,
Associate Professor at the Department of Software Engineering,
National Aerospace University "Kharkiv Aviation Institute",
y.kuznetsova@khai.edu
ORCID: 0000-0001-9416-6981

THE ROLE OF MLOPS IN ADVANCING GREEN ENERGY: AN EVALUATION OF TECHNOLOGIES AND PRACTICES

Abstract. Machine learning-driven decision making is essential for the efficient operation of cloud-hosted virtual power plants (VPPs) aggregating hundreds to thousands of distributed energy resources (DERs). However, manually deploying and maintaining ML models at scale introduces delays, inconsistency and high operational overhead. In this paper, we survey six widely adopted MLOps frameworks—Kubeflow, Apache Airflow, MLflow, Azure ML, AWS SageMaker and Google Vertex AI—against four criteria critical to VPP environments: industry adoption, feature-set completeness, interoperability with major ML frameworks and cloud platforms, and licensing or cost constraints. Drawing on public documentation, repository activity, case studies and market research, we identify trade-offs between open-source flexibility and managed-service convenience. Our analysis shows that Apache Airflow offers the most mature and extensible pipeline orchestration for on-premise and multi-cloud VPP deployments, while Kubeflow excels in Kubernetes-native contexts. Managed services like SageMaker and Azure ML deliver faster time-to-value for teams lacking dedicated infrastructure expertise but incur higher costs and vendor lock-in. Finally, we provide domain-tailored recommendations for integrating continuous training, evaluation and monitoring into VPP forecasting workflows, demonstrating how MLOps adoption can improve prediction latency and grid responsiveness.

The goal of this article is to evaluate and compare leading MLOps frameworks—open-source and managed cloud services—against key criteria (adoption, feature completeness, interoperability, and cost) and to recommend the most suitable solutions for cloud-hosted virtual power plants.

Methodology. We selected six MLOps frameworks based on adoption, features, interoperability and cost; extracted data from official docs, repositories and market reports; scored each tool against our criteria; and distilled domain-specific recommendations for cloud-hosted VPPs.

Scientific Novelty. This article explores the under-researched intersection of MLOps and virtual power plants (VPPs), addressing the specific challenges of applying automated ML workflows to large-scale, cloud-hosted VPP systems. It provides the first domain-specific comparison of MLOps tools tailored to the operational and forecasting needs of VPPs.

Conclusion. MLOps can significantly enhance the performance and scalability of virtual power plants. This study identifies the most suitable tools for VPP use cases, highlighting Apache Airflow and Kubeflow as strong open-source options, while managed services may suit teams with limited infrastructure expertise.

Key words: Machine learning, MLOps, Virtual power plant, Distributed energy, Cloud technology, Forecasting.

Артем КОЛОМИЦЕВ, Юлія КУЗНЕЦОВА. РОЛЬ АВТОМАТИЗАЦІЇ РОЗГОРТАННЯ ІНФРАСТРУКТУРИ У ПРОСУВАННІ ЗЕЛЕНОЇ ЕНЕРГЕТИКИ: ОЦІНКА ТЕХНОЛОГІЙ ТА ПРАКТИК

Анотація. Прийняття рішень на основі машинного навчання має важливе значення для ефективної роботи хмарних віртуальних електростанцій (ВЕС), що об'єднують сотні й тисячі розподілених енергоресурсів (ПЕР). Однак ручне розгортання та підтримка моделей машинного навчання в масштабі призводить до затримок, неузгодженості та високих операційних витрат. У цій статті ми проаналізували шість широко розповсюджених фреймворків автоматизації розгортання інфраструктури машинного навчання (MLOps) – Kubeflow, Apache Airflow, MLflow, Azure ML, AWS SageMaker і Google Vertex AI – за чотирма критеріями, критично важливими для середовищ ВЕС: впровадження в галузі, повнота набору функцій, сумісність з основними фреймворками машинного навчання і хмарними платформами, а також ліцензійні або фінансові обмеження. Спираючись на публічну документацію, активність репозиторіїв, тематичні дослідження та дослідження ринку, ми визначили компроміси між гнучкістю відкритого коду та зручністю керованого сервісу. Наш аналіз показує, що Apache Airflow пропонує найбільш зрілу та розширювану оркестровку конвеєрів для локальних та мультихмарних розгортань ВЕС, в той час як Kubeflow чудово працює в контекстах на базі Kubernetes. Керовані сервіси, такі як SageMaker та Azure ML, забезпечують швидшу окупність

© A. Kolomytsev, Yu. Kuznetsova, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

інвестицій для команд, яким не вистачає спеціалістів у галузі інфраструктури, але несуть більші витрати та прив'язку до певного постачальника. Нарешті, ми надаємо рекомендації щодо інтеграції безперервного навчання, оцінки та моніторингу в робочі процеси прогнозування ВЕС, демонструючи, як впровадження автоматизації розгортання інфраструктури машинного навчання може покращити затримку прогнозування та швидкість реагування мережі.

Мета статті: оцінити та порівняти провідні фреймворки автоматизації розгортання інфраструктури машинного навчання – відкриті та керовані хмарні сервіси – за ключовими критеріями (впровадження, повнота функцій, інтероперабельність та вартість) та рекомендувати найбільш підходящі рішення для віртуальних електростанцій, розміщених у хмарі.

Методологія. Ми обрали шість фреймворків автоматизації розгортання інфраструктури машинного навчання на основі впровадження, функцій, сумісності та вартості; витягли дані з офіційних документів, репозиторіїв та ринкових звітів; оцінили кожен інструмент за нашими критеріями; і сформулювали рекомендації для хмарних ВЕС для конкретних доменів.

Наукова новизна. У цій статті досліджується малодосліджена сфера перетину автоматизації розгортання інфраструктури машинного навчання і віртуальних електростанцій (ВЕС), розглядаються конкретні проблеми застосування автоматизованих робочих процесів машинного навчання до великомасштабних хмарних систем ВЕС. Це перше порівняння інструментів автоматизації розгортання інфраструктури машинного навчання, пристосованих до операційних і прогнозних потреб ВЕС, з урахуванням специфіки конкретної галузі.

Висновки. У результаті дослідження доведено, методи автоматизації розгортання інфраструктури машинного навчання можуть значно підвищити продуктивність і масштабованість віртуальних електростанцій. У цьому дослідженні визначено найбільш підходящі інструменти для використання ВЕС, зокрема Apache Airflow та Kubeflow як сильні варіанти з відкритим вихідним кодом, тоді як керовані сервіси можуть підійти командам з обмеженим досвідом роботи з інфраструктурою.

Ключові слова: машинне навчання, автоматизація розгортання інфраструктури, віртуальна електростанція, розподілена енергетика, хмарні технології, прогнозування.

Introduction. The market for renewable energy is experiencing substantial expansion and is rapidly accelerating. Virtual power plants (VPP) are one of the next logical steps in electrical grid development [8; 13] and they are one of the best ways to integrate renewable energy into the grid, especially residential renewable energy production and storage (like rooftop solar panels and rechargeable lithium-ion wall-mount batteries). A single VPP can manage multiple (hundreds-thousands) distributed energy resources (usually connected to the same grid), which allows them to have a bigger impact on the energy market and better financial gain [13], if compared to those DERs accessing the market on their own. VPP collects telemetry from and sends control commands to those DERs.

VPP decisions can be based on various optimisation strategies, but often those strategies are implemented using machine learning [16; 17]. In many companies (especially early-stage startups) ML models are usually managed and deployed by data science engineers manually, which is a slow and error-prone process [3]. The software development industry has already solved a similar problem by automating operations using DevOps methodology, and an identical approach can be applied to ML [15]. ML Operations (usually shortened to MLOps) is the automation of ML model learning and deployment processes. According to [3] the principles of MLOps include CI/CD pipeline automation, workflow automation, reproducibility, versioning, collaboration, continuous ML training and evaluation, and continuous monitoring. A system that follows MLOps principles may consist of multiple components that allow it to achieve those principles.

ML application in VPP has its caveats: ML models should be able to make predictions fast [15] because VPP must be able to continuously send commands to thousands of DERs, and conditions are always changing: wholesale electricity market prices, weather conditions and forecasts, user energy consumption patterns etc. Some configurations of VPP involve running ML models on DER controllers as IoT edge software [10], but this article does not cover such case – it has additional pitfalls and issues: unreliable network connection [19] on the DER side leads to model not having recent market and weather forecast data, IoT edge hardware capabilities are usually limited, so ML model should be simplified [10] (which decreases predictions quality).

Research question: Evaluate existing MLOps technologies and choose those best fitting to be used in a VPP.

Method.

Related work. We analyzed existing publications on the subject and came to the conclusion that there is a scarcity of publications focusing on MLOps applications in the VPP field.

MLOps. Research regarding MLOps has been going on actively for the past few years. Kreuzberger, Kühl and Hirschl [3] have a very detailed definition of MLOps and a list of relevant technologies for each MLOps component but do not discuss domain-specific recommendations and do not actually compare/select the best technologies. Other relevant publications [9; 12] also lack domain-specific issues discussions and limitations.

ML in a renewable energy sector. Implementation of ML models for renewable energy forecasting and software-controlled distributed energy resources is a prominent topic and multiple publications discuss this issue from various angles [15; 16; 17] even considering applying MLOps approach [15], but they do not specifically cover MLOps application to a VPP.

ML in edge IoT. Multiple software and hardware solutions exist to run (and even retrain) ML models on the edge [4; 10], but this area has a list of very specific issues and a different set of issues if compared to a VPP, which is usually backend solution running on public cloud [10; 19] and talking to DERs via REST APIs of their respective vendors [5; 7]. This article does not cover ML on edge IoT use-case and problems, but the authors agree that this may be one of the ways to implement a VPP.

Selection Criteria. The comparison in the article mostly focused on workflow/pipeline automation tools for ML, excluding optional aspects like Feature Store and automatical triggers for training new models on performance degradation. These tools could be added to ML workflow later, as they are less important and should be built on the foundation of a solid automated ML training pipeline [9].

After going through the mentioned articles, online-accessible documentation [1; 11; 18] hosted by public cloud providers six tools were selected for comparison.

Inclusion criteria:

1. Adoption – The tool must be already adopted and used in the industry, successful case studies should exist for it. Integration of a new unproven solution into the platform may be difficult, because of a lack of documentation and features [14]. Good adoption also usually means it's easier to hire expertise.

2. Feature set – the more parts of the MLOps pipeline can be covered by the same toolset/framework, the fewer new dependencies the project would need to maintain.

3. Interoperability – what ML framework tool supports, what cloud provider integrations exist out of the box, etc.

4. Licensing and cost – is it an open-source solution that can be self-hosted or closed-source privately owned solution only available on vendor's cloud?

Exclusion criteria:

1. Obsolete tools – tools that are no longer maintained or replaced.

Evaluation Framework. Open source tools would be compared on initial configuration and maintenance cost, license limitations, public cloud feature integrations suite, and feature completeness.

Official documentation, public git repositories and market research tools (6sense) were used to estimate each metric's value.

Limitations. The Availability of existing case studies of MLOps applications in the VPP market is limited, so some metric values were approximate estimates.

This article does not evaluate the ease of implementation for an edge IoT MLOps solution, mostly focusing on VPP that works as a hosted backend solution on the public cloud.

Results.

Table 1 contains the final results for the comparison based on selected metrics.

Table 1

MLOps tools comparison

Name	Adoption	Feature set	Interoperability	Licensing & cost
Kubeflow	3	4	1	1
Airflow	1	1	1	1
MLFlow	2	1	1	1
Azure ML	4	2	1	3
AWS SageMaker	5	2	1	3
Google Vertex AI	7	2	1	2

Fig. 1 shows heatmap representing the scores.

Fig. 2 shows normalized comparison of selected MLOps frameworks.

Adoption. Each column contains a place in which a corresponding tool would fit in sorted based on the selected metric (lower is better).

For open-source tools adoption was estimated based on development activity on publicly hosted Git repositories [6]. Airflow is the most mature of the selected subset of MLOps automation tools and is still the most popular and actively supported.

For closed-source tools hosted on public cloud 6sense [2] estimates were used. Azure ML is almost twice as popular as AWS Sagemaker as of this writing. Google Vertex AI was launched a few years later than all of the above tools and market adoption data for it is not yet available, so the authors assume the worst adoption for it in comparison to other tools from the list.

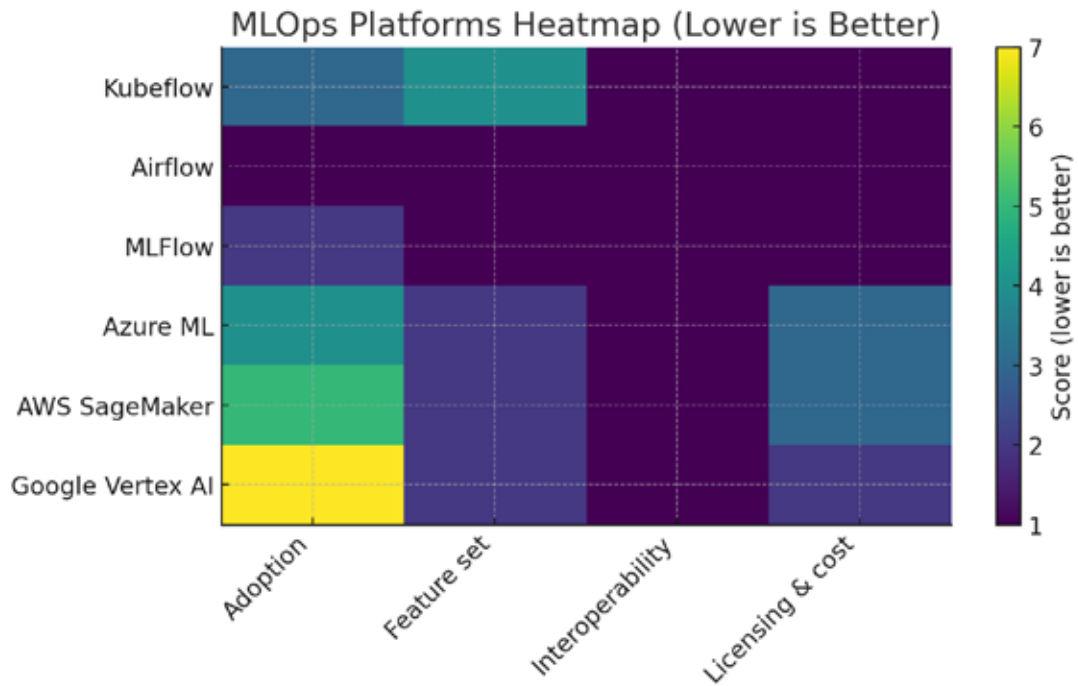


Fig. 1. Heatmap of scores

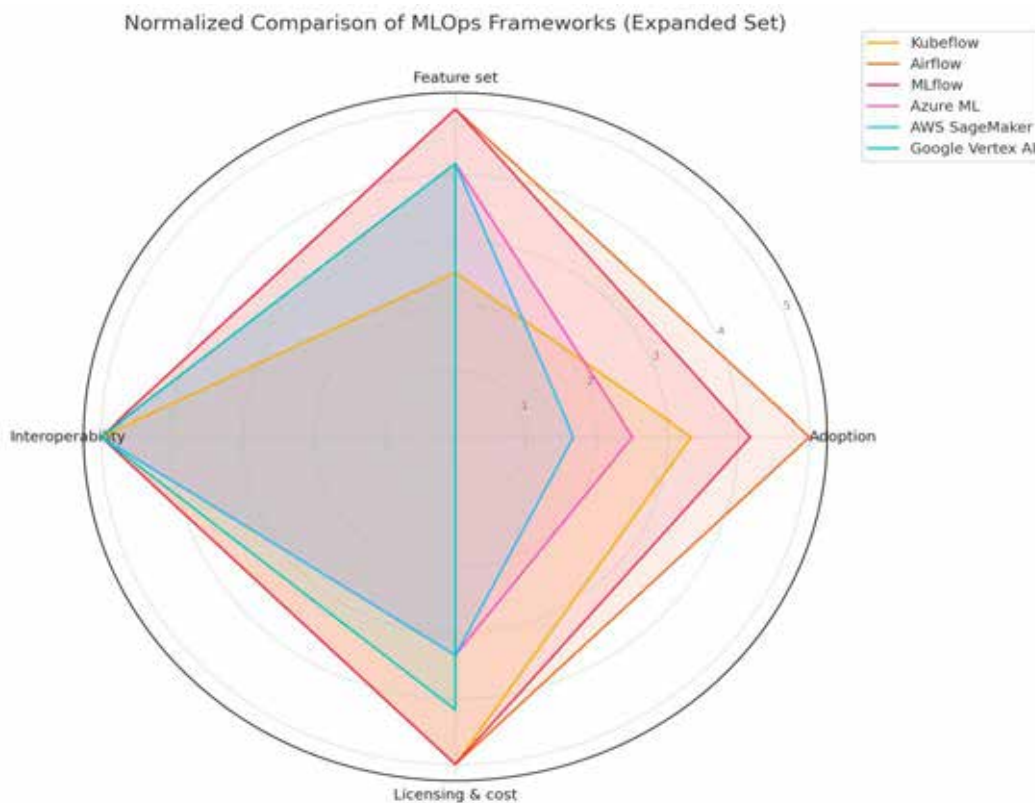


Fig. 2. Normalized comparison of MLOps frameworks

Feature set. Each column contains a place in which a corresponding tool would fit in sorted based on the selected metric (lower is better). Duplicate values represent equal/comparable feature sets.

Airflow and MLFlow are very mature open-source tools and any MLOps-related feature we might need (parametrised training, integration with orchestrators like Kubernetes etc) are implemented.

Azure ML, AWS SageMaker and Google Vertex AI have better integration with their respective cloud, but this leads to worse support for on-premise workloads or workloads on other cloud providers.

Kubeflow is very actively developed and mostly covers the use-case of training models in Kubernetes-managed environment. It's the easiest one to configure on an existing Kubernetes cluster, but if your environment is not running in Kubernetes orchestrator Kubeflow would be useless for you.

Interoperability. All listed MLOps pipeline automation tools support all major frameworks like TensorFlow and PyTorch.

Licensing & cost. Open-source tools allow you to use them free in commercial environments.

When comparing closed-source tools, Google Vertex AI is yet the cheapest to train models on.

Recommendations. Based on the results, the authors of the article recommend using Apache Airflow for most of the MLOps VPP use-cases.

If your environment is completely containerized and orchestrated by Kubernetes Kubeflow could be a good solution, but it's less mature than Airflow and Airflow can be run natively in the Kubernetes on its own.

Azure ML, AWS SageMaker and Google Vertex AI support less features and would be more expensive to host. They should be considered in organisations that do not have mature ops/infrastructure engineering teams that could maintain the MLOps pipeline on their own.

Limitations of Results. Market adoption was not measured based on real telemetry from production systems, but just an estimate.

All metrics were selected without any VPP specificity, albeit authors believe these metrics are the most important in most of the MLOps applications including VPP.

Adoption Barriers and Facilitators. MLOps adoption is mainly slowed by:

- organisational challenges (pure data-science driven team would not be able to achieve MLOps on their own, cross-functional team is required [3]);

- cost (ML is already an expensive endeavour for small or even medium-sized companies and pipeline automation might not be worth the cost for many of them if data-science team is small enough).

But, MLOps adoption can be accelerated by increased awareness and new easier to use tooling being built to abstract existing complexities of MLOps pipelines.

Conclusion. VPP would greatly benefit from MLOps integration, it would improve ML models' efficiency and decision-making processes. But, even considering these benefits MLOps is not yet widely used in VPP. It's mostly caused by the operational complexity of such tools and the cost of their introduction.

Multiple open and closed source tools exist to solve MLOps pipeline automation problems, and most of them cover all major features and support main frameworks.

Different MLOps tools have different limitations and their selection should be based on market adoption, feature set and interoperability with ML frameworks your organisation uses.

Open-source tools like Airflow or MLFlow would be best for most use cases, but their usage requires having a cross-functional data science/operations team. In some organisations it might not be a great fit and cloud-provider tools like AWS SageMaker or Azure ML should be used.

As VPPs continue to advance and assume a more critical role in renewable energy integration, the strategic application of MLOps will be vital.

Bibliography:

1. "Azure Machine Learning – ML as a Service | Microsoft Azure." Accessed: Apr. 13, 2024. URL: <https://azure.microsoft.com/en-gb/products/machine-learning>
2. "Azure Machine Learning vs Amazon SageMaker: Data Science And Machine Learning Comparison," 6sense. Accessed: Apr. 13, 2024. URL: <https://www.6sense.com/tech/data-science-and-machine-learning/azuremachinelearning-vs-amazonsagemaker>
3. D. Kreuzberger, N. Köhl, and S. Hirschl, "Machine Learning Operations (MLOps): Overview, Definition, and Architecture," *IEEE Access*, 2023, vol. 11, pp. 31866–31879, doi: 10.1109/ACCESS.2023.3262138.
4. E. Peltonen and S. Dias, "LinkEdge: Open-sourced MLOps Integration with IoT Edge," in Proceedings of the 3rd Eclipse Security, AI, Architecture and Modelling Conference on Cloud to Edge Continuum, Ludwigsburg Germany: ACM, Oct. 2023, pp. 67–76. doi: 10.1145/3624486.3624496.
5. F. Sattar, A. Husnain, and T. Ghaoud, "Integration of Distributed Energy Resources into a Virtual Power Plant-A Pilot Project in Dubai," 2023 IEEE/IAS Industrial and Commercial Power System Asia (I&CPS Asia), pp. 526–531, Jul. 2023, doi: 10.1109/ICPSAsia58343.2023.10294691.
6. G. Robles, A. Capiluppi, J. M. Gonzalez-Barahona, B. Lundell, and J. Gamalielsson, "Development effort estimation in free/open source software from activity in version control systems," *Empir Software Eng*, 2022, vol. 27, no. 6, p. 135, doi: 10.1007/s10664-022-10166-x.
7. H.-M. Chung, S. Maharjan, Y. Zhang, F. Eliassen, and K. Strunz, "Optimal Energy Trading With Demand Responses in Cloud Computing Enabled Virtual Power Plant in Smart Grids," *IEEE Trans. Cloud Comput.*, 2022, vol. 10, no. 1, pp. 17–30, doi: 10.1109/TCC.2021.3118563.

8. J. Rodríguez-García, D. Ribó-Pérez, C. Álvarez-Bel, and E. Peñalvo-López, "Novel Conceptual Architecture for the Next-Generation Electricity Markets to Enhance a Large Penetration of Renewable Energy," *Energies*, 2019. vol. 12, no. 13, Art. no. 13, doi: 10.3390/en12132605.
9. K. Salama, J. Kazmierczak, and D. Schut, "Practitioners guide to MLOps: A framework for continuous delivery and automation of machine learning".
10. M. Merenda, C. Porcaro, and D. Iero, "Edge Machine Learning for AI-Enabled IoT Devices: A Review," *Sensors*, 2020. vol. 20, no. 9, p. 2533, doi: 10.3390/s20092533.
11. "Machine Learning Service – Amazon SageMaker – AWS," Amazon Web Services, Inc. Accessed: Apr. 13, 2024. URL: <https://aws.amazon.com/sagemaker/>
12. "MLOps: Continuous delivery and automation pipelines in machine learning. Cloud Architecture Center," Google Cloud. Accessed: Mar. 16, 2024. URL: <https://cloud.google.com/architecture/mlops-continuous-delivery-and-automation-pipelines-in-machine-learning>
13. N. Naval and J. M. Yusta, "Virtual power plant models and electricity markets – A review," *Renewable and Sustainable Energy Reviews*, 2021. vol. 149, p. 111393, doi: 10.1016/j.rser.2021.111393.
14. R. S. Kenett, X. Franch, A. Susi, and N. Galanis, "Adoption of Free Libre Open Source Software (FLOSS): A Risk Management Perspective," 2014 IEEE 38th Annual Computer Software and Applications Conference, pp. 171–180, Jul. 2014, doi: 10.1109/COMPSAC.2014.25.
15. R. Subramanya, S. Sierla, and V. Vyatkin, "From DevOps to MLOps: Overview and Application to Electricity Market Forecasting," *Applied Sciences*, 2022. vol. 12, no. 19, Art. no. 19, doi: 10.3390/app12199851.
16. T.-Y. Kim and S.-B. Cho, "Predicting residential energy consumption using CNN-LSTM neural networks," *Energy*, 2019. vol. 182, pp. 72–81, Sep. doi: 10.1016/j.energy.2019.05.230.
17. V. Omelchenko and V. Manokhin, "Optimal Balancing of Wind Parks with Virtual Power Plants," in *Frontiers in Energy Research*, Nov. 2021, p. 665295. doi: 10.3389/fenrg.2021.665295.
18. "Vertex AI," Google Cloud. Accessed: Apr. 13, 2024. URL: <https://cloud.google.com/vertex-ai>
19. W. Shi, J. Cao, Q. Zhang, Y. Li, and L. Xu, "Edge Computing: Vision and Challenges," *IEEE Internet Things J.*, vol. 3, no. 5, pp. 637–646, Oct. 2016, doi: 10.1109/JIOT.2016.2579198.

Дата надходження статті: 24.06.2025

Дата прийняття статті: 30.06.2025

Опубліковано: 23.09.2025

УДК 004.932.2:004.056.55

DOI <https://doi.org/10.32689/maup.it.2025.2.13>

Євген ЛАНСЬКИХ

кандидат технічних наук, доцент,
доцент кафедри інформаційних технологій проектування,
Черкаський державний технологічний університет,
y.lanskykh@chdtu.edu.ua

ORCID: 0000-0003-3389-5720

Олексій НАУМОВ

аспірант кафедри інформаційних технологій проектування,
Черкаський державний технологічний університет,
o.m.naumov.asp23@chdtu.edu.ua

ORCID: 0009-0004-7573-3871

СУЧАСНІ МЕТОДИ ЦИФРОВОГО ВОДЯНОГО МАРКУВАННЯ ЗОБРАЖЕНЬ: КЛАСИФІКАЦІЯ, АНАЛІЗ І ТЕНДЕНЦІЇ РОЗВИТКУ

Анотація. У статті представлено комплексний огляд сучасних методів цифрового водяного маркування зображень, що використовуються для забезпечення авторських прав, перевірки автентичності та відстеження джерел розповсюдження цифрового контенту.

Мета роботи полягає в систематизації підходів до цифрового маркування зображень, зокрема традиційних, гібридних і заснованих на глибокому навчанні, а також у виявленні їх технічних характеристик, сильних і слабких сторін, практичної ефективності та науково-технологічних викликів у контексті цифрової безпеки.

Методологія дослідження базується на системному аналізі актуальних наукових джерел, включаючи рецензовані публікації, оглядові статті, експериментальні дослідження та прикладні реалізації алгоритмів цифрового водяного маркування. Для класифікації методів використано низку ключових технічних і функціональних критеріїв, зокрема: стійкість до атак, прозорість маркування, складність реалізації, обчислювальна ефективність, можливість сліпого виявлення (blind detection), адаптивність до типу зображень, а також широта сфер застосування. Порівняння методів здійснювалося з урахуванням їх здатності забезпечувати баланс між цими характеристиками, що дозволяє визначити оптимальні підходи для практичного застосування в умовах зростаючих вимог до цифрової безпеки, приватності та захисту інтелектуальної власності.

Наукова новизна роботи полягає у цілісному узагальненні еволюції методів водяного маркування із фокусом на сучасні технологічні тенденції, серед яких виокремлюються: інтеграція глибокого навчання для підвищення автономності алгоритмів; поєднання водяного маркування з криптографією та блокчейном; розробка адаптивних рішень, здатних масштабуватися під різні формати зображень і умови використання.

Висновки. Цифрове водяне маркування залишається стратегічно важливим інструментом захисту цифрового контенту. Проаналізовані методи демонструють широкий спектр рішень із різним ступенем прозорості, стійкості до атак та обчислювальної ефективності. Традиційні підходи, засновані на обробці у просторі або частотній області, є добре вивченими, але мають обмежену адаптивність. Гібридні методи поєднують переваги класичних рішень, проте потребують оптимізації для протидії сучасним атакам. Методології на основі глибокого навчання забезпечують новий рівень функціональності, зокрема можливість маркування без доступу до оригіналу (blind watermarking), однак стикаються з викликами, пов'язаними з високими ресурсними витратами та браком стандартів. У контексті цифрової трансформації суспільства актуальними стають дослідження, спрямовані на розробку надійних, масштабованих і нормативно врегульованих систем водяного маркування, здатних ефективно функціонувати в умовах зростаючих вимог до безпеки, приватності та інтероперабельності цифрових платформ.

Ключові слова: цифрове водяне маркування, захист авторських прав, стійкість до атак, гібридні методи, глибоке навчання, сліпе маркування, цифрові зображення.

Yevhen LANSKYKH, Oleksii NAUMOV. MODERN METHODS OF DIGITAL IMAGE WATERMARKING: CLASSIFICATION, ANALYSIS AND DEVELOPMENT TRENDS

Abstract. This article presents a comprehensive review of contemporary digital image watermarking methods used to ensure copyright protection, verify authenticity, and trace the sources of digital content distribution.

The purpose of the study is to systematize approaches to digital image watermarking, including traditional, hybrid, and deep learning-based methods, and to identify their technical characteristics, strengths and weaknesses, practical efficiency, and scientific and technological challenges in the context of digital security.

The methodology is based on a systematic analysis of current scientific sources, including peer-reviewed publications, review articles, experimental studies, and applied implementations of digital watermarking algorithms. The classification of methods was conducted using a set of key technical and functional criteria, including: robustness against attacks, watermark

© Є. Ланських, О. Наумов, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

transparency, implementation complexity, computational efficiency, the possibility of blind detection, adaptability to image types, and the breadth of application areas. The comparative analysis focused on assessing the balance between these characteristics, which enables the identification of optimal approaches for practical application in conditions of increasing demands for digital security, privacy, and intellectual property protection.

The scientific novelty of the study lies in the integrated generalization of the evolution of watermarking methods, with a particular focus on current technological trends. These include the integration of deep learning to increase the autonomy of algorithms; combining watermarking with cryptographic and blockchain technologies; and the development of adaptive solutions capable of scaling across different image formats and application scenarios.

Conclusions. Digital watermarking remains a strategically important tool for protecting digital content. The analyzed methods demonstrate a wide range of solutions with varying degrees of transparency, robustness against attacks, and computational performance. Traditional approaches, based on spatial or frequency domain processing, are well-studied but have limited adaptability. Hybrid methods combine the advantages of classical techniques but require further optimization to withstand sophisticated attacks. Deep learning-based methodologies offer a new level of functionality, including blind watermarking (embedding and extraction without access to the original image), yet face challenges such as high computational costs and the lack of standardized evaluation protocols. In the context of the ongoing digital transformation of society, it is crucial to develop reliable, scalable, and legally regulated watermarking systems capable of functioning effectively under growing requirements for security, privacy, and interoperability across digital platforms.

Key words: digital watermarking, copyright protection, attack robustness, hybrid methods, deep learning, blind watermarking, digital images.

Постановка проблеми у загальному вигляді та її зв'язок із важливими науковими чи практичними завданнями. Захист авторських прав, автентичність цифрових зображень та протидія несанкціонованому розповсюдженню мультимедійного контенту є критично важливими завданнями в умовах сучасного інформаційного середовища. Поширення цифрових технологій у поєднанні з легкістю копіювання даних значно ускладнило захист інтелектуальної власності, що актуалізувало потребу в надійних і універсальних засобах цифрового маркування.

Цифрові водяні знаки (ЦВЗ) виступають ключовим інструментом для забезпечення захисту зображень, оскільки дають змогу вбудовувати приховану інформацію без істотного впливу на візуальну якість. Серед найважливіших наукових та практичних завдань, які розв'язуються через використання ЦВЗ, варто виділити:

- ідентифікацію та підтвердження авторства зображень;
- виявлення несанкціонованих змін або редагування;
- простежування джерела поширення контенту;
- використання в цифрових архівах, судово-експертній практиці та блокчейн-системах.

Однак впровадження таких технологій супроводжується низкою технічних проблем: забезпечення стійкості до атак (обрізання, фільтрація, масштабування), збереження якості зображення, швидкість обробки та адаптивність до різних умов.

Аналіз останніх досліджень і публікацій. Аналіз сучасних публікацій дозволяє виокремити кілька провідних напрямів розвитку методів цифрового водяного знакування (ЦВЗ). Одним із таких напрямів є оглядові та узагальнюючі праці, в яких проаналізовано десятки підходів до ЦВЗ. Зокрема, в оглядах [1], [2], [7] розглянуто методи, що базуються на дискретному вейвлет-перетворенні (DWT), косинусному перетворенні (DCT), найменш значущому біті (LSB) та застосуванні хаотичних карт.

Ряд публікацій [6; 13; 15] зосереджується на комплексному аналізі сучасних технологій цифрового маркування. Так, Hind M. Al-Dabbas та співавтори систематизують наявні рішення та можливі атаки на ЦВЗ. У дослідженні Shweta Wadhera і колег представлено напрями застосування ЦВЗ у сфері охорони здоров'я, освіти та інформаційної безпеки. У свою чергу, Wang та співавтори пропонують інтеграцію цифрового маркування зі стеганографією на основі глибокого навчання, що відображає міждисциплінарний підхід до вирішення проблем захисту цифрового контенту.

Інший напрям досліджень охоплює традиційні та гібридні підходи. У роботах [3] і [4] розглядаються як класичні, так і хаотичні алгоритми маркування, зокрема із використанням відображення Хенона та карти kota Арнольда. Такі рішення спрямовані на підвищення стійкості до атак, зберігаючи водяний знак у зміненому середовищі. Дослідження [5] підтверджує ефективність гібридних схем, у яких поєднуються методи сингулярного розкладання (SVD) із блочними перетвореннями, що дозволяє досягти стійкості без суттєвого зниження прозорості. У статті [16] описано чотириступеневий алгоритм, який комбінує DWT, DCT, SVD та скремблювання, демонструючи високу надійність навіть у випадках зниження роздільності зображення.

Окремий напрям становлять методи, засновані на глибокому навчанні. У роботі Zhong та ін. [17] представлено архітектуру, здатну здійснювати вбудовування водяного знаку в режимі «blind», тобто без доступу до оригінального зображення. Запропоноване рішення забезпечує стійкість до широкого спектру атак, включаючи навіть фотозйомку екрана. Схожий підхід описано в [10], де акцент зроблено

на підвищення геометричної стійкості маркування. Тренована нейромережева модель демонструє здатність зберігати водяний знак після трансформацій, таких як зміна кута, масштабування чи обертання.

Таким чином, наведені джерела демонструють як широту наукових підходів до цифрового водяного маркування, так і нагальну потребу в уніфікованих, ефективних і універсальних рішеннях, здатних функціонувати в умовах складного й змінного цифрового середовища.

Мета статті – систематизація сучасних методів цифрового маркування зображень; аналіз технічних та алгоритмічних підходів (традиційних, гібридних і на базі глибокого навчання); виявлення сильних і слабких сторін кожної групи методів; а також виявлення наукових і технологічних викликів, пов'язаних із захистом контенту в умовах швидкого розвитку цифрових платформ. Особлива увага приділяється стійкості алгоритмів до атак, збереженню візуальної якості зображення, обчислювальній ефективності методів та можливостям застосування у практичних галузях – кібербезпека, цифрова ідентифікація, захист контенту.

Виклад основного матеріалу. Методи цифрового водяного маркування класифікують за різними критеріями, зокрема за способом вбудовування, рівнем перетворення, способом отримання інформації, типом носія, а також за використанням інтелектуальних технологій. З метою систематизації сучасних підходів до цифрового водяного маркування (ЦВМ) зображень доцільно виділити три основні категорії: традиційні, гібридні та методи на основі глибокого навчання (deep learning, DL). Такий поділ базується як на історії розвитку методів, так і на технічних особливостях реалізації алгоритмів, що дозволяє більш об'єктивно провести їх порівняльний аналіз та оцінити перспективи практичного застосування. В оглядовій літературі [1; 2; 6; 7; 13; 15] і практичних реалізаціях можна виокремити три основні категорії методів.

Перша група методів – це традиційні методи цифрового водяного маркування, які ґрунтуються переважно на використанні класичних математичних перетворень або маніпуляціях на рівні пікселів. До них належать просторові методи, що передбачають модифікацію бітів зображення (наприклад, LSB-вставка) [7], частотні методи, що використовують дискретні перетворення (DCT, DWT, FFT), де водяний знак вбудовується у спектральні компоненти [1; 2], а також хаотичні системи, що застосовують детермінований хаос для шифрування або позиціонування водяного знака [3; 4]. Традиційні методи відзначаються відносною простотою реалізації та низькими обчислювальними витратами, що робить їх придатними для пристроїв із обмеженими ресурсами [1; 2; 4; 6]. Зокрема, метод на основі логістичного хаотичного відображення продемонстрував ефективність у забезпеченні початкового рівня стійкості та безпеки вбудованого маркування [4]. Ці підходи залишаються актуальними для задач, де потрібна швидка і легка реалізація без високих вимог до безпеки, зокрема в захисті авторських прав у відкритих базах даних. Однак основним недоліком цієї групи методів є їх низька адаптивність до складних атак, зокрема геометричних спотворень, обтинання та втрат стиснення JPEG.

Друга група – це гібридні методи, які поєднують переваги кількох трансформаційних доменів і додаткових механізмів, що дозволяє підвищити як стійкість, так і візуальну якість [3; 5; 16]. Також гібридні підходи інтегрують додаткові техніки – хаотичні перетворення, криптографічні алгоритми, статистичні моделі, алгоритми оптимізації (генетичні, ролі частинок) [3; 15; 16]. Наприклад, комбінація перетворень (DWT+DCT+SVD) дозволяє балансувати між прозорістю та стійкістю [5; 16], а комбінування хаотичних карт і частотних алгоритмів підвищує стійкість і ускладнює несанкціоноване вилучення знака. Так, у роботі [3] запропоновано комбіноване використання хаотичних перетворень для генерації стійкого водяного знака, що ускладнює зворотну інженерію та підвищує криптостійкість. Гібридні методи часто мають вищу продуктивність, їх перевагами є можливість налаштування під конкретні потреби користувача та досягнення балансу між прозорістю, стійкістю та продуктивністю. Ці методи ефективно застосовуються у сфері цифрової ідентифікації, медичних зображень, судової експертизи. Проте складність реалізації таких систем значно вища, як і обчислювальна складність, особливо у випадках використання багаторівневих перетворень, адже вони потребують ретельного налаштування.

Третя група – це методи цифрового маркування, що базуються на глибокому навчанні, яке останнім часом активно розвивається. Цей напрям є найновішим і передбачає використання нейронних мереж, таких як автоенкодера, згорткові нейронні мережі (CNN), GAN-архітектури та трансформери для навчання вбудовування та виявлення водяного знака [8; 10; 11; 13; 14; 15; 17]. Перевагою цих методів є автоматична адаптація до контенту зображення, висока ємність маркування та стійкість до складних атак, зокрема до геометричних трансформацій, шуму та перекодування [8; 10; 14]. Також вони використовуються для вбудовування та видобування знака у «blind» режимі. Наприклад, у роботі [10] запропоновано сліпу схему маркування зображень на основі глибокого навчання, яка демонструє стійкість до змін масштабу, обертання та інших геометричних деформацій. Разом з тим, такі методи

є ресурсоемними, потребують великих обсягів даних для навчання, а також характеризуються низькою інтерпретованістю, що може ускладнювати їх використання в критично важливих додатках (наприклад, цифрова судова експертиза). Також існують мережі, стійкі до геометричних викривлень, що дозволяють зберегти водяний знак після атак типу обертання, масштабування, обрізання [10]. Загалом ці підходи демонструють високу адаптивність до різних форматів зображень і атак, відкривають нові можливості для інтеграції ЦВМ у розумні системи, захист AI-контенту, deepfake-детекцію, NFT тощо, але потребують великих обсягів даних і ресурсів для навчання.

Класифікація методів цифрового водяного маркування на традиційні, гібридні та засновані на глибокому навчанні дозволяє систематизувати сучасні підходи відповідно до їхніх технічних характеристик, ефективності та застосовності в умовах динамічного цифрового середовища. Такий поділ також сприяє кращому розумінню сильних і слабких сторін кожної групи методів, що є необхідним для обґрунтованого вибору рішень у практичних задачах – від захисту авторських прав до протидії нелегальному тиражуванню цифрового контенту. Така класифікація дозволяє ефективно орієнтуватися в поточному стані досліджень, порівнювати алгоритми за релевантними критеріями та прогнозувати вектори подальшого розвитку технологій цифрового захисту зображень.

З урахуванням швидкого розвитку алгоритмів та зростання вимог до захисту мультимедійного контенту, порівняльний аналіз основних підходів до цифрового водяного маркування є ключовим для виявлення їх сильних і слабких сторін. Для забезпечення всебічної та об'єктивної оцінки сучасних методів було обрано сім ключових критеріїв, які охоплюють як технічні характеристики, так і здатність адаптуватися до прикладних потреб цифрового середовища.

Одним із базових параметрів є стійкість до атак. Вона передбачає здатність водяного знака зберегтися після типових спотворень, таких як стиснення JPEG і JPEG2000, фільтрація, обертання, масштабування, обтинання або додавання шуму. Цей критерій є критично важливим для захисту авторських прав та довготривалого збереження маркування у змінному цифровому середовищі [6; 7, 14].

Не менш важливою є прозорість, яка визначає ступінь впливу водяного знака на візуальну якість зображення. Ідеальним вважається маркування, яке є повністю невидимим для людського ока, що особливо актуально для використання у високоякісних зображеннях, наприклад у медицині, поліграфії чи цифровому мистецтві [16; 17].

Складність реалізації також є суттєвим фактором, оскільки вона впливає на можливість впровадження методу у реальні системи. Вона охоплює як програмну реалізацію, так і алгоритмічну складність, потребу в зовнішніх модулях, а також можливість роботи у розподілених або апаратно обмежених середовищах [4; 13].

Обчислювальна ефективність відіграє ключову роль у застосуваннях, де потрібна обробка великої кількості зображень або функціонування в реальному часі. Вона враховує час, необхідний для вбудовування та виявлення знака, використання пам'яті та можливість реалізації на мобільних чи вбудованих пристроях [8; 15].

Окремо розглядається можливість сліпого виявлення, тобто здатність алгоритму розпізнавати водяний знак без необхідності у збереженні або доступі до оригінального зображення. Такий підхід є особливо важливим для відкритих або розподілених систем, де автентифікація має здійснюватися без зайвих залежностей [3; 5].

Ще одним критичним аспектом є адаптивність до типу зображень. Висока гнучкість у роботі з текстуrowаними, рівномірними, контрастними чи кольоровими зображеннями дозволяє значно розширити сферу застосування цифрового маркування в різних галузях [2; 10].

Останнім критерієм виступає широта сфер застосування. До них належать захист авторських прав, цифрова ідентифікація, інформаційна безпека, автентифікація контенту, використання в NFT або в цифрових паспортах об'єктів тощо [11; 18].

Таким чином, зазначені критерії забезпечують комплексне уявлення про потенціал кожної категорії методів – традиційних, гібридних та заснованих на глибокому навчанні. Вони дозволяють оцінити не лише технічні характеристики, але й практичну доцільність впровадження відповідних рішень у конкретних умовах.

У таблиці нижче узагальнено характеристики традиційних, гібридних та методів на основі глибокого навчання за критичними параметрами, що визначають їхню практичну цінність.

Аналіз показує, що хоча традиційні методи залишаються актуальними завдяки простоті, високій якості зображень після маркування та невеликим обчислювальним витратам, вони поступаються іншим за стійкістю до складних атак та адаптивністю. Гібридні методи демонструють найкращий баланс між прозорістю, стійкістю та універсальністю. Вони добре масштабуються та часто забезпечують підтримку сліпого виявлення, проте їх реалізація є більш складною. Методи на основі глибокого навчання

мають найвищий потенціал, особливо в контексті автоматизації, адаптивного вбудовування та стійкості до складних атак. Вони демонструють найвищу стійкість та гнучкість, проте вимагають значних ресурсів для навчання і розгортання. Їхні обмеження – висока обчислювальна складність та потреба в навчальних даних. Таким чином, вибір оптимального підходу залежить від конкретних вимог до захисту зображень, доступних обчислювальних ресурсів та умов використання.

Таблиця 1

Порівняльний аналіз методів цифрового водяного маркування

Критерій	Традиційні методи	Гібридні методи	Методи на основі глибокого навчання
Стійкість до атак	Середня. Добре працюють проти простих атак, вразливі до геометричних трансформацій	Висока. Забезпечують надійність до широкого спектра атак, включаючи фільтрацію, стиснення, масштабування	Дуже висока. Навчені мережі можуть зберігати знак навіть після фотозйомки або перекодування
Прозорість (якість зображення)	Висока (особливо в частотних методах)	Висока або середня (залежно від схеми комбінування)	Висока, але може залежати від архітектури моделі
Складність реалізації	Низька або середня. Реалізація не потребує значних обчислювальних ресурсів	Середня або висока. Потребують ретельного налаштування та комбінування алгоритмів	Висока. Необхідне навчання моделей, підбір гіперпараметрів та використання GPU
Обчислювальна ефективність	Висока. Швидка обробка навіть на слабкому обладнанні	Середня. Залежить від кількості етапів та складності перетворень	Залежить від моделі: може бути швидким, але навчання – ресурсоємне
Можливість сліпого виявлення (blind detection)	Часткова підтримка	Так, реалізується в багатьох схемах	Так. Один з основних плюсів DL-підходів
Адаптивність до типу зображень	Обмежена. Потребують ручного налаштування під кожний тип зображення	Краща адаптивність завдяки комбінуванню підходів	Висока. DL-моделі можуть навчатися на різних типах даних
Сфери застосування	Захист авторських прав, цифрові архіви	Судово-експертні системи, стійке маркування медіа, блокчейн	Кібербезпека, маркування в соцмережах, протидія фейкам, цифрова ідентифікація

Сфера цифрового водяного маркування динамічно розвивається під впливом новітніх технологій, зокрема штучного інтелекту, блокчейну та генеративних моделей. Попри значний прогрес, низка наукових і прикладних проблем залишається актуальною, а нові виклики постають у зв'язку з ускладненням цифрового середовища. Сучасні тенденції розвитку цифрового водяного маркування (ЦВМ) відкривають нові можливості, але водночас потребують переосмислення підходів, методів і архітектур.

Одним із ключових напрямів є застосування глибокого навчання як рушійної сили інтелектуального маркування. Алгоритми на основі DL забезпечують високий рівень стійкості до складних атак і здатні зберігати вбудований водяний знак навіть після суттєвих спотворень зображення [8; 14]. Проте висока ефективність таких методів супроводжується значними обчислювальними витратами, що ускладнює їх використання на мобільних або вбудованих пристроях. Окрім того, моделі DL мають високу складність, ризик переадаптації до специфічних даних і потребують ретельного навчання [8].

Іншою важливою проблемою стало протистояння генеративним моделям, зокрема diffusion-based генераторам і deepfake-контенту. Сучасні генератори здатні регенерувати зображення настільки глибоко, що традиційні водяні знаки зникають. У відповідь на це досліджується маркування у латентному просторі моделей, яке дозволяє впроваджувати знаки, що зберігаються навіть після генеративних атак. Також розробляються підходи до модифікації самих генераторів з метою вбудованої автентифікації [10; 11].

У контексті децентралізованих технологій ЦВМ знаходить нове застосування в екосистемі Web3 та NFT. Тут цифрові знаки працюють разом із криптографічними сертифікатами для підтвердження авторства, права власності та валідації об'єктів через блокчейн [18]. Поєднання ЦВМ і NFT підвищує прозорість походження зображень, але водночас створює нові виклики, зокрема щодо масштабованості, сумісності з платформами та стійкості до повторного кодування або републікації.

Окрему увагу дослідників привертає етичний аспект цифрового маркування, особливо в чутливих сферах – медицині, юриспруденції, соціальних мережах. Виникає потреба у балансі між захистом і приватністю. Одним із інноваційних рішень у цьому напрямі є концепція *revocable watermarking* – маркування, яке можна відкликати, деактивувати або редагувати за ініціативою правовласника [9].

Цифрове водяне маркування дедалі ширше застосовується для автоматизованого виявлення порушень авторських прав. Сучасні гібридні системи, що поєднують традиційні методи з компонентами глибокого навчання, здатні виявляти маркування навіть у потоковому відео або на платформах з масовим обміном контентом (YouTube, Facebook, Instagram) [6]. Перспективним напрямом стає розвиток *watermark-as-a-service (WaaS)* – хмарних сервісів, інтегрованих із системами керування цифровим контентом.

Зростання роздільної здатності зображень і потреба в обробці в реальному часі висувують нові вимоги до обчислювальної ефективності. Традиційні складні DL-архітектури замінюються спрощеними моделями, *wavelet*-перетвореннями, хаотичними функціями та методами енергоефективного аналізу. Технології *TinyML* та *Edge AI* дозволяють реалізовувати ЦВМ на пристроях із обмеженими ресурсами, зберігаючи належну якість та стійкість [12].

Таким чином, подальший розвиток цифрового водяного маркування відбуватиметься у напрямі збалансованої інтеграції новітніх технологій, таких як DL, blockchain, генеративні моделі та *edge-computing*.

Підсумовуючи викладене, можна виокремити низку ключових викликів, що наразі стоять перед дослідницькою спільнотою у сфері ЦВМ. Однією з найбільш гострих проблем залишається відсутність уніфікованих стандартів. Недостатня узгодженість форматів, метрик якості та процедур верифікації ускладнює інтеграцію маркування в міжплатформенні системи та юридичні середовища.

Крім того, зростає потреба у забезпеченні стійкості до новітніх типів атак, зокрема до *deepfake*, адаптивних фільтрів та AI-реставрації. Незважаючи на досягнення гібридних та DL-підходів, ці загрози становлять серйозну перешкоду для гарантування автентичності цифрового контенту. Ще одним не вирішеним питанням залишається баланс між прозорістю та стійкістю. Досягнення високих показників обох характеристик без втрати візуальної якості зображення є технічно складним завданням, особливо у випадку високодинамічного або деталізованого контенту.

Значною проблемою також є обмеження обчислювальних ресурсів. Сучасні методи на основі глибокого навчання демонструють високу ефективність, проте вимагають значних обчислювальних потужностей для тренування і розгортання, що обмежує їх використання в енергоефективних або вбудованих системах. Крім того, нові моделі штучного інтелекту вже здатні виконувати «очищення» зображень від водяних знаків, включаючи навіть складні маркування. Це формує нагальну потребу в розробці рішень, здатних протистояти нейромережевим методам фільтрації.

На тлі вказаних викликів можна окреслити перспективні напрями подальшого розвитку. По-перше, інтеграція цифрового маркування з криптографічними протоколами та блокчейн-технологіями дозволяє забезпечити трасованість авторства і фіксацію змін контенту. По-друге, перспективними є адаптивні системи, засновані на самонавчанні, здатні реагувати на зміну типів атак та умов функціонування в режимі реального часу.

Ще один важливий вектор розвитку пов'язаний із поширенням ЦВМ на мультимодальні дані. Сучасні дослідження поступово виходять за межі маркування лише зображень, зосереджуючись також на відео, аудіо та 3D-контенті, що потребує універсальних алгоритмів. У правовому аспекті зростає значення цифрових водяних знаків як доказового інструмента, що актуалізує потребу в їх формальній валідації, експертній оцінці та правовому визнанні.

Нарешті, розвиток концепції *explainable AI (XAI)* у сфері DL-маркування сприяє підвищенню прозорості роботи глибоких моделей. Це, у свою чергу, дозволяє краще оцінювати ризики, пояснювати процеси вбудовування та виявлення маркування, а також формувати довіру до таких технологій з боку користувачів і регуляторів.

Висновки. У цій статті було проведено систематизацію сучасних методів цифрового водяного маркування зображень, здійснено їх порівняльний аналіз та окреслено актуальні тенденції розвитку. На основі вивчених джерел можна зробити такі узагальнення.

Цифрове водяне маркування є ключовим інструментом захисту цифрових зображень, оскільки забезпечує ідентифікацію авторства, верифікацію автентичності та простеження джерел поширення контенту. Методи цифрового водяного маркування за своїми характеристиками та підходами поділяються на традиційні, гібридні та засновані на глибокому навчанні.

Традиційні методи працюють у просторі або частотній області, тоді як гібридні поєднують переваги різних підходів, демонструючи баланс між прозорістю, стійкістю та обчислювальною ефективністю,

хоча й залишаються чутливими до складних атак. У свою чергу, методи на базі глибокого навчання відкривають нові можливості, зокрема для маркування в умовах відсутності оригіналу зображення (blind watermarking), проте вимагають значних ресурсів і поки що не мають загальноприйнятих стандартів.

У зв'язку з цим подальші дослідження доцільно спрямовувати на розробку адаптивних, масштабованих і стійких рішень, зокрема із використанням блокчейн- та криптографічних технологій, а також на правову стандартизацію застосування цифрових водяних знаків у доказових процесах.

Таким чином, цифрове водяне маркування продовжує залишатися активною та стратегічно важливою сферою наукових досліджень, із суттєвим потенціалом до практичного застосування в умовах цифрової трансформації суспільства.

Список використаних джерел:

1. Гавриленко В. О. Дослідження методів побудови цифрових водяних знаків для захисту зображень. 2016. URL: <https://dglb.nubip.edu.ua/handle/123456789/1868>.
2. Лозинський С. С. Формування цифрового водяного знаку для захисту авторських прав на графічні об'єкти. *Наукові праці НУХТ*, 2020. URL: <https://openarchive.nure.ua/handle/document/30248>.
3. Сімонов С. О., Іванов С. М., Романенко О. В. Побудова цифрового водяного знака з використанням хаотичних перетворень. *Східно-Європейський журнал передових технологій*. 2021. № 6/9 (114), с. 6–16. DOI: 10.15587/1729-4061.2021.246641.
4. Чекан І. І., Піскорський В. В. Метод вбудовування водяного знака у зображення на основі логістичного хаотичного відображення. *Інформація і керування*. 2021. № 3, с. 121–127. DOI: 10.20998/2522-9052.2021.3.15.
5. A. F. Eldaoushy et al. Efficient Hybrid Digital Image Watermarking. *Journal of Optics*, 2023. DOI: 10.1007/s12596-023-01144-7.
6. Al-Dabbas H. M., Azeez R. A., Ali A. E. Digital Watermarking, Methodology, Techniques, and Attacks: A Review. *Iraqi Journal of Science*. 2023. Vol. 64, No. 8, pp. 4146–4160. DOI: 10.24996/ij.s.2023.64.8.37.
7. Al-khafaji H., Al-Himyari B. A., Hussain N. F. Techniques for Digital Watermarking in Images: A Review. *Journal of University of Babylon for Pure and Applied Sciences*. 2024. DOI: 10.29196/ycdrt588.
8. Ben Jabra, S., Ben Farah, M. Deep Learning-Based Watermarking Techniques Challenges: A Review of Current and Future Trends. *Circuits, Systems, and Signal Processing*. 2024. 43, 4339–4368. DOI: 10.1007/s00034-024-02651-z.
9. Gaur S., Barthwal V. An Extensive Analysis of Digital Image Watermarking Techniques International. *Journal of Intelligent Systems and Applications in Engineering*. 2024. 12(1), 121–145. DOI: 10.31201/ijisae.2024.12145.
10. Mareen H., Antchougov L., Delvaux J., et al. Blind Deep-Learning-Based Image Watermarking Robust Against Geometric Transformations. *arXiv preprint*, 2024. 12 c. URL: <https://arxiv.org/abs/2402.09062>.
11. Padhi S. K., Ali S. S. DLOVE: A New Security Evaluation Tool for Deep Learning Based Watermarking Techniques. *arXiv preprint*, 2024. DOI: 10.48550/arXiv.2407.06552.
12. Rani S., Kaur P. Optimized and Secure Digital Image Watermarking Technique Using Henon Mapping in Redundant Domain. *Multimedia Tools and Applications*. 2024. 83, 81151–81166. DOI: 10.1007/s11042-024-18643-9.
13. Wadhwa S., Kamra D., Bedi P., et al. A Comprehensive Review on Digital Image Watermarking. *arXiv preprint*, 2022. URL: <https://arxiv.org/abs/2207.06909>.
14. Wang X., Zhang Y., Li H., et al. A Brief, In-Depth Survey of Deep Learning-Based Image Watermarking. *Applied Sciences*, 2023, 13(21), 11852. DOI: 10.3390/app132111852.
15. Wang Z., Byrnes O., Bansal R., et al. Data Hiding with Deep Learning: A Survey Unifying Digital Watermarking and Steganography. *arXiv preprint*, 2021. URL: <https://arxiv.org/abs/2107.09287>.
16. Xie Y., Wang Y., Ma M. Design of a Hybrid Digital Watermarking Algorithm with High Robustness. *Journal of Web Engineering*. 2023. DOI: 10.13052/jwe1540-9589.19567.
17. Zhang J., Yu J., Ma J., Wang Y., Jiang H., Wu Z. A Robust Image Watermarking Method Based on Deep Neural Networks. *arXiv preprint*, 2020. 10 c. URL: <https://arxiv.org/abs/2007.02460>.
18. Zhao, L., Gui, X., Shao, Y., Dai, H. Survey on Multipurpose Digital Image Watermarking. *Journal of Computer-Aided Design & Computer Graphics*, 2024, 36(2), 195–222. DOI: 10.3724/SPJ.1089.2024.20078.

Дата надходження статті: 17.06.2025

Дата прийняття статті: 23.06.2025

Опубліковано: 23.09.2025

УДК 336.71:004.8:005.334:006.83

DOI <https://doi.org/10.32689/maup.it.2025.2.14>

Єлизавета ЛИТЮГА

здобувач вищої освіти,

Сумський державний університет, ye.lytiuha@student.sumdu.edu.ua

ORCID: 0000-0002-6009-5131

Валерій ЯЦЕНКО

кандидат технічних наук, доцент кафедри економічної кібернетики,

Сумський державний університет, v.yatsenko@biem.sumdu.edu.ua

ORCID: 0000-0003-2316-3817

РОЗРОБКА TELEGRAM-БОТА ДЛЯ АВТОМАТИЗАЦІЇ ОБЛІКУ КРИПТОТРАНЗАКЦІЙ

Анотація. Актуальність дослідження зумовлена зростанням обсягів криптовалютних операцій і потребою ефективного контролю обліку та прозорості. В умовах підвищених ризиків помилок і компрометації фінансових даних особливого значення набувають інноваційні методи автоматизації. У статті обґрунтовано ефективність Telegram-ботів у внутрішньому фінансовому контролі криптообмінників як засобу, що мінімізує ручне введення, автоматизує звітність та підвищує оперативність і безпеку операцій.

Мета роботи полягає у проектуванні та розробці Telegram-бота, який забезпечує повний цикл обліку транзакцій з криптовалютами: від збору й валідації даних до формування аналітичної звітності в режимі реального часу, підвищуючи точність, швидкість та масштабованість процесів.

Методологія. Використано порівняльну оцінку ефективності чат-ботового рішення; тестові сценарії охоплювали паралельні запити, стрес-тести і вимірювання часу відповіді. також перевірено сумісність Android, iOS, Web і стійкість до несанкціонованого доступу незареєстрованих користувачів.

Наукова новизна. Вперше проведено комплексне дослідження Telegram-ботів як інструменту автоматизації внутрішнього фінконтролю криптообмінників; обґрунтовано їхню здатність знижувати ризик помилок, прискорювати звітність, підвищувати безпеку й інтегруватися у наявні IT-системи. Запропоновано архітектурну модель з Google Sheets та SQLite.

Висновки. Експериментальне впровадження підтвердило, що Telegram-бот суттєво оптимізує фінансовий контроль криптообмінників: середній час відповіді <2 с, втрати даних не зафіксовано, трудомісткість обліку скорочено удвічі. Система має резерв масштабування й може слугувати прототипом для інших фінтех-сервісів із високою транзакційною активністю.

Ключові слова: криптовалюти, фінансовий контроль, автоматизація, чат-бот, Telegram API, фінансова звітність.

Yelyzaveta LYTIUHA, Valerii YATSENKO. USE OF CHATBOTS IN THE INTERNAL FINANCIAL CONTROL OF A CRYPTO EXCHANGE

Abstract. The relevance of this study stems from the growing volume of cryptocurrency transactions and the ensuing need for effective accounting oversight and transparency. Under heightened risks of errors and financial-data breaches, innovative automation methods acquire particular importance. The article substantiates the effectiveness of Telegram bots in the internal financial control of cryptocurrency exchanges, demonstrating their capacity to minimize manual data entry, automate reporting, and enhance both the timeliness and security of operations.

The aim of this work is to design and develop a Telegram bot that delivers a full transaction-accounting cycle—from data collection and validation to real-time analytical reporting—thereby enhancing accuracy, processing speed, and system scalability.

Methodology. A comparative evaluation of the bot's effectiveness was conducted; test scenarios included parallel-request handling, stress testing, and response-time measurement. Compatibility across Android, iOS, and web clients was verified, as was resistance to unauthorized access by unregistered users.

Scientific novelty. This work presents the first comprehensive investigation of Telegram bots as a means of automating internal financial control in cryptocurrency exchanges, substantiating their capacity to lower error risk, accelerate reporting, enhance security, and integrate seamlessly with existing IT infrastructures. An architectural model incorporating Google Sheets and SQLite is proposed.

Conclusion. Experimental deployment confirmed that the Telegram bot markedly optimizes financial control in cryptocurrency exchanges: average response time is under 2 s, no data loss was detected, and accounting workload was halved. The system retains scalability reserves and can serve as a prototype for other fintech services with high transaction throughput.

Key words: cryptocurrencies, financial control, automation, chatbot, Telegram API, financial reporting.

Постановка проблеми в загальному вигляді та її зв'язок із важливими науковими чи практичними завданнями. Криптовалюти, засновані на блокчейні й криптографії, кардинально змінили фінансові операції, забезпечивши децентралізовані перекази, швидкі розрахунки й нижчі транзакційні витрати. Зі зростанням їх обсягів підприємствам-криптообмінникам потрібен ефективний внутрішній фінансовий контроль: оперативна обробка транзакцій, прозорий облік і надійний захист даних, інакше помилки та затримки можуть спричинити значні втрати й регуляторні санкції.

Аналіз останніх досліджень і публікацій. Кількість наукових публікацій, присвячених використанню чат-ботів у різних сферах, постійно зростає. Згідно з базою даних Scopus за останні дев'ять років опубліковано 11407 наукових робіт на цю тему (рис. 1).

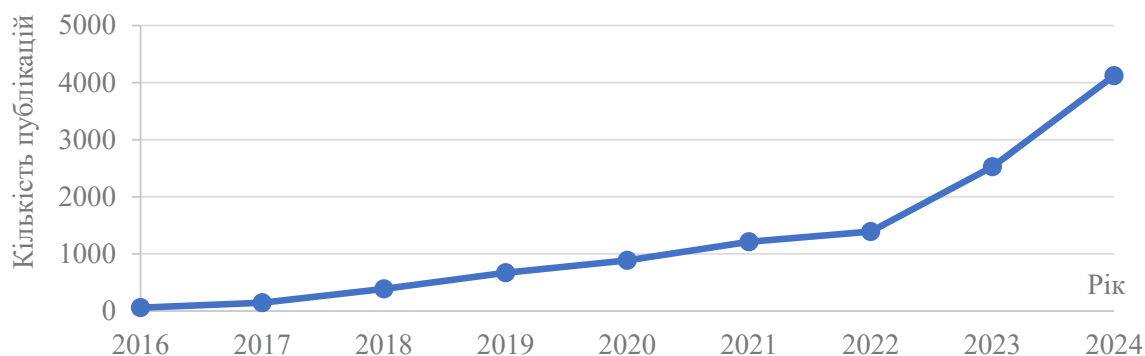


Рис. 1. Динаміка кількості наукових публікацій з тематики чат-ботів в журналах, що індексуються БД Scopus з 2016 по 2024 р.

Найбільшу кількість публікацій в даний період має США – 2003 роботи, на другому місці розташувалась Індія із 1902 працями, а третю позицію посідає Китай – 720 досліджень (рис. 2). Тематика розробки і застосування чат-ботів активно вивчається і українськими науковцями – 97 робіт.

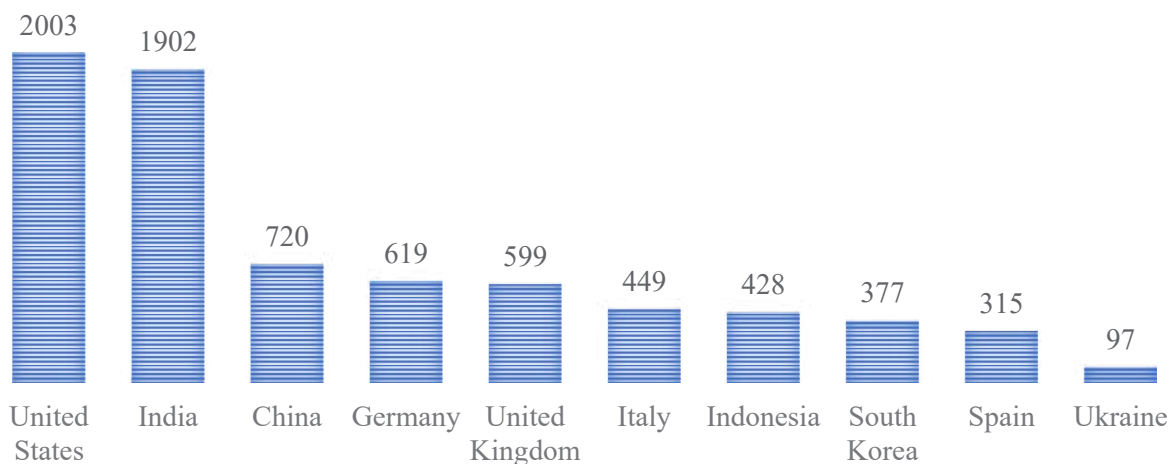


Рис. 2. Кількість публікацій за країнами з тематики чат-ботів за 2016–2024 рр.

Аналіз наукових публікацій бази Scopus за допомогою програмного інструменту для побудови та візуалізації бібліометричних мереж VOSviewer дозволив виокремити шість ключових тематичних кластерів наукових публікацій, пов'язаних із тематикою чат-ботів (рис. 3).

У табл. 1 узагальнено результати бібліометричної кластеризації публікацій, що відображає провідні напрями застосування чат-ботів.

На цьому підґрунті природно постає питання застосування чат-ботів безпосередньо у фінансовому контролі: сучасна наукова література та галузеві публікації все активніше розглядають такі рішення як ефективний інструмент внутрішнього аудиту. Зокрема, дослідження Павелеску (Белбе) А. Е показало, що їх упровадження у державний аудит підвищує якість перевірок і полегшує комунікацію [9],

методологічні переваги чат-ботів, а також розробити та протестувати модель Telegram-бота з інтеграцією Google Sheets і SQLite для обробки операцій, звітності та моніторингу.

Виклад основного матеріалу. Державна політика щодо криптовалюти суттєво відрізняється між юрисдикціями: частина країн офіційно визнає та регулює обіг цифрових активів, тоді як інші запроваджують законодавчі заборони або жорсткі обмеження. Аналітичні огляди показують, що розвинуті держави, зокрема США, Канада й Велика Британія, дозволяють використання біткоїна та інших криптовалют, тоді як Китай, Марокко й Саудівська Аравія повністю заборонили їх обіг [10]. Україна, своєю чергою, обрала шлях активного правового врегулювання: у березні 2022 р. набув чинності Закон «Про віртуальні активи», який визначає статус криптовалют, порядок їх обігу, права власників і вимоги до постачальників послуг [1].

Для обліку криптовалютних транзакцій обмінники застосовують різноманітний інструментарій, що охоплює: комплексні white-label-платформи керування біржею (HollaEx, OpenDAX, AlphaPoint, Binance Broker API); хмарні бухгалтерсько-податкові SaaS-сервіси з криптоінтеграцією (CoinTracking, Koinly, CoinLedger, AccountingSuite); блокчейн-аналітичні та AML/KYC-системи (Chainalysis, Elliptic, CipherTrace); внутрішні чи кастомізовані ERP-рішення. На українському ринку цей перелік доповнюють ad-hoc інструменти – Google Sheets / Excel, конфігурації 1C/BAS і Telegram-боти. Узагальнену класифікацію та ключові функціональні характеристики наведено в (табл. 2).

Таблиця 2

Порівняльна характеристика програмно-технічних рішень для обліку операцій у криптовалютних обмінниках

Категорія	Призначення / функції	Приклади ПЗ	Ключові особливості
Управління криптообмінником	Операційний облік, гарантії, комісії, AML/KYC, звітність	HollaEx, OpenDAX, AlphaPoint, Binance Broker API	Біржова функціональність, гнучка інтеграція
Бухгалтерські та облікові системи	Автоматизований облік транзакцій, формування податкових звітів	CoinTracking, Koinly, CoinLedger, AccountingSuite + криптоінтеграція	API-синхронізація, шаблони декларацій
Blockchain-аналітика / AML / KYC	Відстеження та оцінка ризиків транзакцій	Chainalysis Reactor/KYT, Elliptic, CipherTrace	Ризик-скоринг, аудит, регуляторний комплаєнс
Кастомні ERP-рішення [5; 7]	Комплексний фінансовий і операційний облік у великих біржах	Внутрішні CRM/ERP-системи	Повна масштабованість, налаштованість
Табличні інструменти загального призначення	Швидкий старт, базова аналітика	Excel, Google Sheets	Безкоштовно/умовно безкоштовно, спільна робота онлайн [13]
1C/BAS із криптоінтеграцією	Управлінський облік для малого та середнього бізнесу	Конфігурації 1C / BAS + скрипти	Локальне налаштування під українські вимоги
Повноцінні ERP-системи	Міжнародний облік, інтеграція підрозділів, мінімізація помилок	SAP ERP, Microsoft Dynamics 365, Odoo	Висока вартість, підходить для великих компаній

Малі криптообмінники та ФОП переважно використовують Google Sheets або Excel: нульові витрати, миттєве розгортання й достатня функціональність для обмеженого обсягу транзакцій. White-label біржі та ERP-модулі виправдані лише коли навантаження чи регуляторні вимоги виходять за межі можливостей табличних редакторів.

Для подальшої автоматизації запропоновано Telegram-бот, який інтегрується з біржовими API й базами даних, мінімізує людський фактор і забезпечує оперативну аналітику при мінімальних витратах. Початковий етап дослідження вже ідентифікував недоліки чинного обліку та сформулював системні вимоги; їх стисло узагальнено в (табл. 3).

Наступний етап – розроблення архітектури та реалізації Telegram-бота (рис. 4). Структура системи гарантує надійність, масштабованість і простоту експлуатації.

Рисунок 4 ілюструє ключові компоненти Telegram-бота, що забезпечують реєстрацію, обробку й зберігання фінансових операцій криптообмінника.

Користувачські запити з клієнта Telegram через Telegram API надходять до контролера telegram.py (Python, telebot), який парсить команди, формує інтерфейс і перетворює дані на структурований запит. Далі запит обробляє обчислювальне ядро analytics.py: воно виконує валідацію й класифікацію операцій,

розрахунок показників і формування звітів, звертаючись до гібридного сховища – Google Sheets (хмарна аналітика й резерв) та SQLite (швидке локальне OLTP). Неперервну роботу забезпечує сервісне середовище Python-server, що циклічно опитує Telegram API та синхронно обробляє транзакції.

Таблиця 3

Порівняння недоліків Google Sheets та переваг Telegram-бота для внутрішнього фінансового контролю криптообмінника

Недоліки використання Google Sheets	Можливості Telegram-бота для їх нейтралізації
Ручне введення даних ⇒ помилки	Автоматизоване введення знижує ризик помилки
Відсутність автоматичного оновлення курсу валют	Динамічний перерахунок курсу за середньозваженою формулою на основі поточних операцій
Складність оперативної звітності	Миттєве формування звітів на запит користувача
Немає швидкого перегляду залишків	Швидкий запит залишку коштів безпосередньо в чаті
Високе навантаження на персонал	Системні нагадування й вбудована логіка істотно скорочують ручну працю
Низький рівень захисту даних	Авторизація й контроль доступу через Telegram-ID підвищують безпеку
Відсутність інтеграції з іншими сервісами	Єдине централізоване середовище з потенціалом масштабування та інтеграції

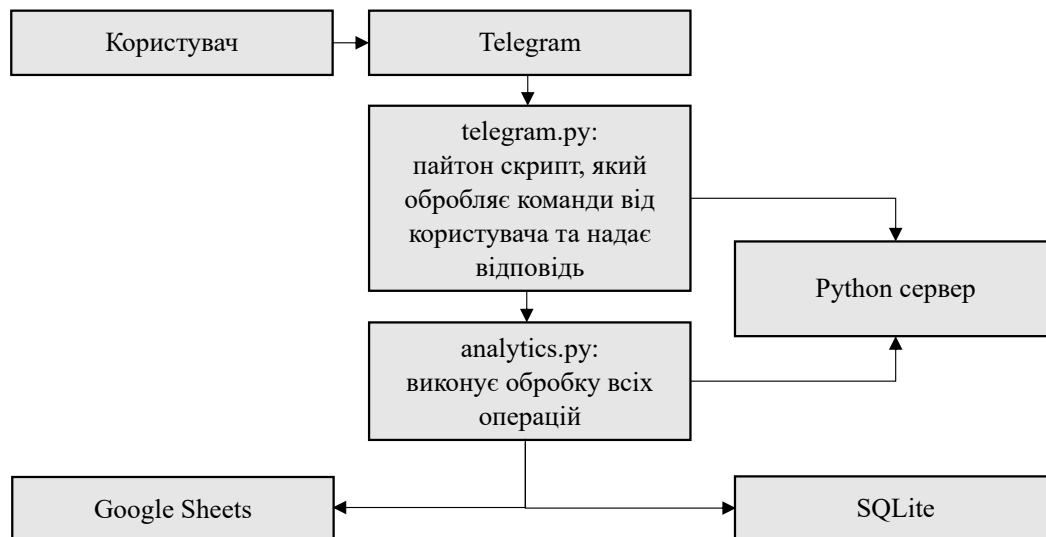


Рис. 4. Архітектурна схема модульної системи Telegram-бота

У якості системи зберігання даних використано локальну реляційну СУБД SQLite та хмарний сервіс Google Sheets, докладні характеристики відповідних компонентів наведено в (табл. 4).

Таблиця 4

Основні компоненти системи зберігання даних Telegram-бота

Компонент	Структура та поля	Призначення
SQLite	Таблиця users (id, username, role); operations (id, line_type, line_id, json_data, timestamp); referrals (id, referrer_name, line_id, referral_type, referral_value, timestamp); sqlite_sequence (name, seq)	Зберігання облікових записів і ролей користувачів, журнал транзакцій у вигляді JSON, інформація про рефералів, управління автоінкрементом
Google Sheets	Аркуш "Transactions" (date, balance_cash_eur, balance_cash_usd, balance_cash_pln, balance_bn_eur, balance_bn_pln, balance_usdt, exchange_rate_eur, exchange_rate_usd, exchange_rate_pln, debt1...debt7, total_referral, operation_type, comment, referral_name, referral_type, referral_value, line_id)	Хмарне резервне копіювання транзакцій, поточних залишків і боргів; зовнішня аналітика; колективний доступ до даних

На наступному етапі розгорнуто користувацький інтерфейс Telegram-бота, який у режимі реально-го часу обробляє команди й повідомлення, відображає головне меню у форматі кнопкової клавіатури та керує багатокроковими сценаріями введення даних. Через обробники callback-подій бот послідовно запитує тип операції, валюту, категорію боргу та суму, після чого передає зібрані параметри до обчислювального ядра. Сукупність реалізованих функцій наведено в (табл. 5).

Таблиця 5

Функціональні можливості Telegram-бота

Функція / Модуль	Тип взаємодії з користувачем	Кодові функції / методи	Опис
Нова операція	Кнопкова навігація	handle_new_operation() → process_operation(data)	Збір даних транзакції й передача на обчислення
Витрата	Текстове введення	handle_expense() → new_expense(data)	Реєстрація витрати, оновлення балансів і збереження
Показати залишки	Кнопка "Show Balances"	show_balances(message)	Формування звіту про поточні залишки
Останні операції	Кнопка "Recent Operations" + Inline-кнопки	handle_recent_operations() → show_transactions() → delete_operation()	Перегляд історії транзакцій і видалення записів
Борги	Кнопка "Debts" + послідовні запити	handle_debts() → show_debts() → add_debt() → rest_debt()	Додавання, перегляд і погашення заборгованостей
Адміністратори	Кнопка "Admin" + підменю	handle_admin() → admin_income() → admin_operations_count()	Показ прибутку й кількості угод за обраний період
Модуль операцій	–	process_operation(data)	Валідація, розрахунок результату й формування структури для збереження
Модуль витрат	–	new_expense(data)	Розрахунок впливу витрати, запис у БД та Google Sheets
Модуль боргів	–	show_debts() → add_debt(data) → rest_debt(data)	Логіка додавання і погашення боргів
Модуль рефералів	–	register_referral() → add_referral_entry(data)	Реєстрація рефералів і розрахунок комісій
Модуль аналітики	–	count_income(period) → count_operations(period)	Обчислення прибутку й статистики операцій
Ініціалізація БД	–	initialize_db()	Створення і підготовка таблиць у SQLite
Синхронізація	–	load_data_from_google_sheets() → append_row() → delete_rows() → batch_format()	Обмін даними між Google Sheets і локальною БД

Доступ до адміністративної панелі обмежено жорстко визначеним переліком Telegram-юзернеймів, що унеможливорює виконання керівних дій неавторизованими особами. Оповіщення про критичні події (створення операції, погашення боргу тощо) автоматично ретранслюються у задану службову групу Telegram, забезпечуючи оперативне інформування відповідальних співробітників.

Після реалізації архітектури, серверної логіки й клієнтського інтерфейсу Telegram-бота проведено валідаційне тестування, спрямоване на перевірку працездатності системи в умовах реальної експлуатації, виявлення логічних помилок і оцінювання її надійності та продуктивності. Коректність, стабільність і пропускну здатність програмного комплексу оцінено за результатами, наведено у в (табл. 6).

Таблиця 6

Результати апробації функціональної та продуктивної надійності системи

Категорія	Обсяг перевірки	Підсумок
Кросплатформність	UI та підключення на Windows/Android/iOS	Успішно
Функціональність	Команди: операції, витрати, борги, реферали, статистика	Успішно
Інтеграція модулів	Узгодженість telegram.py ↔ analytics.py ↔ SQLite ↔ GSheets	Успішно
Доступ	Рольова авторизація; блокування фальшивих адмін-запитів	Успішно
Нотифікації	Автоалерти про операції, борги, виплати в службовий чат	Успішно

Висновки. Таким чином, розроблений Telegram-бот повністю реалізує мету дослідження: автоматизує збір, обробку й зберігання даних про транзакції, витрати, борги та реферальні виплати криптообмінника. Застосування telebot, SQLite і інтеграції з Google Sheets забезпечило безшовну взаємодію з користувачем без ручного втручання в структуру даних. Модульна архітектура (операції, витрати, борги, аналітика, звітність) надала системі масштабованості й гнучкості. Отже, Telegram-бот є ефективним інструментом автоматизації фінансового обліку невеликих криптообмінників, зменшуючи трудові витрати, підвищуючи точність і оптимізуючи бізнес-процеси.

Список використаних джерел:

1. Про віртуальні активи. Закон України, № 2074-IX від 17.02.2022. URL: <https://zakon.rada.gov.ua/go/2074-20>. (дата звернення: 25.04.2025).
2. Bang J., Choi M.-J. Design of Real-Time Transaction Monitoring System for Blockchain Abnormality Detection. *Lecture Notes in Electrical Engineering*. 2019. DOI: 10.1007/978-3-030-34365-1_25.
3. Basri N. A., Ahmad N. A., Krishnan V., Razeman H., Ibrahim S. Z., Sultan J. M. Electronic Medicine Organizer Box with Telegram Bot. *AIP Conf. Proc.* 2023. 2579 (1): 020028. DOI: <https://doi.org/10.1063/5.0112804>.
4. Deloitte launches innovative 'DARTbot' internal chatbot. URL: <https://www2.deloitte.com/us/en/pages/about-deloitte/articles/press-releases/deloitte-launches-innovative-dartbot-internal-chatbot.html>
5. IBM. Enterprise Resource Planning (ERP) Advantages & Disadvantages. 2023. URL: <https://www.ibm.com/think/insights/enterprise-resource-planning-advantages-disadvantages> (дата звернення: 25.04.2025).
6. Louro B., Abreu R., Cabral J., Sequeiros J. B. F., Inácio P. R. M. Analysis of the Capability and Training of Chat Bots in the Generation of Rules for Firewall or Intrusion Detection Systems. *ARES '24: Proceedings of the 19th International Conference on Availability, Reliability and Security*. Article No.: 123, pp. 1–7. DOI: <https://doi.org/10.1145/3664476.367090>.
7. NetSuite. Pros and Cons of Using an ERP. *NetSuite*. 2020. URL: <https://www.netsuite.com/portal/resource/articles/erp/erp-pros-cons.shtml> (дата звернення: 25.04.2025).
8. Parlika R., Atmaja P. W. Realtime Monitoring of Bitcoin Prices on Several Cryptocurrency Markets using Web API, Telegram Bot, MySQL Database, and PHP-Cronjob. *Proc. of the 2020 6th Information Technology International Seminar (ITIS)*. 2020.
9. Pavelescu (Belbe) A. E. Contributions to the Improvement of the Process of Internal Audit at Public Educational Entities by Using Chatbots. *Ovidius University Annals: Economic Sciences Series*. 2023. XXIII (2). pp. 833–840.
10. Rasure E. Countries Where Bitcoin Is Legal and Illegal. *Investopedia*. 2024. URL: <https://www.investopedia.com/articles/forex/041515/countries-where-bitcoin-legal-illegal.asp> (дата звернення: 25.04.2025).
11. Sabry F., Labda W., Erbad A., Malluhi Q. M. Cryptocurrencies and Artificial Intelligence: Challenges and Opportunities. *IEEE Access*. 2020. Vol. 8. P. 175840–175858. DOI: 10.1109/ACCESS.2020.3025211.
12. Suárez-Díaz J. L., Rodríguez-Barroso N., Gómez-Trenado G. Interactive Telegram Bot as a Support Tool for Teaching in Higher Education. *Lecture Notes in Networks and Systems*. 2024. Vol. 957. pp. 352–361. DOI: https://doi.org/10.1007/978-3-031-75016-8_33.
13. The Advantages of Google Sheets vs. Excel. 2023. URL: <https://www.toptal.com/management-consultants/excel-experts/google-sheets-advantages> (дата звернення: 25.04.2025).

Дата надходження статті: 05.06.2025

Дата прийняття статті: 13.06.2025

Опубліковано: 23.09.2025

УДК 004.42:519.85

DOI <https://doi.org/10.32689/maup.it.2025.2.15>

Юрій ЛУК'ЯНЧУК

кандидат технічних наук, доцент,

доцент кафедри комп'ютерних наук,

Луцький національний технічний університет, iuriilukianchuk87@gmail.com

ORCID: 0000-0001-9690-6197

Лев БОЙКО

кандидат технічних наук,

доцент кафедри інженерії програмного забезпечення,

Луцький національний технічний університет, lewlazarus@gmail.com

ORCID: 0009-0001-0117-8551

ІННОВАЦІЙНЕ ПРОГРАМНЕ ЗАБЕЗПЕЧЕННЯ ДЛЯ ФОРМУВАННЯ МОРАЛЬНО-ЕТИЧНИХ ЦІННОСТЕЙ ПІДЛІТКІВ У ЦИФРОВУ ЕПОХУ

Анотація. У сучасному суспільстві, що перебуває під сильним впливом цифрових технологій, особливої актуальності набуває питання морального й етичного виховання підлітків. Цифрове середовище, зокрема соціальні мережі, часто стає джерелом кібербулінгу, фейкових новин, фішингу, що негативно впливає на психоемоційний стан молоді.

Метою цього дослідження є розробка програмного забезпечення, що сприятиме формуванню моральних орієнтирів у підлітків віком 12–17 років, з акцентом на розвиток критичного мислення, навичок цифрової безпеки та толерантності.

Методологічно дослідження ґрунтується на поєднанні теоретичного аналізу наукової літератури та емпіричних методів, статистичних даних про рівень кіберзагроз та булінгу серед молоді. Також застосовано метод проектування програмного продукту та його експериментальне впровадження в освітньому середовищі.

Наукова новизна полягає у використанні гейміфікації та адаптивного навчання у форматі мобільного застосунку, інтерфейс якого наближений до соціальних мереж. Це дозволяє підліткам легше сприймати навчальний матеріал і застосовувати його в реальному житті.

Результати дослідження показали ефективність запропонованого інструменту для виховання морально стійких і соціально відповідальних молодих людей. Програмний продукт отримав позитивні відгуки з боку освітніх установ, що свідчить про доцільність його подальшого масштабування.

Ключові слова: моральне виховання, цифрові технології, підлітки, кібергігієна, гейміфікація, критичне мислення, програмне забезпечення.

Iurii LUKIANCHUK, Lev BOIKO. INNOVATIVE SOFTWARE FOR DEVELOPING MORAL AND ETHICAL VALUES AMONG TEENAGERS IN THE DIGITAL AGE

Abstract. In today's society, which is deeply influenced by digital technologies, the issue of moral and ethical education of adolescents is gaining critical importance. The digital environment, particularly social media, often becomes a source of cyberbullying, fake news, and phishing, negatively impacting the emotional and psychological well-being of young people.

The purpose of this study is to develop software aimed at fostering moral values in teenagers aged 12–17, with a focus on enhancing critical thinking, digital security skills, and tolerance.

The methodology is based on a combination of theoretical analysis of relevant academic literature and empirical methods, including statistical data on the prevalence of cyber threats and bullying among adolescents. The study also employs design methods for software development and experimental testing within educational environments.

The scientific novelty of the project lies in the application of gamification and adaptive learning techniques within a mobile application whose interface mimics familiar social media platforms. This approach makes the learning content more accessible and engaging for the target audience.

The results demonstrate the effectiveness of the proposed digital tool in fostering moral resilience and social responsibility among youth. The application has received positive feedback from educational institutions, which confirms the feasibility of its further development and wider implementation.

Key words: moral education, digital technologies, teenagers, cyber hygiene, gamification, critical thinking, software.

Постановка проблеми. Морально-етичне виховання підлітків завжди було однією з найважливіших складових освіти. Воно формує не лише особистісні якості, а й визначає, як молодь взаємодіє з навколишнім світом і будує стосунки з оточуючими. У сучасному світі, де технології стали невід'ємною частиною життя, якісна освіта набуває нових вимірів. Особливо це стосується таких актуальних

© Ю. Лук'янчук, Л. Бойко, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

проблем, як кібергігієна та булінг, які дедалі більше впливають на психічне здоров'я та соціальну адаптацію підлітків.

Україна, як і багато інших країн, стикається з викликами, пов'язаними зі зростаючим впливом цифрового середовища на молодь. Підлітки проводять багато часу в Інтернеті, де на них може підстерігати небезпека: від кібербулінгу до впливу шкідливого контенту. Крім того, булінг у навчальних закладах залишається серйозною проблемою, яка торкається тисяч дітей, впливаючи на їхню самооцінку, успішність у навчанні та загальний психологічний стан.

Морально-етичне виховання у таких умовах має бути спрямоване не лише на формування традиційних цінностей, а й на навчання підлітків безпечній поведінці в мережі, повазі до оточуючих, а також вмінню протистояти агресії в реальному житті та в онлайн-середовищі. Важливо донести до дітей розуміння того, що їхні слова та вчинки в Інтернеті мають реальні наслідки, а повага до інших є основою здорового суспільства.

Аналіз останніх досліджень і публікацій. На сьогодні особливу увагу слід приділяти кібергігієні (етичній поведінці в мережі), толерантності в суспільстві та зниженню ризиків булінгу в навчальних закладах. Також важливо розвивати критичне мислення, щоб вміти відрізнити фейкові новини та контент, створений за допомогою штучного інтелекту.

Згідно з дослідженням Rakuten Viber [10], у 2024 році понад 62 % українців стикалися з кіберзлочинцями, що призвело до масових витоків інформації та персональних даних. Серед підлітків цей показник ще вищий, оскільки навчальні заклади недостатньо уваги приділяють етиці спілкування в соцмережах.

Також існують дані про булінг у навчальних закладах серед підлітків [2]. Його причинами можуть бути як наслідування поведінки в сім'ї, так і неадекватна реакція підлітка на аморальне ставлення до нього. У 2024 році поліція зафіксувала 219 повідомлень про булінг, переважно серед учнів 14–16 років (рис. 1).



Рис. 1. Кількість випадків булінгу в освіті

Джерело: сформовано авторами на підставі [2]

Особливий резонанс отримала інформація про катування 12-річної дівчинки у м.Біла Церква [5]. Попри оперативні дії правоохоронців, моральний стан дитини було підірвано назавжди, а винуватців притягнуто лише до адміністративної відповідальності (рис. 2).

Окремо варто відзначити, що складні життєві обставини в Україні впливають на психічний стан підлітків. Наприклад, ще до повномасштабного вторгнення, згідно зі звітом UNICEF Ukraine [1], 13 % дітей та підлітків 10–19 років мали психічні розлади. Сьогодні цей показник зріс у кілька разів.

Проблема морального й етичного виховання підлітків у цифрову епоху набула актуальності лише в останні роки, коли спостерігається стрімке зростання використання цифрових технологій не лише в освіті, а й у повсякденному житті молоді. У класичних педагогічних концепціях моральне виховання розглядалося як складова формування особистості через авторитети, культуру, традиції, проте сьогодні роль цих механізмів суттєво знижена. Замість родини та школи джерелами формування цінностей для молоді дедалі частіше стають соціальні мережі, відеоплатформи, ігрові середовища.

Аналіз наукових джерел підтверджує високий рівень занепокоєння дослідників щодо впливу цифрового середовища на особистість підлітка. О. Пометун [6; 7] підкреслює, що основою моральної саморегуляції у сучасному суспільстві має бути критичне мислення, здатність до етичного аналізу інформації та усвідомлення наслідків онлайн-поведінки. Також повідомляється, що культура діалогу, довіра і толерантність є ключовими для формування цифрової етики у школі, однак їх слід підсилювати цифровими інструментами.



Рис. 2. Приклад підліткового булінгу

Джерело: сформовано авторами на підставі [5]

В дослідженнях С. Горбачова [8] розглядається кібербулінг як морально-психологічна загроза, що вимагає системної інтеграції цифрової безпеки у навчальний процес, починаючи з початкових класів.

Міжнародні дослідження підтверджують аналогічні тенденції. У звіті London School of Economics під керівництвом Sonia Livingstone [14] вказано, що підлітки, які проводять понад 4 години на день у мережі без контролю, мають вищий рівень агресії, знижену здатність до емпатії й слабшу ідентифікацію моральних норм. У цьому ж дослідженні підкреслюється ефективність гейміфікованих освітніх платформ у процесі морального самозростання.

Щодо вітчизняного контексту, окремі ініціативи, наприклад цифрові курси з кібергігієни на платформі «Освіторія» [12] або освітні матеріали про інформаційну безпеку на «Дія.Освіта» [4], є важливими кроками, однак вони зосереджені здебільшого на технічній стороні захисту, а не на моральному чи етичному контексті.

Дослідження О. Спіріна [9] аналізують ефективність гейміфікації в освіті, звертаючи увагу на її потенціал у розвитку навичок і мотивації до навчання. Проте проблема ціннісного наповнення таких платформ лишається відкритою.

Таким чином, невирішеними залишаються:

- брак комплексних програм, що поєднують моральне виховання, психологічну підтримку й цифрову компетентність;
- відсутність інструментів, адаптованих до національного й культурного контексту України;
- недостатній акцент на етичному аспекті онлайн-взаємодії у формальних освітніх програмах;
- обмежене практичне впровадження ігрових механік у виховання етичних орієнтирів.

Метою статті є обґрунтування доцільності та ефективності використання інноваційного програмного забезпечення для формування морально-етичних цінностей у підлітків в умовах цифрового середовища. Автори прагнуть довести, що застосування цифрових інструментів із елементами гейміфікації, критичного мислення та кібергігієни може значно підвищити рівень усвідомлення етичних норм серед молоді, зробити виховний процес сучасним, цікавим та інтерактивним.

Для реалізації цієї мети було поставлено такі завдання:

- проаналізувати основні ризики цифрового середовища для підлітків з моральної точки зору;
- дослідити педагогічні підходи до формування етичної свідомості в умовах онлайн-комунікації;
- розробити та апробувати програмне забезпечення, орієнтоване на формування моральних установок через гейміфікацію;
- оцінити його вплив на поведінкові й ціннісні орієнтації учнів середнього шкільного віку.

Виклад основного матеріалу дослідження. Об'єктом цього дослідження є морально-етичне виховання підлітків в Україні в умовах сучасної цифрової середовища. Це складний і багатогранний процес, що включає формування моральних норм, соціальних цінностей, навичок взаємодії з іншими людьми, а також здатності приймати відповідальні рішення. Особлива увага приділяється тому, як морально-етичне виховання впливає на поведінку підлітків у реальному та віртуальному світі, їхню здатність протистояти негативним явищам, зокрема булінгу та інформаційним загрозам.

Предметом дослідження є використання програмного забезпечення як засобу морально-етичного виховання підлітків. У роботі розглядається, як спеціалізовані цифрові технології можуть сприяти розвитку моральних цінностей, формуванню толерантності, навчанню кібергігієни та критичному

мисленню. Зокрема, досліджується ефективність інтерактивних освітніх платформ, що використовують гейміфікацію та адаптивні методи навчання у формуванні відповідального ставлення підлітків до власної поведінки як у соціальних мережах, так і в повсякденному житті.

Дослідження використовує комплексний підхід, що поєднує аналіз теоретичних основ морально-етичного виховання, оцінку соціальних викликів та практичну реалізацію сучасних цифрових технологій в освітньому процесі.

Держава активно сприяє морально-етичному вихованню школярів. Відповідно, було запропоновано низку правил для формування моральної культури учнів [11]. Однак на практиці надзвичайно важко донести інформацію до всіх безпосередніх учасників освітнього процесу та вимагати від них раціонального дотримання толерантності. Крім того, через різні обставини вчителі не можуть повноцінно проводити виховну роботу та пояснювати всі ризики аморальної поведінки. Саме тому було вирішено розробити спеціалізоване програмне забезпечення для виховання морально-етичних принципів серед молоді в ігровій формі. Основний акцент зроблено на кібергігієні та безпечній поведінці в мережі. Важливо також навчати толерантності до інших людей. У цьому випадку йдеться не лише про інших підлітків, а й про людей з інклюзивністю, оскільки їх щодня стає все більше в повсякденному житті. Також виховується та пропагується повага до старшого покоління, оскільки вони сьогодні є найменш захищеними. Окремо розглядається розвиток критичного мислення, яке стане основою для подальшого розширення пізнавальних здібностей.

За основу було взято фреймворк Flutter [13]. Це програмне забезпечення з відкритим вихідним кодом для створення програм для платформ Android та iOS, а також для веб-додатків, розроблене Google.

Дизайн розроблений у формі популярних платформ соціальних медіа, оскільки він найкраще сприймається користувачами віком 12–17 років. Освітня платформа також розроблена з використанням методології гейміфікації. Це дозволяє утримувати увагу користувача та зробити весь процес використання максимально зручним. Крім того, гейміфікація освітнього процесу показала відмінні результати при вивченні нової теми [3]. Гейміфікація в освіті – це використання ігрових елементів та прийомів в освітньому процесі для створення інтерактивного та мотивуючого середовища. Вона використовується для підвищення залученості та інтересу студентів, стимулювання їх до активної участі в навчанні та формування позитивного ставлення до освітнього процесу. Дослідження впливу впровадження ігрових механік на освітній процес показали, що цей підхід подвоює швидкість засвоєння нових навичок студентами. Цей метод навчання не є повноцінною освітньою грою як такою. Використання гейміфікації в освіті передбачає застосування окремих елементів, які роблять процес навчання більш захопливим та мотивують учасників.

Дизайн логотипу для додатка був розроблений відповідно до стандартів та практик фахівців (рис. 3). Він має приємну кольорову гаму та зрозумілий інтерфейс. Назва додатка також була обрана за рекомендаціями фахівців і була визначена як «Cyber Scout». Це дає розуміння, що користувач досліджуватиме та аналізуватиме нові теми, пов'язані з мережевою безпекою. Наразі додаток реалізовано українською мовою. У майбутньому, після отримання позитивних відгуків та доопрацювання програмного забезпечення, його буде перекладено англійською.



Рис. 3. Логотип програмного застосунку «Cyber Scout»

Джерело: розроблено авторами

Логіка роботи програмного застосунку полягає в тестуванні користувача простими питаннями та виборі відповіді лише з двох варіантів: так чи ні (рис. 4). Після проходження всіх питань користувач

отримує інформацію про кількість правильних відповідей (рис. 5). Всі питання відсортовані за категоріями, що передбачають опанування знань з тем (рис. 6):

- фішинг та шахрайство;
- захист персональних даних;
- виявлення бот-ферм;
- як захистити свої дані під час війни;
- як розпізнати фейки;
- розпізнавання зображень, згенерованих штучним інтелектом.

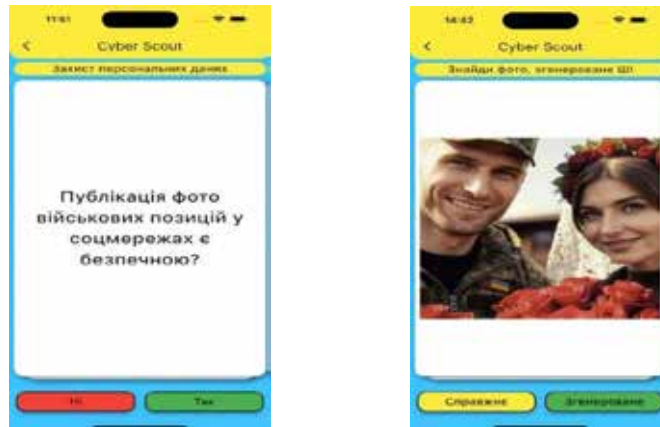


Рис. 4. Приклад тестових карток

Джерело: розроблено авторами

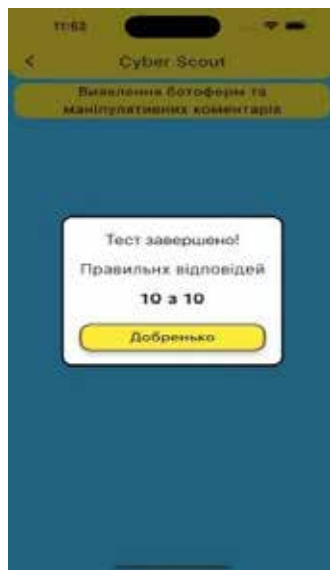


Рис. 5. Приклад результатів випробувань

Джерело: розроблено авторами



Рис. 6. Зображення початкового екрана з категоріями опитування

Зручний інтерфейс пропонує варіанти отримання відповіді натисканням кнопки або проведенням пальцем, щоб перегорнути картку питання у відповідному напрямку (рис. 7).

Наразі програмне забезпечення, що розробляється, все ще перебуває в процесі доопрацювання та наповнення контентом. Однак перші позитивні відгуки про важливість та необхідність такого застосування вже отримані від навчальних закладів та громадських організацій. Тому тестування відбудеться найближчим часом.

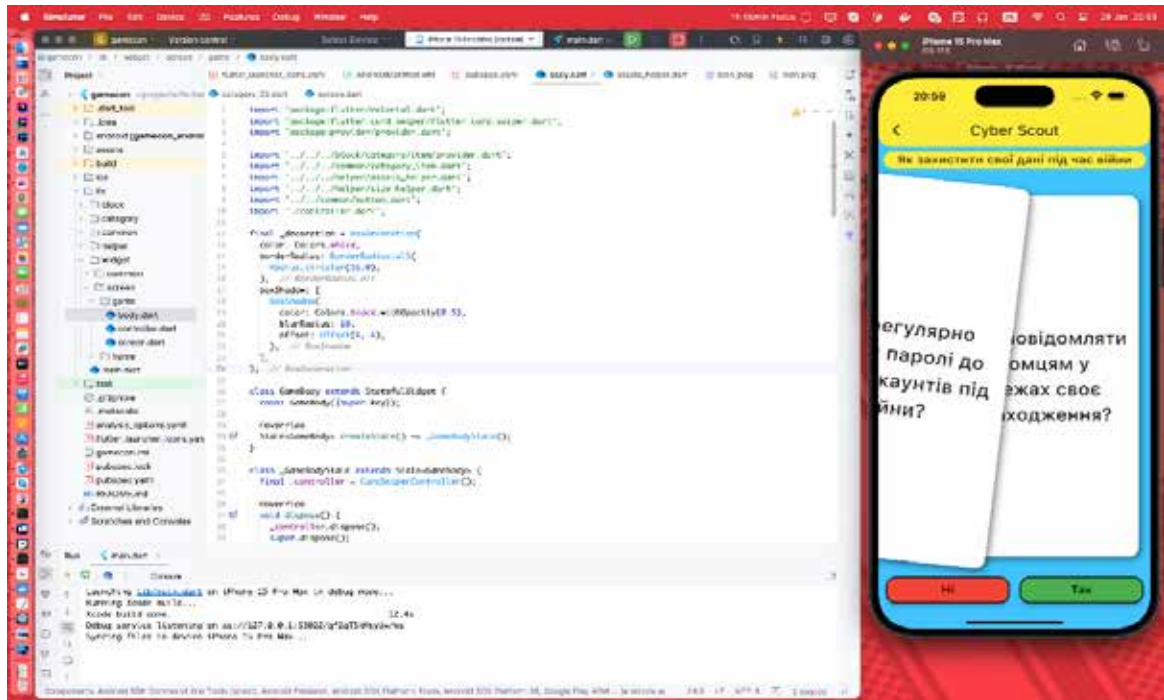


Рис. 7. Програмний код та зображення інтерфейсу застосунку

Джерело: розроблено авторами

Висновки. Розробка програмного забезпечення для морально-етичного виховання підлітків є важливим кроком до формування нового покоління, яке зможе ефективно функціонувати в умовах сучасного цифрового світу. Програма, яка поєднує елементи кібергігієни, толерантності, критичного мислення та гейміфікації, відповідає на сучасні виклики, з якими стикаються підлітки в Україні та за кордоном.

У майбутньому планується додавання нових модулів, які б охоплювали більше аспектів морально-етичного виховання, таких як екологічна свідомість, фінансова грамотність тощо. Також активно триває процес адаптації програми для інших мов, що дозволить розширити її аудиторію та використувати в інших країнах. Крім того, розглядається можливість інтеграції програми з існуючими освітніми платформами для полегшення її використання в школах та інших навчальних закладах.

Список використаних джерел:

- 13% усіх дітей та підлітків віком 10–19 років мають психічні розлади. URL: <https://nus.org.ua/2021/10/05/13-usih-ditej-ta-pidlitkiv-vikom-10-19-rokiv-mayut-psyhichni-rozlady-zvit-yunisef/> (дата звернення 29.05.2025).
- В Україні зафіксували понад 200 випадків булінгу серед підлітків у 2024 році. URL: <https://www.rbc.ua/rus/news/ukrayini-zafiksuvali-ponad-200-vipadkiv-bulingu-1734604436.html> (дата звернення 27.05.2025).
- Гейміфікація в освіті як засіб підвищення її ефективності. URL: <https://www.bestcleverslms.com/piznavalne/shcho-take-heyimifikatsiya-v-osviti/> (дата звернення 10.06.2025).
- Дія. Освіта. Онлайн-безпека. URL: <https://osvita.diia.gov.ua/catalog/topic/online-security> (дата звернення 06.06.2025).
- Побиття 12-річної дівчинки у Білій Церкві: поліція встановила вже 10 учасників інциденту. URL: <https://sensor.net/ua/news/3531561/politsiya-rozprovila-novi-detali-pobyttya-12-richnoyi-divchynky-u-biliiy-tserkvi> (дата звернення 29.05.2025).
- Пометун О. І. Критичне мислення як педагогічний феномен. *Український педагогічний журнал*. 2018. № 2. С. 89–98. URL: http://www.irbis-nbuv.gov.ua/cgi-bin/irbis_nbuv/cgiirbis_64.exe?I21DBN=LINK&P21DBN=UJRN&Z21ID=&S21REF=10&S21CNR=20&S21STN=1&S21FMT=ASP_meta&C21COM=S&2_S21P03=FILA=&2_S21STR=ukrpj_2018_2_14 (дата звернення 01.06.2025).
- Пометун О. І. Розвиток критичного мислення учнів засобами шкільного підручника історії. О. І. Пометун, Н. М. Гупан. *Проблеми сучасного підручника*. 2018. Вип. 20. С. 327–338. URL: http://www.irbis-nbuv.gov.ua/cgi-bin/irbis_nbuv/cgiirbis_64.exe?I21DBN=LINK&P21DBN=UJRN&Z21ID=&S21REF=10&S21CNR=20&S21STN=1&S21FMT=ASP_meta&C21COM=S&2_S21P03=FILA=&2_S21STR=psp_2018_20_32 (дата звернення 01.06.2025).
- Про кібербулінг: куди звертатися та як протидіяти. URL: <https://vseosvita.ua/c/news/post/15707> (дата звернення 03.06.2025).
- Спірін О. М. Інновації в освіті: нові підходи та технології. URL: <https://vseosvita.ua/library/naukovo-metodychna-stattia-innovatsii-v-osviti-novi-pidkhody-ta-tekhnohohii-958448.html> (дата звернення 08.06.2025).

10. У 2024 році понад 62% українців контактували з кіберзлочинцями. URL: https://dev.ua/news/kontaktuvaly-iz-kiberzlovmysnykamy?utm_source=lg&utm_medium=msg&utm_campaign=news_dev_ua_channel&utm_content=kiberzlovmysnykamy (дата звернення 27.05.2025).
11. Формування моральної культури учнів. URL: <https://vseosvita.ua/blogs/formuvannia-moralnoi-kultury-uchniv-vo-sukhomlynskyi-pro-moralne-vykhovannia-85197.html> (дата звернення 08.06.2025).
12. Як захистити дитину в інтернеті. URL: <https://osvitoria.media/experience/yak-zahystyty-dytynu-v-interneti/> (дата звернення 06.06.2025).
13. Flutter framework. URL: <https://uk.wikipedia.org/wiki/Flutter> (дата звернення 10.06.2025).
14. Livingstone S., Byrne J. Risks and opportunities of digital life for adolescents. URL: <https://www.lse.ac.uk/media-and-communications/research/research-projects/eu-kids-online> (дата звернення 03.06.2025).

Дата надходження статті: 17.06.2025

Дата прийняття статті: 23.06.2025

Опубліковано: 23.09.2025

УДК 004.056.55:004.738.5

DOI <https://doi.org/10.32689/maup.it.2025.2.16>

Ігор МАРТИНЮК

аспірант,

Державний торговельно-економічний університет,

i.p.martyniuk@gmail.com

ORCID: 0009-0005-2815-1104

Альона ДЕСЯТКО

доктор філософії в галузі «Комп'ютерні науки», доцент,

завідувач кафедри інженерії програмного забезпечення та кібербезпеки,

Державний торговельно-економічний університет, desyatko@gmail.com

ORCID: 0000-0002-2284-3418

ЗАГРОЗА АТАКИ 51 % ДЛЯ НИЗЬКОХЕШРЕЙТОВИХ ЦИФРОВИХ ВАЛЮТ

Анотація. Атаки 51 % являють собою серйозну загрозу для безпеки блокчейнів, особливо для цифрових валют (ЦВ) з низьким хешрейтом, що робить їх уразливими до маніпуляцій і втрати довіри з боку користувачів.

Мета. Дослідження загроз атаки 51 % у блокчейнах із низьким хешрейтом, зокрема для ЦВ і криптовалютних бірж (КВБ), з метою розробки ефективних стратегій протидії з використанням інтелектуальних інформаційних систем (ІІС) і теорії ігор.

Методологія. Проведений у межах статті огляд і аналіз попередніх досліджень показав, що ця загроза потребує комплексного підходу до захисту. У статті обґрунтовано застосування теорії ігор як інструменту аналізу стратегій обох сторін, що може значно поліпшити розуміння динаміки взаємодії між захисниками і зловмисниками, що в перспективі дасть змогу розробити більш ефективні проактивні заходи захисту ЦВ і КВБ, включно з використанням потенціалу ІІС.

Наукова новизна. ІІС можуть включати модуль оцінки стратегій сторін протистояння атакам, а також фінансових ресурсів сторін протистояння, що допоможе прогнозувати можливість реалізації атак 51 % або, навпаки, ефективно протидіяти їм. Показано, що жоден із методів захисту не є універсальним, тому комбінування різних підходів є критично необхідним для створення стійких рішень, здатних ефективно реагувати на нові загрози, як для ЦВ, так і для КВБ.

Висновки. Майбутні дослідження мають зосередитися на розробці та тестуванні нових ігрових моделей для обчислювального ядра ІІС, які інтегруватимуть наявні підходи та враховуватимуть еволюцію методів атак і захисту в криптовалютних системах, що дасть змогу не лише підвищити рівень безпеки, а й забезпечити довіру користувачів до цифрових валют, що, в свою чергу, є важливим для їхнього успішного впровадження та використання у фінансових системах.

Ключові слова: цифрові валюти, криптовалютні біржі, атака 51 %, низький хешрейт, теорії ігор, стратегія сторін.

Ihor MARTYNIUK, Alona DESIATKO. THREAT OF 51 % ATTACK FOR LOW-HASHRATE CRYPTOCURRENCIES

Abstract. 51 % attacks pose a serious threat to the security of blockchains, especially for digital currencies (DCs) with a low hash rate, which makes them vulnerable to manipulation and loss of trust by users.

Purpose. Research into the threats of 51 % attacks in blockchains with low hashrate, in particular for CCs and cryptocurrency exchanges (CCEs), with the aim of developing effective countermeasure strategies using intelligent information systems (IIS) and game theory.

Methodology. The review and analysis of previous research conducted within the article showed that this threat requires a comprehensive approach to protection. The article substantiates the use of game theory as a tool for analyzing the strategies of both parties, which can significantly improve the understanding of the dynamics of interaction between defenders and attackers, which in the future will allow developing more effective proactive measures to protect DCs and cryptocurrency exchanges (CEs), including using the potential of intelligent information systems (IIS).

Scientific novelty. IIS may include a module for assessing the strategies of the parties to resist attacks, as well as the financial resources of the parties to the confrontation, which will help predict the possibility of implementing 51 % attacks or, conversely, effectively counteract them. It is shown that none of the protection methods is universal, therefore, combining different approaches is critically necessary to create sustainable solutions that can effectively respond to new threats, both for DCs and CEs.

Conclusion. Future research should focus on the development and testing of new game models for the computational core of the IIS, which will integrate existing approaches and take into account the evolution of attack and protection methods in cryptocurrency systems, which will not only increase the level of security, but also ensure user trust in digital currencies, which, in turn, is important for their successful implementation and use in financial systems.

Key words: digital currencies, cryptocurrency exchanges, 51 % attack, low hash rate, game theory, strategy of opposing parties.

© І. Мартинюк, А. Десятко, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

Вступ. Останніми роками на тлі популярності Bitcoin, Binance Coin (BNB), тощо, спостерігається значне зростання інтересу й до інших цифрових валют (ЦВ) та криптовалютних систем, що зумовлено їхнім потенціалом у галузі децентралізованих фінансів. Однак разом із цим зростанням виникає низка загроз, здатних підірвати безпеку та стабільність криптовалютних систем. Однією з актуальних є загроза атаки 51 %, яка є ризиком, пов'язаним із концентрацією обчислювальної потужності в руках одного або декількох учасників мережі, а в умовах низького хешрейту ця загроза набуває особливої значущості, оскільки мінімальні ресурси, необхідні для здійснення подібної атаки, стають доступними для значної кількості зловмисників. Зауважимо, що як було показано в низці досліджень, зокрема, в [1–3], ЦВ з низьким хешрейтом являють собою легшу мішень для атак 51 %. У низці джерел до таких валют можна, наприклад, віднести Chia (XCH), Cardano (ADA), Algorand (ALGO), PancakeSwap (CAKE) [2–4].

Як було показано в [2–4] низький хешрейт може бути небезпечним для ЦВ, оскільки він робить мережу більш вразливою до атак, зокрема атаки 51 %, оскільки зловмисник може отримати контроль над більш ніж 50 % хешрейту і почати маніпулювати блокчейном, наприклад, проводити подвійні витрати. А крім того, низький хешрейт може призвести до збільшення часу підтвердження транзакцій, що ускладнить використання ЦВ для платежів, оскільки користувачі можуть чекати довше, ніж зазвичай, для підтвердження своїх транзакцій, а якщо хешрейт падає, це також може спричинити занепокоєння серед користувачів та інвесторів, і призведе до зниження інтересу до ЦВ і, як наслідок, до падіння її вартості.

Зауважимо, що багато країн розглядають можливість запровадження власних ЦВ, відомих як центральні банки цифрових валют (CBDC) [5; 6]. І ці нові ЦВ можуть зіткнутися з аналогічними загрозами, включно з ризиком атаки 51 %, оскільки на початкових етапах впровадження CBDC можуть мати відносно низький хешрейт, особливо якщо їхні мережі будуть не повністю завантажені або використані обмеженим числом учасників. Багато CBDC розробляються з акцентом на централізоване управління, що може знизити рівень децентралізації і зробити систему більш вразливою для атак, а крім того, CBDC можуть використовувати нові технології, які ще не пройшли достатнього тестування на безпеку, що може породити уразливість перед атаками. Отже, зловмисники можуть побачити в атаках на CBDC можливість отримати економічну вигоду, особливо якщо такі ЦВ будуть пов'язані з великими фінансовими системами.

Таким чином, на підставі вищесказаного, можна констатувати, що актуальність досліджень у цій галузі зумовлена не тільки зростаючими інвестиціями в ЦВ, а й необхідністю розроблення ефективних механізмів захисту від атак, здатних завдати шкоди як користувачам, так і інфраструктурі криптовалютних бірж (КВБ). Сучасні інтелектуальні інформаційні системи (ІС), засновані на алгоритмах машинного навчання (МН) і аналітики даних, можуть істотно підвищити рівень захисту ЦВ, даючи змогу не тільки виявляти аномалії в поведінці мережі, а й передбачати можливі атаки, зокрема атаку 51 %, що відкриває нові горизонти для гарантування безпеки у сфері блокчейн-технологій.

У цій статті ми проведемо аналіз попередніх досліджень і наукових публікацій щодо загроз, пов'язаних із реалізацією атаки 51 % на ЦВ із низьким хешрейтом, і розглянемо можливості використання сучасних ІС для розроблення проактивних заходів протидії цим загрозам.

Постановка проблеми. Актуальність дослідження загроз, пов'язаних з атакою 51 % на ЦВ з низьким хешрейтом, зумовлена збільшенням інтересу до ЦВ та їхньою інтеграцією у фінансові системи. Однак деякі ЦВ, володіючи обмеженими ресурсами, стають уразливими для атак, які можуть призвести до значних втрат як для користувачів, так і для інфраструктури. Незважаючи на наявність низки теоретичних досліджень, багато аспектів безпеки ЦВ залишаються недостатньо вивченими, а наявні підходи до захисту ЦВ і КВБ часто не враховують специфічних ризиків, пов'язаних із низьким хешрейтом, що, відповідно, породжує потребу в глибшому аналізі, який дасть змогу виявити вразливості та визначити ефективні заходи протидії. І в зв'язку з цим, розвиток сучасних ІС може надати нові можливості для аналізу загроз, що посприє створенню більш стійких і безпечних ЦВ. І важливим кроком на шляху розв'язання зазначеної проблеми є систематизація наявних знань та виявлення напрямів для подальших досліджень у цій галузі.

Мета дослідження. Метою цього аналітичного дослідження є систематизація та критичний аналіз наявних наукових публікацій, присвячених загрозам атаки 51 % на ЦВ з низьким хешрейтом, а також оцінювання потенціалу сучасних ІС для розроблення проактивних заходів протидії подібним загрозам. Дослідження спрямоване на виявлення ключових вразливостей, що існують у межах цієї проблематики, і визначення напрямів для розв'язання низки конкретних завдань у сфері забезпечення безпеки ЦВ.

Методика дослідження. У роботі застосовано аналітичну методику, що включає огляд і систематизацію наявних наукових публікацій, статей, дисертацій та інших досліджень, присвячених загрозам

атаки 51 % для ЦВ. Основними етапами методики є систематичний пошук актуальних наукових публікацій у спеціалізованих базах даних і наукових журналах, присвячених проблематиці безпеки ЦВ, а також критичний аналіз обраних публікацій з метою виявлення основних тем, підходів і результатів, що стосуються загроз атаки 51 %, а також оцінки застосовуваних методів захисту. Крім того, виконано порівняння різних підходів і рішень, представлених у літературі, для виявлення найефективніших стратегій протидії загрозам реалізації атаки 51 % і виявлено недостатньо вивчені аспекти та напрямки для подальших досліджень, що може сприяти поліпшенню безпеки ЦВ.

Результати дослідження. У світлі вищевикладеного, завдання протидії атакам 51 % на ЦВ з низьким хешрейтом є відносно новим і перебуває на стадії активного дослідження. Останніми роками спостерігається зростання кількості наукових публікацій, присвячених цій проблематиці, що свідчить про зростаючий інтерес до питань безпеки ЦВ і необхідних заходів захисту. Таким чином, подальший аналіз наявних публікацій дасть змогу виявити ключові підходи та напрямки для подальших досліджень у цій галузі.

Так, у [1] розглядається загроза «51 % атаки» на блокчейни, коли одна група майнерів контролює більше половини обчислювальної потужності мережі, при цьому автори проводять дослідження поведінки майнерів, щоб зрозуміти, як такі атаки можуть бути здійснені та які чинники впливають на їхню ймовірність.

У [2] автори розглядають безпеку блокчейнів з фокусом на виявлення атак 51 %. У роботі подано аналіз вразливостей блокчейн-систем і обговорено методи виявлення та запобігання таким атакам. Автори наводять приклади, що ілюструють наслідки атак 51 %, пропонуючи підходи для поліпшення захисту блокчейнів від подібних загроз і акцентуючи увагу на важливості моніторингу та оцінки безпеки блокчейн-мереж.

У [3] розглядається проблема атаки 51 % для мережі Біткойн. Автори надають огляд цієї загрози та описують механізми, які можуть бути використані для її здійснення, а також наслідки, які така атака може мати для мережі та користувачів. Також авторами обговорюються методи захисту та поліпшення безпеки Біткойн-мережі, щоб знизити ризики, пов'язані з цією атакою.

У дисертації [4] досліджується вплив різних атак на ринки ЦВ. Дисертант аналізує, як різні види атак, як-от зломи КВБ і атаки 51 %, впливають на ціни та волатильність ЦВ. У роботі розглядаються також механізми реагування ринків на такі загрози та їхні довгострокові наслідки для інвесторів та екосистеми ЦВ загалом.

У [7] розглядаються особливості атаки 51 %, включно з механізмами її реалізації та потенційними наслідками для блокчейн-систем, при цьому автор аналізує різні аспекти вразливості блокчейнів до цієї атаки, а також пропонує методи для захисту від неї.

У [8] аналізується метод запобігання атаці 51 % на Біткойн з використанням стохастичного аналізу двофазного механізму proof-of-work. Автор аналізує, як цей підхід може поліпшити безпеку мережі, знижуючи ймовірність того, що одна група майнерів зможе контролювати більш ніж половину обчислювальної потужності, а також обговорюються переваги та недоліки запропонованого методу, а також його потенційний вплив на стійкість блокчейна до подібних атак.

У [9] досліджуються сприйняття атак 51 % на ЦВ з алгоритмом proof-of-work у соціальних медіа. Автори застосовують методи обробки природної мови (NLP) для аналізу обговорень і думок користувачів у соціальних мережах, щоб зрозуміти, як ці атаки впливають на суспільне сприйняття і довіру до ЦВ.

У [10] автори сфокусувалися на тому, як атаки, включно з атаками 51 %, впливають на прибутковість ЦВ. Автори аналізують ринкові дані та намагаються зрозуміти, які зміни в цінах відбуваються під час атак, а також виявляють закономірності та фактори, які можуть впливати на реакцію ринку у відповідь на кіберзагрози.

У [11] досліджуються різні аспекти проектування нових блокчейнів. Автори розглядають, як інноваційні підходи можуть бути застосовані для створення більш стійких і ефективних блокчейнів, а також як вони можуть допомогти у вирішенні існуючих проблем, таких як безпека і консенсус.

Робота [12] присвячена оцінці механізмів консенсусу і безпеки блокчейна в розрізі загрози реалізації атаки 51 %. Автори аналізують наявні методи захисту, їхню ефективність і вразливість, пропонуючи рекомендації щодо поліпшення безпеки блокчейнів для запобігання подібних атак.

У [13] обговорюється, що відбувається під час атак 51 % і як вони впливають на мережі блокчейна. Автор пояснює механізми атаки 51 %, її наслідки для користувачів і мережі загалом, а також у висновку сформульовано рекомендації поради щодо мінімізації ризиків, пов'язаних із подібними загрозами.

У [14] представлено інструмент під назвою BlockScore, який призначений для виявлення та дослідження вразливостей, що виникають унаслідок форків блокчейн-проектів. Автори аналізують, як уразливості можуть поширюватися і впливати на безпеку та стабільність форкнутих блокчейнів, пропонуючи методи для їх виявлення та усунення.

Робота [15] аналізує маніпуляції з ЦВ, відомі як «pump and dump». Автори досліджують методи виявлення та аналізу таких маніпуляцій на ринках ЦВ, пропонуючи підходи для захисту інвесторів і підвищення прозорості торгових практик.

У [16] розглядається вплив інтернет-спільноти на ЦВ на прикладі порівняння Dogecoin і Litecoin у Twitter. Автори аналізують, як обговорення і думки користувачів у соціальних мережах можуть впливати на сприйняття і цінність цих ЦВ, досліджуючи динаміку взаємодії інтернет-спільноти з цими проектами.

У [17] представлено ефективний підхід до захисту блокчейнів з алгоритмом proof-of-work від атак 51 % з використанням інформації про минулі блоки. Автори пропонують схему, яка враховує історію майнінгу та обчислювальної потужності, щоб підвищити безпеку мережі та зробити атаки менш імовірними.

Дисертація [18] присвячена розробці гібридної моделі консенсусу, що об'єднує механізми proof-of-work (PoW) і proof-of-stake (PoS) для захисту криптовалютної системи від атак 51 %. Автор досліджує, як така комбінація може підвищити стійкість мережі до атак, покращуючи безпеку і знижуючи ризики для учасників.

У дисертації [19] розглядаються атаки 51 % на блокчейни та пропонується архітектура рішення для забезпечення безпеки Інтернету речей (IoT) з використанням механізму proof-of-work. Автор досліджує, як можна адаптувати блокчейн-технології для захисту IoT-пристроїв від атак, забезпечуючи безпеку та цілісність даних.

Робота [20] присвячена гібридному підходу для мінімізації ризику атак 51 % на ЦВ. Автори пропонують методи та стратегії, які комбінують різні механізми консенсусу та безпеки, щоб підвищити стійкість криптовалютних систем до таких атак, обговорюючи їхню ефективність та застосування.

У [21] автор аналізує проблеми безпеки, з якими стикаються КВБ, а також розглядає фактори, що сприяють кіберзагрозам, наприклад, такі як уразливості в системах безпеки, соціальна інженерія тощо.

У [22] розглядаються різні ризики, з якими стикаються компанії, що працюють у сфері ЦВ, а також обговорюються фінансові, операційні, технологічні та правові ризики, характерні для криптовалютного бізнесу. Також аналізуються причини виникнення цих ризиків, включно з нестабільністю ринку, змінами в законодавстві та кіберзагрозами.

На підставі виконаного огляду попередніх досліджень сторонніх авторів можна стверджувати, що для ефективної протидії атакам 51 % важливо задіяти потенціал сучасних інтелектуальних інформаційних систем (ІС), оскільки такі ІС можуть посприяти розробленню проактивних заходів, спрямованих на виявлення та запобігання атакам на ранніх стадіях, аналізуючи великі обсяги даних про мережеву активність і моделюючи поведінку учасників. Використання ІС дасть змогу не тільки підвищити безпеку криптовалютних мереж, а й створити більш стійкі механізми консенсусу, що в кінцевому підсумку посилить довіру користувачів до блокчейн-технологій.

Крім того, важливо в таких ІС задіяти модулі оцінки фінансових ресурсів сторін протистояння, адже для проведення такої атаки зловмиснику знадобляться значні фінансові ресурси для реалізації своїх стратегій з контролю над ЦВ. І якщо ми згадали, що для реалізації атаки 51 % сторонам необхідно задіяти фінансові ресурси, то сторони, відповідно, можуть дотримуватися певних стратегій, причому обидві сторони можуть зазнати як фінансових втрат, так і отримати виграш, якщо обрана стратегія для сторони протистояння виявиться правильною.

Наприклад, зловмисники можуть вкласти значні кошти в обладнання для майнінгу та електроенергію, щоб збільшити свою частку хешрейту. Однак, така стратегія пов'язана з ризиком, оскільки високі витрати на обладнання та можливість невдачі в атаці можуть призвести до фінансових втрат.

Також обидві сторони можуть об'єднувати свої ресурси з іншими майнерами для збільшення шансів на успішну атаку. Такий варіант також пов'язаний з ризиками, оскільки можливий розпад альянсу і поділ прибутку, що може призвести до фінансових втрат для всіх учасників.

Слід згадати й те, що зловмисники (атакуюча сторона) може маніпулювати ринковими цінами ЦВ, створюючи дефіцит або розпускаючи чутки. Однак, ця стратегія теж ризикована, оскільки непередбачувана реакція ринку може призвести до значних збитків для зловмисника.

Атака з виведенням активів:

Сторона захисту щонайменше може інвестувати в дослідження вразливостей блокчейну для виявлення шляхів атаки, хоча ця стратегія і пов'язана з витратами, які можуть не окупитися.

Вважаємо, що для продовження дослідження для розгляду стратегій сторін доцільно задіяти потенціал теорії ігор, оскільки в даній постановці задачі можна виокремити двох основних гравців – захисника, який прагне захистити ЦВ від атаки 51 %, і зловмисника, який намагається реалізувати цю атаку.

Використовуючи теорії ігор, можна під час майбутніх досліджень розробити для обчислювального ядра ІІС математичні моделі, що описують стратегії обох сторін, їхні вибори та потенційні результати, а також оцінювати ризики та вигоди для кожної зі стратегій, аналізуючи, які дії призводять до оптимальних результатів для кожного гравця.

Однак, для ефективної протидії атакам 51 % на блокчейни важливо розглянути різні підходи, які можуть допомогти в розробці стратегій захисту. У (табл. 1) представлено порівняння підходу на основі теорії ігор з іншими методами, включаючи їхні переваги та недоліки.

Таблиця 1

Порівняння різних підходів, їхніх переваг і недоліків для розв'язання задачі аналізу протидії атакам 51 % для низькохешрейтових ЦВ

Підхід для розв'язання задачі	Переваги	Недоліки
Теорія ігор.	Дозволяє моделювати взаємодії двох гравців. Аналізує стратегії та їхні результати. Оцінка ризиків і вигод для обох сторін.	Вимагає складних математичних моделей. Обмежена передбачуваність поведінки гравців.
Аналіз вразливостей.	Дозволяє виявляти слабкі місця в системі. Допомогає розробити конкретні заходи захисту.	Може бути недостатньо актуальним. Не враховує поведінку зловмисників.
Прогностичні моделі (наприклад, машинне навчання).	Забезпечує передбачення на основі великих даних. Дозволяє виявляти аномалії в реальному часі.	Залежність від якості вхідних даних. Може не враховувати нові, невідомі загрози.
Фінансовий аналіз і ризик-менеджмент.	Оцінка витрат і потенційних збитків. Допомогає ухвалювати обґрунтовані рішення на основі фінансових даних.	Може ігнорувати технічні аспекти. Не завжди враховує динаміку загроз.

Джерело: складено авторами на підставі даних літературних джерел [1–10]

Виконаний у таблиці 1 аналіз різних підходів до розв'язання задачі протидії атакам 51 % демонструє, що кожен метод має свої унікальні переваги та обмеження, що додатково наголошує на необхідності комплексного підходу, який поєднує в собі елементи теорії ігор та інші стратегії, щоб забезпечити максимальний захист і стійкість ЦВ і КВБ.

Висновки. Проведений огляд і аналіз попередніх досліджень засвідчив, що атаки 51 % становлять серйозну загрозу для безпеки блокчейнів, особливо для ЦВ із низьким хешрейтом. Обґрунтовано, застосування теорії ігор як інструменту аналізу стратегій обох сторін, що допоможе значно поліпшити розуміння динаміки взаємодії між захисниками та зловмисниками, даючи змогу розробити більш проактивні заходи захисту, зокрема й на основі задіяння потенціалу інтелектуальних інформаційних систем (ІІС), що можуть містити модуль оцінювання стратегій сторін протистояння та їхніх фінансових ресурсів для реалізації атаки 51 % або протидії їй.

Показано, що жоден із методів не є універсальним, тому важливо комбінувати їх для створення стійких рішень, здатних ефективно реагувати на нові загрози для ЦВ і КВБ.

Майбутні дослідження мають зосередитися на розробці та тестуванні нових ігрових моделей для обчислювального ядра ІІС, що інтегрують наявні підходи та враховують еволюцію методів атаки та захисту в криптовалютних системах.

Список використаних джерел:

- Гонак І. М. Ризики функціонування криптовалютного бізнесу. *Науковий вісник Херсонського державного університету. Серія "Економічні науки"*, 2021. (44), 81–86.
- Гребельний А. Кіберризки на криптовалютних біржах: причини, наслідки та їх мінімізація. *Modern engineering and innovative technologies*, 2024. (34-01), 106–113.
- Amin M. R. 51% Attacks on Blockchain: Architecture of a Solution for a Proof-of-Work Blockchain to Protect IoT. Bachelor's thesis, International University of Business Agriculture and Technology, Dhaka, Bangladesh. 2020.
- Aponte-Novoa F. A., Orozco A. L. S., Villanueva-Polanco R., Wightman P. The 51% Attack on Blockchains: A Mining Behavior Study. *IEEE access*, 2021. 9, 140549–140564.
- Baruwa Z., Bhattacharjee S., Chandnani S. R., Zhu Z. Perceptions of Social Media Perceptions of 51% Attacks on Proof-of-Work Cryptocurrencies: A Natural Language Processing Approach. arXiv preprint arXiv:2310.14307. 2023.
- Bastiaan M. Preventing the 51%-Attack: A Stochastic Analysis of Two-Phase Proof-of-Work in Bitcoin. Access by link. 2015, January. URL: <https://fmt.ewi.utwente.nl/media/175.pdf>.
- Hao Y. Research of the 51% Attack Based on Blockchain. In 2022 3rd International Conference on Computer Vision, Image and Deep Learning & International Conference on Computer Engineering and Applications (CVIDL & ICCA). IEEE. 2022, May. pp. 278–283.

8. Kim S. K., Yeun C. Y., Damiani E., Al-Hammadi Y. Diverse Perspectives in New Blockchain Design Using Inventive Problem Solving Theory. *Y IEEE International Conference on Blockchain and Cryptocurrency*, Seoul, South Korea. 2019, May.
9. La Morgia M., Mei A., Sassi F., Stefa J. Dogecoin on Wall Street: Analyzing and Detecting Cryptocurrency Pump-and-Dump Manipulations. *ACM Transactions on Internet Technology*, 2023. 23(1), 1–28. DOI: 10.1145/356130
10. Lansiaux E., Tchagaspanian N., Forget J. Social Influence on Cryptocurrency: A Comparative Example on Twitter between Dogecoin and Litecoin. *Frontiers in Blockchain*, 2022. 5, 829865. DOI:10.3389/fbloc.2022.829865
11. Lee S. A. Investigating the impact of cyber security attacks on cryptocurrency markets (doctoral dissertation, Macquarie University). 2022.
12. Nabilou H. Testing the waters of the Rubicon: the European Central Bank and central bank digital currencies. *Journal of Banking Regulation*, 2020. 21(4), 299–314.
13. Niranjani V., Kamachi P. S., Siddharth S., Venkatchalam V., Vishal H. A hybrid approach to minimise 51% attack in cryptocurrencies. In *2022 8th International Conference on Advanced Computing and Communication Systems (ICACCS) IEEE*. 2022, March. Vol. 1, pp. 2100–2103.
14. Rahman A. Implementation of hybrid PoW–PoS to prevent 51% attack in cryptocurrency system (Ph.D. dissertation, United International University). 2019.
15. Raju R. S., Gurung S., Rai P. Огляд атаки 51% в мережі Bitcoin. *Contemporary Issues in Communication, Cloud and Big Data Analytics: Proceedings of CCB 2020*, 2022. 39–55.
16. Ramos S., Pianese F., Leach T., Oliveras E. The Great Crypto Turbulence: Understanding Cryptocurrency Profitability under Attacks. *Blockchain: Research and Applications*, 2021. 2(3), 100021.
17. Sayeed S., Marco-Gisbert H. Assessing Blockchain Consensus and Security Mechanisms against the 51% Attack. *Applied sciences*, 2019. 9(9), 1788. DOI:10.3390/app9091788
18. Sethaput V., Innet S. Blockchain application for central bank digital currencies (CBDC). *Cluster Computing*, 2023. 26(4), 2183–2197.
19. Taylor C. What Happens in a 51% Attack? CoinMarketCap Academy. URL: <https://coinmarketcap.com/academy/article/what-happens-in-51-attacks>
20. Yang X., Chen Y., Chen X. An Efficient Scheme against 51% Blockchain Attacks in Proof-of-Work Systems with weighted historical information. In *2019 IEEE International Conference on Blockchain (Blockchain)* (pp. 261–265). IEEE.
21. Ye, C., Li, G., Cai, H., Gu, Y., & Fukuda, A. (2018, September). Analysis of security in blockchain: Case study in 51%-attack detecting. In *2018 5th International conference on dependable systems and their applications (DSA)*. IEEE. 2019, July. pp. 15–24.
22. Yi X., Fang Y., Wu D., Jiang L. BlockScope: detection and analysis of common vulnerabilities in forked blockchain projects. 2022. arXiv preprint arXiv:2208.00205.

Дата надходження статті: 23.06.2025

Дата прийняття статті: 01.07.2025

Опубліковано: 23.09.2025

УДК 621.391.83(045)
DOI <https://doi.org/10.32689/maup.it.2025.2.17>

Ігнат МИРОШНИЧЕНКО

аспірант кафедри інформаційних технологій,
Державний університет «Київський авіаційний інститут», ignat.myr@gmail.com
ORCID: 0000-0002-7810-2678

ІНТЕГРОВАНІ НАВІГАЦІЙНІ ТА КОМУНІКАЦІЙНІ ПРОТОКОЛИ ДЛЯ АВТОНОМНИХ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ: АРХІТЕКТУРИ, МОДЕЛІ ТА ВИКЛИКИ РЕАЛІЗАЦІЇ В ДИНАМІЧНИХ СЕРЕДОВИЩАХ

Анотація. У роботі представлено підхід до розробки автономних інтелектуальних систем (АІС), що здатні діяти самостійно в динамічних середовищах без постійного втручання оператора. Такі системи складаються з груп інтелектуальних агентів, об'єднаних у мережу зі спільною метою – виконання складних місій.

Метою статті є аналіз обмежень автономного функціонування в частині комунікаційних та навігаційних можливостей елементів системи; формулювання ключових вимог до реалізації зв'язку та навігації в автономних інтелектуальних групах; а також розробка і оцінка відповідної архітектури і моделей, здатних забезпечити стабільну, ефективну та координовану діяльність системи під час виконання поставлених завдань.

Методологія дослідження ґрунтується на системному аналізі функціональних властивостей автономних систем із залученням моделей інформаційної взаємодії для оцінки ефективності комунікаційних і навігаційних процесів. Було застосовано моделювання логічної архітектури та протоколів взаємодії агентів у динамічному середовищі для визначення оптимальних структур управління.

Наукова новизна. Особливу увагу приділено розробці моделей навігації та зв'язку, які виступають основними функціональними компонентами забезпечення автономної мобільності. Автор розкриває функціональні вимоги до систем зв'язку та навігації, включаючи повноту, швидкодію, енергоефективність, узгодженість і стійкість. Наведено аналіз топологій зв'язку (централізовані, децентралізовані, ієрархічні, гібридні) та особливості їх впливу на ефективність взаємодії агентів.

Висновки. У сфері навігації основний акцент зроблено на інтеграцію глобальних супутникових (GNSS) та інерціальних (INS) систем як основи для підвищення точності й стійкості в реальних умовах. У статті запропоновано протокол консенсусної навігації, який дозволяє усунути розбіжності між агентами шляхом обміну даними, забезпечуючи когерентність рішень. Запропоновано архітектурну модель протокольної структури автономної мобільності, яка об'єднує зв'язок і навігацію в логічні шари з високою адаптивністю до конкретних завдань. Такий підхід сприяє створенню стійких, масштабованих та універсальних автономних систем, придатних для застосування в різних галузях, від безпілотних технологій до рятувальних операцій.

Ключові слова: автономні системи, автономні мережі, системи розподіленого інтелекту, навігація, зв'язок.

Ihnat MYROSHNYCHENKO. INTEGRATED NAVIGATION AND COMMUNICATION PROTOCOLS FOR AUTONOMOUS INTELLIGENT SYSTEMS: ARCHITECTURES, MODELS AND IMPLEMENTATION CHALLENGES IN DYNAMIC ENVIRONMENTS

Abstract. The paper presents an approach to the development of Autonomous Intelligent Systems (AIS) capable of operating independently in dynamic environments without continuous operator intervention. Such systems consist of groups of intelligent agents united in a network with a common goal – to carry out complex missions.

Purpose. To analyze the constraints of autonomous operation in terms of the communication and navigation capabilities of system elements; to formulate key requirements for the implementation of communication and navigation in autonomous intelligent groups; and to develop and evaluate the corresponding architecture and models capable of ensuring stable, efficient, and coordinated system performance during task execution.

Methodology. The research methodology is based on a systems analysis of the functional properties of autonomous systems using models of information interaction to assess the effectiveness of communication and navigation processes. Logical architecture modeling and agent interaction protocols in a dynamic environment were applied to determine optimal control structures.

Scientific novelty. Special attention is paid to the development of navigation and communication models, which serve as the main functional components ensuring autonomous mobility. The author outlines functional requirements for communication and navigation systems, including completeness, responsiveness, energy efficiency, consistency, and resilience. The paper analyzes communication topologies (centralized, decentralized, hierarchical, hybrid) and their impact on the efficiency of agent interaction.

Conclusions. In the field of navigation, the focus is on integrating Global Navigation Satellite Systems (GNSS) and Inertial Navigation Systems (INS) as a basis for increasing accuracy and robustness in real-world conditions. The article proposes a consensus-based navigation protocol that resolves discrepancies between agents through data exchange, ensuring coherence in decision-making. An architectural model of a protocol structure for autonomous mobility is proposed, integrating communication and navigation into logical layers with high adaptability to specific tasks. This approach promotes the creation of robust, scalable, and versatile autonomous systems suitable for use in various domains, from unmanned technologies to rescue operations.

Key words: autonomous systems, autonomous networks, distributed intelligence systems, navigation, communication.

© І. Мирошніченко, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

Постановка проблеми в загальному вигляді та її зв'язок із важливими науковими чи практичними завданнями. Актуальність теми зумовлена сучасними глобальними викликами, що потребують високого рівня технологічної готовності до реагування в умовах підвищеної невизначеності, небезпеки та обмеженого людського доступу. Автономний інтелект стає ключовим напрямом розвитку інтелектуальних систем, оскільки він забезпечує не лише виконання задач без участі людини, але й підвищує рівень адаптивності, безпеки та оперативності в реальному часі.

Особливої важливості це набуває у сферах, де критичне значення має швидке отримання точних навігаційних і комунікаційних даних: під час ліквідації наслідків стихійних лих, проведення пошуково-рятувальних операцій, моніторингу екологічних змін у важкодоступних місцях, спостереження за лісовими пожежами та забрудненням природних ресурсів. Впровадження ефективних протоколів обміну інформацією між автономними елементами забезпечує не лише стабільну взаємодію в межах системи, але й підвищує загальну ефективність виконання місій. У цьому контексті дослідження принципів та моделей навігаційної і комунікаційної взаємодії між агентами є критично важливим для забезпечення надійної та скоординованої поведінки автономних мобільних платформ.

Аналіз останніх досліджень і публікацій. У науковій літературі проблематика організації автономних груп інтелектуальних елементів є об'єктом активного вивчення як вітчизняними, так і зарубіжними дослідниками, які пропонують численні моделі та стратегії організації автономних груп інтелектуальних елементів. Але такі аспекти, як вимоги до функцій зв'язку та навігації, що забезпечують ефективну та надійну діяльність автономних систем в умовах динамічного середовища, а також концепціям об'єктно-орієнтованого сприйняття, мобільності та адаптивної взаємодії агентів все ж не до кінця досліджені. Так, у своїй роботі Р. Бабає та А. Ехяей [4] дослідили керування групою безпілотних апаратів для транспортування вантажів, запропонувавши алгоритм на основі розширеного фільтра Кальмана та теорії Мура-Пенроуза для стійкого керування. Лю Лю в своїй роботі [10] розробила протокол для розподіленої системи BPLA, що підтримує зміну формації при фіксованій топології. Уолкер Голлник [6] запропонував методологію проектування міської повітряної мобільності з оцінкою її технічної й експлуатаційної доцільності. Кучеров Д., Фу М., Козуб А. [2] запропонували навігаційні та комунікаційні протоколи для ефективної взаємодії елементів автономних систем у динамічних середовищах.

Тенедоріо та співавтори [15] узагальнили інновації в геоінформаційних системах та принципах «зеленого» урбанізму для сталого розвитку міст. Сяохай Юй ін. [16] запропонували протокол «лідер-послідовник» для транспортної платформи на базі Arduino з урахуванням кута нахилу дороги. Лі Юна та Дженг Чана [9] розробили кооперативне керування морськими буксирами з інтелектуальним протоколом для стабільної навігації. К. Шарма та Н. Гонді [14] здійснили огляд таких протоколів Інтернету, як CoAP, MQTT, XMPP та інші, що застосовуються на різних рівнях архітектури, і проаналізували їхню ефективність та надійність за критеріями енергоефективності, безпеки та простоти впровадження. М. Сакіб та співавтори [13] досліджували питання безпеки та поточні атаки, характерні для кожного рівня архітектури Інтернету речей. Основи фундаментальних комунікаційних протоколів – TCP, UDP, IP, каналного та фізичного рівнів – були розглянуті в роботі Д. Г. Мораїса [11] та представлено аналіз стеків протоколів, що використовуються в сучасних системах передачі даних, зокрема в контексті бездротових технологій, таких як 5G NR, Wi-Fi 6 та Bluetooth LE 5. Автор детально розглядає функціонування протоколів на різних рівнях мережевої моделі, що є важливим для розуміння архітектури бездротового зв'язку у смартфонах.

Еом Тхе-Іль та ін. [5] описали стек протоколів через MSC-діаграми, а Д. Херцог [8] використав OSI та TCP/IP як приклади складних моделей протоколів для надійного проектування. Р. Олфаті-Сейбер [12] запропонував нові методи керування інтелектуальними системами на основі поведінкових реакцій, а Жчанг К. [17] створив розподілений протокол для мобільних транспортних засобів, що підтримує групову конфігурацію.

Аналіз досліджень вказує на існування обмежень сучасних автономних інтелектуальних систем: недостатній рівень автономії, вузька спеціалізація, обмежений інтелектуальний потенціал у розподілених сценаріях і низька масштабованість через залежність від людських ресурсів. Це ускладнює їх застосування у складних динамічних середовищах.

Метою статті є вивчення обмежень автономної роботи щодо комунікаційних та навігаційних функцій елементів; визначення основних вимог до комунікаційних та навігаційних рішень автономних інтелектуальних груп; та на основі висунутих вимог, пропонування та оцінка рішень, таких як архітектури та моделі комунікаційних та навігаційних функцій, які забезпечать ефективну та надійну роботу групи при виконанні її місії.

Виклад основного матеріалу. Автономну групу можна визначити як взаємодіючу групу інтелектуальних елементів, здатних приймати рішення, які мають визначену місію, зокрема, з чітко визначеними

критеріями успіху; потребують взаємодії елементів під час виконання; здатні працювати без постійного контролю з боку керуючого органу (оператора) у будь-який час або протягом тривалого часу.

Розподілена автономна група також вирізняється умовою, що всі елементи групи мають подібний рівень інтелекту, тобто всі здатні приймати рішення з однаковими або подібними вхідними факторами, на відміну від моделі центрального інтелекту, де один або підмножина елементів має вищий рівень інтелекту та приймає рішення, що виконуються іншими моделями (також відома як модель «головний-підлеглий» або «лідер-наслідувач»).

Проблему розробки та впровадження надійних та ефективних протоколів зв'язку та навігації можна розглядати як один з критично важливих елементів ефективної роботи групи. Розглянемо основні вимоги, яким повинні відповідати ефективні та надійні комунікаційні вузли та протоколи групи, на основі сформульованих вимог опишемо та розглянемо моделі автономних комунікаційних мереж та оцінімо запропоновані моделі в рамках цілей автономної роботи та місії.

Операційна модель автономних груп може відрізнятися від операційної моделі систем та мереж, керованих людиною, кількома способами, зокрема:

- елементи повинні приймати власні рішення під час роботи, не покладаючись на команди оператора;
- елементи повинні підтримувати постійний та надійний зв'язок зі своїми колегами для обміну важливою інформацією, необхідною для надійної роботи;
- протоколи систем зв'язку та навігації повинні бути надійними та стійкими до факторів навколишнього середовища та несприятливих впливів, як природних, так і техногенних;
- системи повинні працювати в межах обмежень на ресурси, що накладаються умовами автономної мобільності та функціонування;
- на найвищому рівні автономності система повинна бути здатною працювати в повністю автономному режимі протягом тривалого часу без будь-якого контролю з боку оператора.

Моделі автономної роботи, описані вище, диктують наступні основні вимоги до комунікаційних та навігаційних (Comm&Nav, C&N) компонентів автономних інтелектуальних систем. Для забезпечення ефективного функціонування автономних інтелектуальних систем у груповій взаємодії необхідно дотримуватись низки критичних вимог до інформаційного обміну та комунікаційних процесів. Насамперед, важливою є повнота отримуваної інформації: кожен елемент системи повинен мати доступ до всіх необхідних вхідних даних для прийняття рішень як через власні сенсори, так і через обмін інформацією з іншими учасниками групи. Ефективність комунікацій визначається раціональним використанням ресурсів, зокрема енергії, що є особливо важливо в умовах автономного функціонування. Крім того, швидкодія передавання даних повинна відповідати часовим обмеженням, накладеним динамікою середовища та необхідністю оперативного реагування.

Ще однією важливою вимогою є узгодженість та послідовність інформації між елементами системи, що гарантує єдність ситуаційного сприйняття та скоординованість дій. Надійність і стійкість комунікаційної інфраструктури до впливу несприятливих зовнішніх умов (наприклад, завад, перебоїв зв'язку або обмеженого енергозабезпечення) є критичними для забезпечення стабільної роботи автономних систем у реальному середовищі. Сукупність цих характеристик формує основу для побудови надійних, ефективних та адаптивних колективних АІС.

Одним із ключових завдань проектування автономних інтелектуальних систем (АІС), здатних відповідати вимогам до критичних систем у складі автономних груп, є організація внутрішньогрупової комунікації та характер інформаційного обміну між елементами. Від структури взаємодії значною мірою залежать ефективність прийняття рішень, узгодженість дій та загальна надійність функціонування системи в цілому. На концептуальному рівні виділяють кілька базових моделей організації АІС. Централізована модель (типу «головний – підлеглий») передбачає прийняття рішень окремою підмножиною елементів, які формують команди для інших учасників групи. Децентралізована або повністю розподілена модель ґрунтується на автономному прийнятті рішень кожним елементом на основі даних, отриманих із середовища та від інших учасників. Гібридна модель поєднує риси централізованого та децентралізованого підходів: певні рішення делегуються уповноваженим елементам, тоді як інші приймаються децентралізовано. Ієрархічна модель передбачає поділ системи на множину автономних підгруп, що координуються вищим рівнем управління із застосуванням відповідних протоколів координації та управління.

Конфігурації організаційних та комунікаційних функцій автономних груп проілюстровано на (рис. 1).

Розглянемо логічну організацію таких функцій в автономній інтелектуальній системі та припустимо, що елементи групи оснащені відповідними фізичними інструментами для встановлення та підтримки зв'язку, необхідного для виконання місії та в межах обмежень автономної роботи.

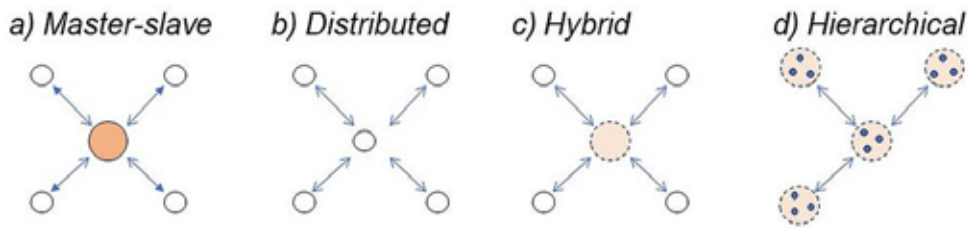


Рис. 1. Організаційні та комунікаційні моделі автономних інтелектуальних систем

Вимоги повноти та зв'язності підкреслюють необхідність того, щоб комунікаційний рівень забезпечував можливість передачі інформації між будь-якими елементами, якими необхідно обмінюватися під час операції, у достатньому обсязі та протягом встановленого часу; це можна інтерпретувати як обмеження безперервної та безперебійної логічної зв'язності, rc ; та мінімальної гарантованої швидкості передачі інформації, ri . Логічну зв'язність можна реалізувати за допомогою низки моделей фізичної зв'язності в динамічній автономній групі: повністю взаємопов'язана сітчаста, точка-точка, центральний диспетчер, ширококомовна, ретрансляційна та інші.

Вибір моделі зв'язку може бути продиктований умовами експлуатації та типом групи/системи. Наприклад, у повністю розподіленій автономній системі протоколи зв'язку можуть виконуватися будь-яким елементом групи, але моделі центрального та диспетчерського типів будуть нездійсненними через спеціальні функції/ролі вибраних елементів.

Далі, потрібно буде оцінити моделі зв'язності щодо основних вимог, викладених раніше. Наприклад, розглянемо взаємозв'язок між потужністю, необхідною для надійного зв'язку, та вартістю програмної обробки запитів на зв'язок в ієрархічній моделі (рис. 1, d).

Припустимо, що мінімальна потужність, необхідна для надійної обробки запиту, дорівнює p_m . Тоді, щоб мати змогу обробляти комунікаційні команди на відстані rc , потужність передавача повинна бути:

$$p_m = \alpha \frac{P_t}{r_c^2}, \quad (1)$$

де P_t – потужність передавача елемента,
 α – додатна константа.

Якщо розмір операційної групи D_g перевищує проміжок зв'язку rc , необхідно ввімкнути режим ретрансляції, який обробляє та повторно передає запит, доки він не досягне пункту призначення. Деякі події повторної передачі можна оцінити як:

$$nt \sim \frac{D_g}{r_c} = \frac{\beta}{r_c}, \quad (2)$$

де β – const, оскільки розмір групи вважається постійним у межах обраної моделі роботи.

Для того щоб зрозуміти, як вибір потужності передачі, а отже, і тривалості зв'язку, вплине на вартість програмних обчислень (циклів процесора) під час обробки запитів на зв'язок. Вартість програмного забезпечення для обробки запиту розраховується за формулою:

$$S_c \propto n_t = \frac{\gamma}{\sqrt{P_t}} \quad (3)$$

З (3) випливає, що зменшення потужності передачі в комунікаційному модулі автономної групи може призвести до збільшення необхідної обчислювальної потужності в утилітах обробки запитів. З іншого боку, з моделлю ширококомовлення, можна отримати такі результати:

$$P_t \sim p_m D_g^2; S_c \sim N_g \quad (4)$$

де N_g – кількість елементів у групі/системі, оскільки кожен елемент має обробляти кожен запит, що надходить від усіх інших елементів. Тоді як потужність передачі, так і вартість програмного забезпечення можна вважати постійними в межах обраної архітектури та конфігурації групи, і балансування як у (3) вище, неможливе.

Очевидно, що врахування таких міркувань під час проектування архітектури та операційних моделей автономних груп може бути важливим і навіть критично важливим для ефективної та надійної роботи систем.

Розглянуті приклади моделей організації комунікаційних функцій в автономних мобільних системах ілюструють застосування суттєвих обмежень на операційні параметри групи, що впливає з основних вимог, сформульованих раніше. Враховуючи обмеження формату, детальний огляд комунікаційних моделей та відповідних обмежень на параметри та компоненти системи, що виникають із вимог автономної роботи, сформульованих раніше, буде наведено далі.

Навігація є важливим компонентом мобільної АІС, який відповідає за виконання дій та функцій, необхідних для місії, у правильному географічному місці та у правильний час. Навігаційна компонента забезпечує основу для наступних функцій, необхідних для базової роботи та повноцінного функціонування автономної системи:

- визначення точного місцезнаходження (у тривимірному просторі), орієнтації, напрямку та інших параметрів траєкторії;
- перевірка правильності, з необхідною точністю, розташування, напрямку та інших параметрів руху;
- функції, процедури та протоколи для підтримки мобільності, пересування та позиціонування в бажаному географічному місці;
- стійкість та відновлення у випадках функціональних та елементних збоїв; та інші.

Наразі глобальні навігаційні супутникові системи (GNSS), такі як GPS та інші, вважаються стандартом та методом вибору для навігаційної підтримки пілотованих та автономних систем, включаючи безпілотні літальні апарати, наземні та морські роботи (з адаптацією технології), використання яких може бути ефективним для виконання різних функцій та складних завдань у цивільній, безпековій та військовій сферах. Однак GNSS, як єдине джерело та основа навігаційної підтримки може мати важливі обмеження: вона може бути недоступна в усіх середовищах, таких як під водою або через перешкоди тощо, і з цієї причини зазвичай шукають додаткові або додаткові методи підвищення надійності навігаційної функції.

Поєднання інерціальної навігаційної системи (INS) та (GNSS) стає поширеним вибором, який може компенсувати обмеження та недоліки ефірних систем. Наприклад, у міських районах система GNSS може бути недоступною через ефекти блокування сигналу, що можна компенсувати навігаційними розрахунками. З іншого боку, інерціальна система може підтримувати задовільну точність лише на коротких відстанях, і система GNSS допомагає їй зменшити накопичення похибок вимірювання з відстанню [7].

Таким чином, спільне використання цих систем може дозволити, з одного боку, обмежити поширення помилок менш точної, але більш інформативної інерціальної системи, а з іншого боку, збільшити швидкість доставки інформації до бортових споживачів, значно покращивши стійкість та зменшивши шумову складову помилок високоточної супутникової системи. З окреслених причин інтегровані інерціально-супутникові навігаційні системи (ІІНС) стають дедалі популярнішими в застосуваннях, що вимагають високої надійності та стійкості навігаційної системи, включаючи автономні інтелектуальні системи, де вона може мати критично важливий вплив.

Цей короткий огляд демонструє, що наразі існує ряд варіантів для проектування та розробки навігаційних компонентів і функцій автономних інтелектуальних систем для широкого спектру завдань, включаючи критично важливі, в різних середовищах: наземних, морських, підводних та авіаційних сферах застосування. Це забезпечує надійну та узгоджену роботу навігації як важливої та критичної логічної функції (рівня) автономної системи.

Для ілюстрації розглянемо приклад. Припустимо, що автономна система, що складається з інтелектуальних агентів $\{a_1 \dots a_N\}$, які використовують модель широкомовного зв'язку, спирається на навігацію. Припустимо, що агенти приймають рішення щодо мобільності під час виконання місії на основі навігаційних вимірювань, виконаних елементами, $M = \{m_1 \dots m_N\}$.

З причин, коротко описаних раніше, немає гарантії, що набір вимірювань $M(t)$ буде завжди узгодженим між елементами. Наприклад, навігаційний модуль на певному елементі може вийти з ладу або зазнати несприятливої статистичної події тощо. Тоді рішення, прийняті таким елементом, можуть бути несумісними з рішеннями інших елементів групи, і він може зрештою втратити узгодженість у роботі, що негативно вплине на успіх місії.

Щоб протидіяти такій можливості, можна реалізувати протокол консенсусу/когерентності на основі моделі широкомовного зв'язку, за допомогою якого агенти періодично обмінюються своїми навігаційними вимірюваннями з найближчими сусідами або всією групою. Кожен агент, наприклад потім k приймає рішення на основі результатів вимірювання: m_k та інформації, отриманої від групи. Якщо вимірювання виявляються узгодженими в межах певного запасу, дозволеного функцією групової операції, може бути виконана звичайна послідовність рішень; навпаки, якщо вимірювання елемента

виявляється несумісним з вимірюваннями групи, може бути викликано альтернативне (резервне) консенсусне рішення на основі даних, сумісних для більшості агентів.

Іншими словами, навігаційне рішення $D(m_k)$ модифікується за допомогою функції прийняття рішень $\tilde{D}(m_k, M)$, яка враховує вимірювання інших агентів та призводить до консенсусного рішення $C(M)$ у разі розбіжностей:

$$\tilde{D}(m_k, M) = D(m_k), m_k \notin M; \text{else } \tilde{D}(m_k, M) = C(M) \quad (5)$$

Обмін навігаційними даними в автономній групі як основа для консенсусу, що забезпечує узгодженість у прийнятті рішень щодо навігації та мобільності, проілюстровано на (рис. 2).

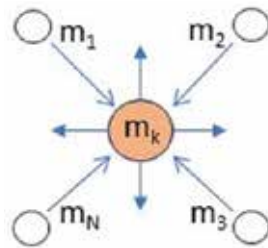


Рис. 2. Протокол навігації на основі консенсусу/когерентності в автономній мобільній системі

Наведений вище приклад протоколу інтелектуальної когерентної навігації демонструє можливість проектування інтелектуальної взаємодії між елементами автономних інтелектуальних систем, які підтримують узгоджену, надійну та стійку роботу в несприятливих умовах та у випадках несправності та збоїв.

Функції зв'язку та навігації в протокольному структурі автономної інтелектуальної мобільності, що обговорюються в цій роботі, забезпечують критично важливі базові послуги для автономної інтелектуальної системи, дозволяючи інтелектуальним агентам обмінюватися інформацією, необхідною для виконання функцій та операцій, визначати місцезнаходження, виконувати завдання напрямку та руху тощо.

З точки зору ефективної організації функцій різних типів в інтелектуальних системах високої складності, інкапсуляція таких функцій та утиліт у логічні групи або «шари», визначені відповідними протоколами, видається природним рішенням, яке привернуло увагу дослідницької спільноти Еом Тхелль та ін [5], Херцог Д. [8], Шмельова Т. [3], Кучеров Д. [1]. Запровадження інкапсульованої протокольної структури може забезпечити низку суттєвих переваг порівняно з жорстко закодованими («запрограмованими до програми») рішеннями. Зокрема, спостерігається спрощення процесу проектування та впровадження за рахунок використання сумісних стандартних функцій, що сприяє скороченню повного циклу – від розробки до експлуатації. Ізоляція функціональних рівнів у межах інкапсульованої архітектури забезпечує підвищену надійність системи, зменшуючи ризики міжрівневих збоїв. Крім того, така структура забезпечує високу гнучкість дизайну та функціональну універсальність, що дозволяє динамічно адаптувати систему до різних умов експлуатації та специфіки завдань.

Висновки і перспективи подальших досліджень. Узагальнюючи викладене, встановлено, що автономні інтелектуальні системи (AIC) залишаються актуальним об'єктом наукових досліджень. Попри досягнутий прогрес, вони й досі характеризуються обмеженнями щодо автономії, масштабованості, функціональної гнучкості та адаптації до складних середовищ. Проаналізовано сучасні підходи до побудови протоколів і моделей мобільних систем з автономним інтелектом, з акцентом на критично важливі аспекти комунікації та навігації. Запропоновано методологію аналізу дій агентів, що базується на структурній організації групової поведінки, а також визначено вимоги до систем зв'язку та навігації й окреслено ключові компоненти архітектури автономних систем. Найбільш перспективними визнано гібридні навігаційні рішення (наприклад, GNSS/INS), які дозволяють підвищити точність та ефективність. Для підвищення функціональності AIC запропоновано універсальну протокольну структуру автономної мобільності, що забезпечує модульність, надійність, сумісність і спрощення проектування завдяки стандартизованим функціям. Очікується, що стандартні структури спільних утиліт та функцій слугуватимуть основою для появи простору та екосистеми мобільних автономних інтелектуальних систем з різноманітним та універсальним застосуванням у різних завданнях та областях.

Список використаних джерел:

1. Кучеров Д., Долгих С., Мирошніченко І., Пошивайло О. Протокольні та логічні моделі для забезпечення надійності та стійкості автономних систем БПЛА. Матеріали 13-ї Міжнародної конференції з надійних систем, послуг і технологій. Афіни, Греція, 2023. С. 1–7. doi: 10.1109/DESSERT61349.2023.10416491 (дата звернення: 28.04.2025)

2. Кучеров Д., Фу М., Козуб А. Синтез законів керування рухом групи БПЛА з природними перешкодами. Автоматизовані системи в авіаційній та аерокосмічній галузях: зб. наук. праць / ред. Т. Шмельова, Ю.В. Шмельова, Н. Сікирда, Н. Різун, Д. Кучеров, К. Дергачов. Суми : Вид-во СумДУ, 2019. С. 193–219.
3. Шмельова Т., Ковальов Ю. Н., Долгих С., Бурлака О. Оптимізація польоту автономних груп БПЛА на основі геометричного моделювання. *Матеріали 5-ї Міжнародної конференції «Актуальні проблеми розробки безпілотних літальних апаратів»*, Київ, Україна, 2019. С. 79–82. doi: 10.1109/APUAVD47061.2019.8943856. (дата звернення 01.05.2024)
4. Babaie R., Ehyaei A. Cooperative Control of Multiple Quadrotors for Transporting a Common Payload. *AUT Journal of Modeling and Simulation*. 2018. № 50(2). P. 147–156. doi: 10.22060/miscj.2018.14252.5100. (date of access: 15.04.2025)
5. Eom T.-I., Lee W.-Y., Lee D.-Y., Kim J.-H., Hur W.-H. Procedure-based development platform for communication protocol stack software. In *Proceedings of the IEEE Conference on Advanced Microsystems for Automotive Applications (CAMAD) 2016*. P. 12–17. doi: <https://doi.org/10.1109/CAMAD.2016.7790323>. (date of access: 17.04.2025)
6. Gollnick M., Niklass M., Swaid M. Methodology and initial results of the assessment of UAM concept potential. *Proceedings of Aerospace Europe Conference, 2020. Hamburg*, 1–11. URL: <https://www.academia.edu/66696287> (date of access: 20.04.2025)
7. Hansen J. M., Fossen T. I., Johansen T. A. Nonlinear observer design for GNSS-aided inertial navigation systems with time-delayed GNSS measurements. *Control Engineering Practice*, 2017. 60, 39–50. doi: <https://doi.org/10.1016/j.conengprac.2016.11.016>. (date of access: 18.04.2025)
8. Hercog D. Protocol stack. In *Communication protocols: Principles, methods, and specifications. Springer International Publishing*. 2020. P. 67–81. doi: https://doi.org/10.1007/978-3-030-50405-2_5 (date of access: 20.04.2025)
9. Li Y., Zheng J. Design and implementation of cooperative turning control for the towing system of unpowered facilities. *IEEE Access*. 2018 № 99. P. 1–5. doi: <https://doi.org/10.1109/ACCESS.2018.2819692> (date of access: 21.04.2025)
10. Liu L., Lian, X., Zhu C., He L. Distributed cooperative control for UAV swarm formation reconfiguration based on consensus theory. *2nd International Conference on Robotics and Automation Engineering (ICRAE)*. 2017. P. 264–268. doi: <https://doi.org/10.1109/ICRAE.2017.8291392> (date of access: 1.05.2025)
11. Morais D. H. Data communication systems protocol stacks. *5G NR, Wi-Fi 6, and Bluetooth LE 5 (Chapter 2)*. 2023. doi: https://doi.org/10.1007/978-3-031-33812-0_2 (date of access: 15.04.2025)
12. Olfati-Saber R. Flocking for multi-agent dynamic systems: Algorithms and theory. *IEEE Transactions on Automatic Control*. 2006. № 51(3). P. 401–420. doi: <https://doi.org/10.1109/TAC.2005.864190> (date of access: 01.05.2025)
13. Saqib M., Jasra B., Moon A. H. A systematized security and communication protocols stack review for Internet of Things. *IEEE International Conference for Innovation in Technology (INOCON)*. Bangluru, India. 2020. P. 1–9. doi: <https://doi.org/10.1109/INOCON50539.2020.9298196> (date of access: 28.04.2025)
14. Sharma C., Gondhi N. K. Communication protocol stack for constrained IoT systems. In *2018 3rd International Conference on Internet of Things: Smart Innovation and Usages (IoT-SIU)*. Bhimtal, India. 2020. doi: <https://doi.org/10.1109/IoT-SIU.2018.8519904> (date of access: 28.04.2025)
15. Tenedório A., Estanqueiro R., Henriques C. D. Methods and applications of geospatial technology in sustainable urbanism. Pennsylvania, USA: IGI Global. 2021. doi: <https://doi.org/10.4018/978-1-7998-2249-3> (date of access: 26.04.2025)
16. Yu X., Guo G., Gong J. Cooperative control for a platoon of vehicles with leader-following communication strategy. *Chinese Automation Congress (CAC)*. 2020. P. 6721–6726. URL: <https://ieeexplore.ieee.org/document/8243988>. (date of access: 01.05.2025)
17. Zhang Q., Lapierre L., Xiang X. Distributed control of coordinated path tracking for networked nonholonomic mobile vehicles. *IEEE Transactions on Industrial Informatics*, 2013. № 9(1). P. 472–484. doi: <https://doi.org/10.1109/TII.2012.2219541> (date of access: 17.04.2025)

Дата надходження статті: 13.06.2025

Дата прийняття статті: 17.06.2025

Опубліковано: 23.09.2025

УДК 004.852:519.85

DOI <https://doi.org/10.32689/maup.it.2025.2.18>**Костянтин МІНКОВ**

аспірант кафедри програмного забезпечення комп'ютерних систем,

Чернівецький національний університет імені Юрія Федьковича

ORCID: 0009-0007-3400-050X

ВИКОРИСТАННЯ ACTIVE LEARNING ЯК ІНСТРУМЕНТУ ДЛЯ УТОЧНЕННЯ МІТОК У ЗАДАЧАХ, ЗАСНОВАНИХ НА ПРИНЦИПАХ СЛАБКОГО НАВЧАННЯ

Анотація. Було запропоновано ітеративний метод, який використовує активне навчання для виявлення та корекції помилкових міток у слабо анотованому наборі даних. Вихідним є набір даних, у якому приклади мають лише зашумлені або частково надані анотації. Запропонований підхід поступово уточнює навчальну вибірку шляхом ідентифікації зразків, які з високою ймовірністю містять помилкові помічення, та передбачає залучення цілеспрямованої перевірки, як з боку людини, так і за допомогою евристичних правил із високим рівнем довіри.

Метою статті є підвищення якості міток у слабо анотованому наборі даних із мінімальними витратами на ручну анотацію. Для цього запропоновано ітеративний підхід, що базується на активному навчанні та селективній валідації з боку людини.

Методологія. Розроблено ітеративний метод, у якому дискримінативна модель (DistilBERT) навчається на поточному наборі міток, оцінює невизначеність своїх передбачень на основі варіативності виходів, а для зразків із високою невизначеністю генерує псевдомітки. Ці мітки передаються на швидке підтвердження або виправлення анотаторам. В основу слабого анотування покладено чотири евристичні λ -функції, що формують початкові мітки, а генеративна модель враховує залежності між цими функціями через регуляризовану матрицю зв'язків. Експериментально досліджено ефективність підходу на наборі IMDb Movie Reviews, що містить 50 000 текстових відгуків, з чітким розподілом на навчальну, валідаційну й тестову вибірки.

Наукова новизна. На відміну від традиційних методів активного навчання, які передбачають ручну переанотацію найневпевненіших зразків, запропоновано гібридну стратегію, що поєднує слабе навчання, псевдоанотування та вибірково людську перевірку. Це дозволяє досягати цільової точності моделі в 2–3 рази швидше, ніж класичні підходи без слабких міток, за рахунок ефективного використання обмеженого людського ресурсу та стартової інформації з евристичних правил.

Висновки. Ітеративна стратегія, яка поєднує відбір зразків на основі розбіжності передбачень, автоматичне псевдоанотування та селективну людську валідацію, дозволяє ефективно покращити якість міток у слабо анотованих даних без повної переанотації. Запропонований метод продемонстрував конкурентні результати на реальному корпусі відгуків IMDb, забезпечуючи високу точність моделі класифікації зі зниженими витратами на ручну працю.

Ключові слова: псевдомітки, матриця переходів, евристичні правила, оракул, зважування за значущістю, впевненість моделі.

Kostiantyn MINKOV. THE USE OF ACTIVE LEARNING FOR REFINING LABELS IN WEAK SUPERVISION BASED TASKS

Abstract. An iterative method has been proposed that uses active learning to detect and correct mislabeled data in a sparsely annotated dataset. The starting point is a dataset in which examples have only noisy or partially provided annotations. The proposed approach gradually refines the training sample by identifying examples that are highly likely to contain incorrect labels and involves targeted verification, both by humans and using heuristic rules with a high level of confidence.

The goal of this paper is to improve the quality of labels in a sparsely annotated dataset with minimal manual annotation costs. To do this, we propose an iterative approach based on active learning and selective human validation.

Methodology. An iterative method has been developed in which a discriminative model (DistilBERT) is trained on the current set of labels, evaluates the uncertainty of its predictions based on the variability of outputs, and generates pseudo-labels for samples with high uncertainty. These labels are sent to annotators for quick confirmation or correction. Weak annotation is based on four heuristic λ -functions that form the initial labels, and the generative model takes into account the dependencies between these functions through a regularized connection matrix. The effectiveness of the approach was experimentally investigated on the IMDb Movie Reviews dataset, which contains 50,000 text reviews, with a clear division into training, validation, and test samples.

Scientific novelty. Unlike traditional active learning methods, which involve manual re-annotation of the most uncertain samples, a hybrid strategy is proposed that combines weak learning, pseudo-annotation, and selective human verification. This allows the target model accuracy to be achieved 2–3 times faster than classical approaches without weak labels, due to the efficient use of limited human resources and initial information from heuristic rules.

Conclusions. An iterative strategy that combines sample selection based on prediction divergence, automatic pseudo-annotation, and selective human validation allows for effective improvement of label quality in weakly annotated data without complete re-annotation. The proposed method demonstrated competitive results on a real-world IMDb review corpus, providing high classification model accuracy with reduced manual labor costs.

Key words: pseudo-labels, transition matrix, heuristic rules, oracle, importance weighting, model confidence.

© К. Мінков, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

Вступ. У машинному навчанні одним із головних викликів залишається потреба у великій кількості високоякісних анотованих даних. Повноцінна ручна анотація таких наборів даних є трудомісткою, дорогою та часто нездійсненною у прикладних умовах. Саме тому дедалі більшої уваги набувають методи, що дозволяють ефективно навчатися на частково анотованих або зашумлених вибірках. Зокрема, підходи, що ґрунтуються на слабкому навчанні, дають змогу використовувати евристичні, правила, автоматично згенеровані сигнали чи агреговані джерела як сурогатні мітки.

Проте слабкі мітки можуть бути ненадійними або хибними, що без відповідної корекції призводить до значного зміщення у моделі. Щоб подолати цю проблему, останні дослідження все активніше застосовують активне навчання як механізм, що дозволяє вибірково запитувати інформацію про найбільш невизначені зразки.

Аналіз останніх досліджень і публікацій. Бігель С. з колегами [2] запропонували метод Active WeaSuL, який інтегрує активне навчання у парадигму слабого навчання задля підвищення якості маркування в умовах обмеженого доступу до вручну анотованих даних. У рамках Active WeaSuL передбачено залучення обмеженої кількості справжніх міток, що надаються експертами вибірково, тобто лише для тих зразків, які викликають найбільшу невизначеність у моделі. Модель базується на генеративному підході до об'єднання слабких міток, який модифікується за допомогою додаткової пеналізації до функції втрат. Ця пеналізація змушує модель узгоджувати ймовірнісні мітки зі справжніми мітками, отриманими в межах активного навчання. Замість локального переписування міток, запропонований метод коригує саму логіку агрегації слабких сигналів. Щоб ефективно визначати найбільш інформативні приклади для запиту, використовується стратегія вибірки на основі максимального розходження між ймовірнісним прогнозом моделі і частотним розподілом експертних міток у відповідному кластері (bucket), що оцінюється через дивергенцію Кульбака-Лейблера. Експериментальні дослідження охоплюють як синтетичні дані (двовимірні гаусіани), так і реальні набори даних, тобто виявлення візуальних відносин і класифікацію спаму на YouTube. У всіх випадках модель демонструє стабільну перевагу над базовими стратегіями [2].

В межах дослідження Ванга Ю. та ін. [5] представлено новий підхід до задачі виявлення об'єктів, який поєднує активне навчання з методами як зі слабким так і неповним підкріпленнями. ALWOD реалізує механізм warm-start активного навчання, який ґрунтується на попередньому тренуванні моделей за допомогою невеликого підмноження повністю анотованих зображень і великого слабо анотованого набору, доповненого синтетично згенерованими зображеннями. Суттєвою інновацією є використання генератора зображень, який дозволяє створити допоміжний набір повністю анотованих зображень шляхом вставки об'єктів з невеликої кількості справжніх FA-зображень у фонові сцени. Це дозволяє навчити першу модель в умовах обмеженої анотації, що забезпечує високу стартову продуктивність. На відміну від традиційних підходів, ALWOD не вимагає повного створення нових анотацій: замість цього анотатори виправляють або уточнюють вже запропоновані детектором межі об'єктів та категорії [5].

Розширення активного навчання, яке дозволяє вибирати найінформативніші зразки для анотації, а також контролювати точність цих анотацій, розроблено Олміном А., Дж. Ліндквістом, Свенссоном Л. та Ліндстеном Ф. [1]. Ідея науковців базується на тому, що менш точні анотації є дешевшими, тому їх можна зібрати більше в межах того самого бюджету, розширивши покриття простору ознак. Ключовим внеском є нова функція вибору (acquisition function), яка базується на взаємній інформації між слабкою анотацією та моделлю (а не між анотацією і справжньою міткою, як у попередніх підходах). Це дозволяє уникнути залежності від латентної змінної й адаптувати метод до будь-якого типу анотаційного шуму. Усі обчислення були інтегровані у фреймворк Gaussian Processes, що надає ймовірнісну оцінку та можливість ефективної побудови невизначеності. Результати експериментів на синусоїдальних функціях, наданих UCI та синтетичних задачах класифікації демонструють, що облік точності анотацій дозволяє суттєво підвищити продуктивність моделі за фіксованого бюджету. Особливо сильний ефект спостерігається у випадках, коли низькоточні анотації є значно дешевшими [1].

З огляду на проаналізовану наукову літературу, слід відзначити, що питання застосування підходу Active Learning як засобу уточнення міток у задачах, заснованих на принципах слабого навчання, досі залишається недостатньо розробленим. Недостатність ґрунтовних досліджень у цьому напрямі вказує на потребу в подальшому теоретичному осмисленні та емпіричному вивченні зазначеної проблематики з метою підвищення ефективності методів слабо контрольованого навчання.

Метою статті є підвищення якості міток в слабо поміченому наборі даних із мінімальними витратами на анотацію шляхом ітеративного відбору зразків із високою розбіжністю передбачень, генерації для них псевдоміток та залучення людини лише для підтвердження або корекції цих машинних пропозицій.

Виклад основного матеріалу дослідження. Па початку поточного дослідження у центрі уваги постає слабо помічений набір даних, який виражається як:

$$D = \left\{ \left(x_i, \hat{y}_i \right) \right\}_{i=1}^N \quad (1)$$

в якому кожна мітка $\tilde{y}_i$ може бути шумною, або вилученою за допомогою простих правил, а не мати надійне джерело походження.

Незважаючи на таку невизначеність, проводиться навчання первинної моделі f_0 , яка може бути як глибокою нейронною мережею, так і іншою параметричною моделлю. Навчання проводиться шляхом мінімізації середньої сурогатної втрати слабких міток (weak labels):

$$\hat{\theta} = \arg \min_{\theta} \frac{1}{N} \sum_{i=1}^N \ell(f_{\theta}(x_i), \tilde{y}_i) \quad (2)$$

де $\ell(\cdot, \cdot)$ може представляти як перехресну ентропію для функцій класифікації, так і середнє квадратичне похибки для функцій регресії [6].

Така стратегія «Warm-Start» забезпечує модель першим необробленим відображенням від вхідних даних до шумних цільових значень навіть за неточності певних міток [10].

Задля опису впливу шумних слабких міток та уточнення на рівні псевдоміток вводиться формальна модель шуму та отримуються межі поширення помилок для моделі порогового відсічення за впевненістю (confidence thresholding). Припускається, що кожна слабка мітка $\tilde{y}$ складається з однієї істинної мітки y в каналі шуму, залежного від класу, який описується за допомогою $C \times C$ матриці переходів T , тобто [3]:

$$T_{ij} = \Pr(\tilde{y} = j | y = i), \sum_{j=1}^C T_{ij} = 1 \forall i \quad (3)$$

Така модель охоплює як симетричний шум $T_{ij} = 1 - \eta, T_{ji} = \frac{\eta}{C-1}$ для $j \neq i$, так і асиметричний шум, тобто класово-залежний, який передбачає систематичне змішування певних класів. Згідно з цією моделлю, емпіричний ризик на основі зашумлених міток $R_{шум}(f) = E_{(x,y) \sim D} [\ell(f(x), \tilde{y})]$ залежить від «чистого» ризику $R_{чист}(f) = E[\ell(f(x), y)]$ таким чином:

$$R_{шум}(f) = \sum_{i=1}^C \sum_{j=1}^C T_{ij} E[\ell(f(x), j) | y = i] \quad (4)$$

Якщо матриця T є повноранговою, може застосовуватися T^{-1} , щоб отримати за допомогою зважування за значущістю (importance weighting) незміщену оцінку $R_{чист}(f)$:

$$\hat{R}_{чист}(f) = \frac{1}{N} \sum_{k=1}^N \sum_{i=1}^C [T^{-1}]_{\tilde{y}_k, i} \ell(f(x_k), i) \quad (5)$$

З метою уникнення залежності від певної повнорангової матриці переходів T , для кожного класу c передбачається принаймні одна якірна множина (anchor set) [8]:

$$A_c = \{x : p(y = c | x) = 1\} \quad (6)$$

За таких умов, діагональний елемент T_{cc} може оцінюватися безпосередньо за допомогою:

$$\hat{T}_{cc} = \Pr(\tilde{y} = c | x \in A_c) \approx \frac{1}{|A_c|} \sum_{x \in A_c} 1_{\{\tilde{y}(x) = c\}} \quad (7)$$

Недіагональні елементи T_{ij} потім отримуються шляхом нормування по кожному рядку так, щоб виконувалася умова $\sum_j \hat{T}_{ij} = 1$.

Оскільки дійсні якірні множини можуть бути невідомими, вони можуть бути наближено обрані як сукупність зразків з високою впевненістю моделі, наприклад:

$$\{x : p_{\theta}(c | x) \geq \tau_{якір}\} \quad (8)$$

де $\tau_{якір}$ є близьким 1. Таким чином створюється послідовний оцінювач T , чия похибка зникає з підвищенням порогового значення.

Навіть якщо T є оціненою, вона може бути погано обумовленою (ill-conditioned). Внаслідок цього, прямо інверсія замінюється регуляризованою за Тихоновим псевдоінверсією [9]:

$$\hat{T}_{\lambda}^{-1} = (\hat{T}^{\top} \hat{T} + \lambda I)^{-1} \hat{T}^{\top} \quad (9)$$

при чому $\lambda > 0$ визначається за допомогою перехресного затвердження (cross validation). Таким чином відбувається зменшення посилення оціночного шуму в напрямку малих сингулярних значень та отримується чисельно стабільна та незміщена оцінка ризику з похибкою порядку $O(\lambda)$ [7].

В наступній фазі базова модель f_θ вбудовується у цикл активного навчання задля точного знаходження точок даних, за яких мітки потребують втручання людини. Спочатку, для кожного зразка x_i розраховується міра неточності в межах моделі. У завданнях класифікації для цього застосовується згаданий розподіл $p_\theta(y=c|x_i)$ для всіх класів c . Враховуючи цей розподіл обчислюється ентропія:

$$H(x_i) = -\sum_c p_\theta(y=c|x_i) \log p_\theta(y=c|x_i) \quad (10)$$

яка показує, наскільки широко модель розподіляє свої ймовірності між класами.

До того ж, визначається показник найменшої впевненості (least-confident Score):

$$LC(x_i) = 1 - \max_c p_\theta(y=c|x_i) \quad (11)$$

що визначає, наскільки низькою є впевненість класів ймовірності [4].

Після цього здійснюється сортування слабо помічених зразків за спаданням відповідно до цих обох показників. Найвищу позицію у цьому сортуванні займають дані, за яких модель є найбільш невпевненою.

Якщо призначити псевдомітку $\hat{y} = \operatorname{argmax}_c p_\theta(c|x)$ тільки тоді, коли впевненість моделі дорівнює $C(x) = \max_c p_\theta(c|x) \geq \tau$, тоді підпадають зразки з низьким оціненим рівнем шуму. Таким чином:

$$S_\tau = \{x : C(x) \geq \tau\} \quad (12)$$

що є підмножиною зразків x , для яких впевненість моделі $C(x)$ є не меншою за заданий поріг τ .

У межах класово залежної шумової моделі для кожного зразку $x \in S_\tau$, тобто такого, для якого впевненість моделі $C(x) \geq \tau$, ймовірність помилки псевдомітки (тобто того, що $\hat{y} \neq y$) обмежується зверху наступною оцінкою:

$$\Pr(\hat{y} \neq y | x \in S_\tau) \leq \frac{1-\tau}{\tau} \max_{i \neq j} T_{ij} \quad (13)$$

Ця нерівність показує, що ймовірність неправильного класифікаційного передбачення зменшується лінійно зі зростанням впевненості моделі τ , навіть у присутності шуму, описаного матрицею переходів T .

Внаслідок цього, додаткова частка ризику, яка вводиться за посередництва передбачення тільки найбільш впевнених псевдоміток, обмежується за допомогою

$$R(f_{S_\tau}) - R_{\text{ист}}(f_{S_\tau}) \leq (1-\tau) \max_{i \neq j} T_{ij} \quad (14)$$

де якщо $\tau \rightarrow 1$, шум у псевдомітках зникає.

Загалом ці порогові значення показують, що щонайбільше $(1-\tau) \max_{i \neq j} T_{ij}$ псевдоміток є хибними; та що в результаті зважуваного за значущістю донавчання отримується незміщена оцінка дійсного ризику.

Хоча порогове відсічення за рівнем впевненості обмежує похибку у випадку з псевдомітками, зразки з найменшим рівнем впевненості також мають надсилатися до зовнішнього анотатора, тобто «оракула». Під час кожної анотації циклу активного навчання обираються ті вхідні дані x_i , чиї показники неточності знаходяться за межами встановленого порогу яким потім оракул призначає нові мітки.

Оракул моделюється як анотатора з шумом, що характеризується клас-залежною матрицею помилок:

$$O_{ij} = \Pr(y_{\text{оракул}} = j | y_{\text{дійсн}} = i), \sum_j O_{ij} = 1 \quad (15)$$

Така матриця оцінюється за допомогою малого відкладеного калібрувального набору даних, тобто оракулу надаються зразки з відомою міткою а його помилки – фіксуються. Під час уточнення, кожна підтверджена псевдомітка додаткового зважується відповідно до надійності оракула для цього класу. Таким чином міткам з класів, у випадку з якими оракул часто допускає помилок, призначаються менші вагові коефіцієнти.

Оскільки кожен запит до оракула спричиняє витрати c_i та затримку ℓ_i , визначаються загальний бюджет B і допустима максимальна затримка L . Тому вибір множини запитів формулюється як задача оптимізації:

$$\max_{S \subseteq \{x_i\}} \sum_{x_i \in S} \Delta_i \text{ за умов } \sum_{x_i \in S} c_i \leq B, \max_{S \subseteq \{x_i\}} \ell_i \leq L \quad (16)$$

де Δ_i є очікуваним зменшенням узагальненої помилки внаслідок запиту зразку x_i , що оцінюється за допомогою функцій впливу, або невизначеності урахуванням різноманіття.

Якщо індекс i обирається на ітерації t , його вихідна слабка мітка $\tilde{y}_i$ замінюється новою отриманою міткою та множина навчання оновлюється відповідним чином:

$$D_{t+1} = \left\{ (x_j, \tilde{y}_i) \right\}_{j \in S_t} \cup \left\{ (x_i, y_i^{нов}) \right\}_{i \in S_t} \quad (17)$$

де S_t відповідає множині індексів, запитуваних на часовому кроці t .

Після повторного призначення міток, знову проводиться навчання f_θ на D_{t+1} шляхом мінімізації:

$$\hat{\theta}_{t+1} = \operatorname{argmin}_\theta \frac{1}{|D_{t+1}|} \sum_{(x,y) \in D_{t+1}} \ell(f_\theta(x), y) \quad (18)$$

що поширює експертні або відповідні правилам уточнення по загальному середовищу прийняття рішень моделі.

В межах циклу уточнення з участю людини, кожна слабка мітка не піддається переписуванню новою отриманою від оракула міткою. На практиці, цей процес переписування міток уточнюється шляхом застосування механізму селективного виправлення в ході якого вирішується, чи довіряти мітці, виправляти її чи поєднувати з альтернативним варіантом в залежності від впевненості моделі та схвалення оракула. Після запитування оракула на предмет зразка x_i , оцінений рівень впевненості моделі $C_i = \max_c p_\theta(y = \# x_i)$ порівнюється відносно порогового значення $\tau \in (0, 1)$. Якщо $C_i \geq \tau$, а мітка оракула $y_i^{нов}$ співпадає з вихідною слабкою міткою, остання приймається та розглядається в якості достатньо надійної, згодом не піддаючись зміні. У випадках коли $C_i < \tau$, або якщо судження оракула протирічить слабкому анотуванню, похібний сигнал був достатньо шумним щоб потребувати виправлення.

Крім то, коли як слабка мітка, так і мітка оракула містять взаємодоповнювальну інформацію, тобто, якщо слабкий анотатор є точним для широкої підожини зразків, але систематично помиляється на окремих граничних випадках, ці два сигнали об'єднуються з врахуванням зважування за впевненістю. Позначаючи через $\alpha \in [0, 1]$ ваговий коефіцієнт, призначений початковій слабкій мітці, обчислюється фінальна цільова мітка для навчання:

$$y_i^{фін} = \alpha \tilde{y}_i + (1 - \alpha) y_i^{нов} \quad (19)$$

де α може бути сама функцією C_i . Таким чином, якщо $\alpha = \frac{C_i}{C_i + (1 - C_i)}$, тоді до слабкої мітки при-

значається більше довіри за умови впевненості моделі, а вага змінюється до $y_i^{нов}$ при зменшенні рівня впевненості. Таким чином, механізм уточнення плавно переходить між повністю правилами на базі міток та виправленнями оракула. Після завершення кожної ітерації модель перенавчається на вибірково скоригованих цільових мітках, і таким чином лише справді сумнівні мітки змінюються і достовірність тих зразків зберігається, які вже вважаються надійними.

Після включення скоригованих цільових міток до навчального набору даних модель одразу перенавчається, що таким чином замикає цикл «навчання → запит → корекція → перенавчання». Оновлена оцінка параметрів формується у вигляді:

$$\theta_{t+1} = \operatorname{argmin}_\theta \frac{1}{|D_t|} \sum_{(x_i, y_i) \in D_t} \ell(f_\theta(x_i), y_i) \quad (20)$$

що мінімізує ту саму функцію втрат ℓ , що й на початку. Оскільки кожна щойно виправлена мітка стосується зразку з максимальною невизначеністю, її включення непропорційно сильно впливає на здатність моделі вирішувати подібні неточності в інших зразках. Практика показує, що вже після кількох ітерацій межа прийняття рішень моделі істотно стабілізується, а кількість зразків з високою ентропією в пулі невідомих даних зменшується.

Цикл продовжується доти, доки він не навмисно не буде зупинений. Використовуються два критерії: по-перше, значенням T_{max} обмежується загальна кількість запитів до оракула, щоб уникнути необмежених витрат на анотування; по-друге, спостерігається або функція втрат на валідаційній множині $L_{val}(\theta_t)$ або внутрішня міра якості міток $Q(D_t)$. Таким чином, цикл зупиняється, коли виконується хоча б одна з умов:

$$L_{val}(\theta_{t+1}) - L_{val}(\theta_t) < \varepsilon \text{ або } Q(D_{t+1}) - Q(D_t) < \delta \quad (21)$$

де ε та δ є малими граничними значеннями. У цей момент подальші виправлення міток дають усе менший ефект.

Щоб збалансувати швидкість анотування та якість міток, застосовується політика з подвійним порогом:

– високий поріг впливу τ_H : приклади з невизначеністю вище цього порогу одразу надсилаються до оракула;

– низький поріг впливу τL : приклади з невизначеністю між τL та τH накопичуються для періодичних пакетних запитів.

Цей підхід дозволяє обробляти найважливіші приклади в реальному часі, не перевантажуючи при цьому анотатора менш критичними випадками.

Розглядаючи виправлені мітки як часткову процедуру знешумлення у межах PAC-підходу до навчання (Probably Approximately Correct learning – ймовірісно приблизно коректне навчання), можна показати, що кожне виправлення, підтверджене оракулом, зменшує ефективний рівень шуму η у навчальному наборі, тим самим звужуючи стандартну межу узагальнення. Нехай H позначає гіпотетичний клас моделей із VC-розмірністю (Валника – Червоненкіса) d . Якщо початковий рівень шуму слабких міток дорівнює η_0 , тоді після m селективних виправлень залишковий рівень шуму η_m задовольняє умову:

$$\eta_m \leq \eta_0 \left(1 - \frac{N_m}{N} \right) \quad (22)$$

а з імовірністю не менше $1 - \delta$ узагальнена похибка вдосконаленої моделі f_{θ_m} обмежується так:

$$R(f_{\theta_m}) \leq \hat{R}^{(m)}(f_{\theta_m}) + O(N d \log \left(\frac{\sqrt{\eta_m} + \log(1/\delta)}{N} \right)) \quad (23)$$

де $R^{(m)}$ є емпіричним ризиком на частково виправленій вибірці, а N – її розміром. Таким чином, кожне виправлення звужує границю, зменшуючи η_m , якщо N є досить великим, щоб контролювати дисперсію.

Спираючись на результати з теорії складності вибірки для активного навчання, можна обмежити кількість запитів до оракула, необхідну для досягнення допустимої надлишкової помилки ϵ . Якщо гіпотетичний клас H має коефіцієнт розбіжності θ щодо розподілу зі слабким поміченням, тоді достатньо $O(\theta d \log(1/\epsilon))$ запитів, щоб зменшити загальну похибку до рівня, близького до ідеального безшумного випадку. На практиці θ відображає, як швидко невпевненість моделі концентрується навколо помилково помічених прикладів; менше θ передбачає меншу необхідну кількість запитів.

Для оцінки методу в роботі у реальних умовах був проведений експеримент на наборі даних IMDb Movie Reviews, що містить 50 000 текстових відгуків про фільми, розподілених порівну між двома класами: позитивний і негативний. Для тренування було використано 25 000 прикладів, ще 5 000 – для валідації, і 20 000 – для тестування.

Були побудовані чотири слабкі функції ($\lambda_1 - \lambda_4$), що базуються на евристичних припущеннях, а саме:

- λ_1 – присутність позитивних слів з фіксованого словника;
- λ_2 – частота негативної лексики;
- λ_3 – емоційна забарвленість заголовку відгуку;
- λ_4 – оцінка на IMDb (у разі наявності).

Генеративна модель навчалася з урахуванням залежностей між λ -функціями через регуляризовану матрицю зв'язків. Значення гіперпараметра α було встановлено на 100. Дискримінантну модель реалізовано на базі DistilBERT, тобто попередньо натренованої трансформерної архітектури, адаптованої до завдання двокласової класифікації.

На (рис. 1) показано зміну F1-значення протягом 100 ітерацій активного уточнення міток. Початкова продуктивність моделі після лише слабого маркування становила 0.74. Після 10 активних ітерацій точність зросла до 0.87, а після 25 – до 0.91.

Для порівняння були реалізовані чисте активне навчання на базі логістичної регресії без слабких міток; та метод, у якому слабкі мітки замінюються на істинні після уточнення. Метод активного навчання без слабких міток досяг F1 = 0.91 після 60 ітерацій, а другий підхід – після 80. У порівнянні, запропонований у цьому дослідженні підхід досягає аналогічного рівня точності в 2–3 рази швидше завдяки використанню слабких правил як стартової точки.

Висновки. Було запропоновано ітеративний метод, який використовує активне навчання для виявлення та корекції помилкових міток у слабо анотованому наборі даних. Вихідним є набір даних, у якому приклади мають лише зашумлені або частково надані анотації. Запропонований підхід поступово уточнює навчальну вибірку шляхом ідентифікації зразків, які з високою ймовірністю містять помилкові помічення, та передбачає залучення цілеспрямованої перевірки, або з боку людини, або за допомогою евристичних правил із високим рівнем довіри. На кожній ітерації модель навчається на поточному наборі міток і перевіряється на узгодженість своїх передбачень: оцінюється зміна вихідних прогнозів при варіаціях вхідних даних або протягом епох навчання, при цьому висока варіативність трактується як ознака невизначеності мітки. Замість повної переанотації створюються машинно згенеровані

псевдомітки для сумнівних прикладів і подаються на швидке підтвердження чи виправлення анотаторам. Після інтеграції уточнених міток модель перенавчається, що поступово вдосконалює функцію маркування та дозволяє поетапно усувати найбільш критичні помилки. Протягом кількох циклів ця гібридна стратегія, що поєднує вибір на основі розбіжностей та вибірку людську валідацію, поступово підвищує якість анотацій без необхідності повного перепомічення всього набору даних.

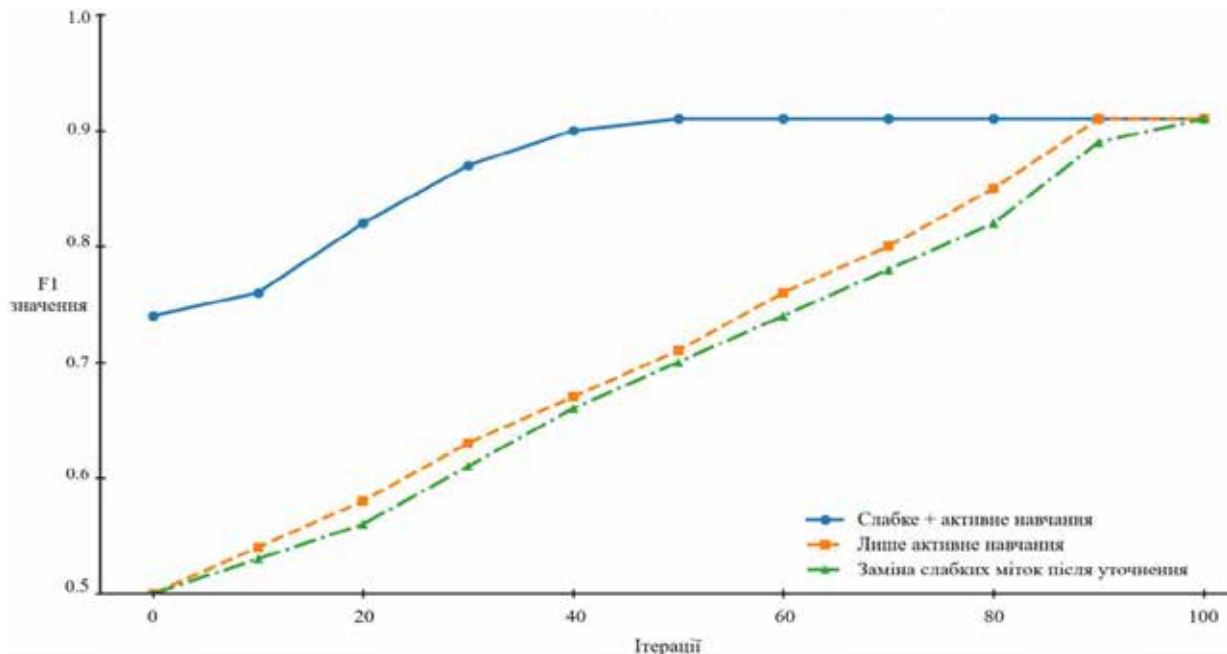


Рис. 1. Порівнянні методів з огляду на F1-значення

Список використаних джерел:

1. Active Learning with Weak Supervision for Gaussian Processes. Olmin A., Lindqvist J., Svensson L., Lindsten F. *ArXiv* : website. 2024. DOI: <https://doi.org/10.48550/arXiv.2204.08335>
2. Active WeaSuL: improving weak supervision with active learning. Biegel S., El-Khatib R., Vilas Boas L. Oliveira O., Baak M., Aben N. *ArXiv* : website. 2021. DOI: <https://doi.org/10.48550/arXiv.2104.14847>
3. Adaptive Confidence Thresholding for Monocular Depth Estimation. Choi H., Lee H., Kim S., Kim S., Kim S., Sohn K., Min D. *ArXiv* : website. 2021. DOI: <https://doi.org/10.48550/arXiv.2009.12840>
4. Agrawal A., Tripathi S., Vardhan M. Active Learning Approach Using a Modified Least Confidence Sampling Strategy for Named Entity Recognition. *Progress in Artificial Intelligence*. 2021. Vol. 10. DOI:10.1007/s13748-021-00230-w
5. ALWOD: Active Learning for Weakly-Supervised Object Detection. Wang Y., Ilic V., Li J., Kisacanin B., Pavlovic V. *ArXiv* : website. 2023. DOI: <https://doi.org/10.48550/arXiv.2309.07914>
6. Auto-generating weak labels for real & synthetic data to improve label-scarce medical image segmentation. Deshpande T., Prakash E., Ross E. G., Langlotz C., Ng A., Valanarasu J. M. J. *ArXiv* : website. 2024. DOI: <https://doi.org/10.48550/arXiv.2404.17033>
7. Cross-Validation Strategy Impacts the Performance and Interpretation of Machine Learning Models. Sweet L.-B., Müller C., Anand M., Zscheischler J. *Artificial Intelligence for the Earth Systems*. 2023. Vol. 2. P. 1–35. DOI:10.1175/AIES-D-23-0026.1
8. Lesci P., Vlachos A. AnchorAL: Computationally Efficient Active Learning for Large and Imbalanced Datasets. *ArXiv* : website. 2024. DOI: <https://doi.org/10.18653/v1/2024.naacl-long.467>
9. Vincent E. Tikhonov regularization approach to solving inverse problems for parameter learning: master's thesis. African Institute for Mathematical Sciences (AIMS), Cameroon; scientific supervisor: Dr Floriane Melo Kue. Cameroon, 2022. 30 May. 34 p. DOI:10.13140/RG.2.2.21667.43043
10. Warm Start Active Learning with Proxy Labels & Selection via Semi-Supervised Fine-Tuning / Nath V., Yang D., Roth H. R., Xu D. *ArXiv* : website. 2022. DOI: <https://doi.org/10.48550/arXiv.2209.06285>

Дата надходження статті: 10.06.2025

Дата прийняття статті: 17.06.2025

Опубліковано: 23.09.2025

УДК 004.42:37.018.43

DOI <https://doi.org/10.32689/maup.it.2025.2.19>

Артем МОЙСЕЄНКО

аспірант кафедри інженерії програмного забезпечення,

Національний аерокосмічний університет «Харківський авіаційний інститут», a.m.moiseienko@khai.edu

ORCID: 0009-0001-7852-4993

Юлія КУЗНЕЦОВА

кандидат технічних наук, доцент,

доцент кафедри інженерії програмного забезпечення,

Національний аерокосмічний університет «Харківський авіаційний інститут», y.kuznetsova@khai.edu

ORCID: 0000-0001-9416-6981

**АНАЛІЗ ПРОБЛЕМНИХ ПИТАНЬ ВПРОВАДЖЕННЯ СУЧАСНИХ ХМАРНИХ ТЕХНОЛОГІЙ
У СИСТЕМИ УПРАВЛІННЯ НАВЧАННЯМ**

Анотація. Стаття присвячена огляду поточного стану розвитку СУН та використання хмарних технологій у цій сфері. Проаналізовано глобальні тенденції росту ринку СУН і впровадження хмарних сервісів, зокрема прискорення цифрової трансформації освіти під впливом пандемії та воєнних умов. Визначено основні переваги використання хмарних обчислень в освітніх платформах, а також окреслено проблемні аспекти та виклики, що виникають при впровадженні СУН.

Мета. Узагальнити сучасні підходи та рішення у галузі СУН та хмарних технологій, виявити актуальні проблеми їх впровадження й експлуатації, особливо у контексті оптимального розгортання сервісів СУН в хмарному середовищі.

Методологія. Застосовано методи аналізу та синтезу наукових публікацій останніх років, статистичних даних ринку освітніх технологій і досвіду впровадження СУН в Україні та світі. Виконано порівняльний аналіз можливостей традиційних та хмарних СУН-рішень, а також контент-аналіз досліджень щодо продуктивності, безпеки та архітектури цих систем.

Наукова новизна. Уточнено сучасний стан розвитку СУН з урахуванням хмарних трендів, систематизовано актуальні виклики (від відсутності уніфікованих методик оцінювання ефективності СУН до питань безпеки та оптимального використання ресурсів). Вперше виокремлено проблему оптимального розгортання сервісів СУН у хмарній інфраструктурі як одну з ключових для підвищення ефективності e-learning.

Висновки. Сучасні СУН активно зростають за популярністю та функціональністю, інтегруючи хмарні сервіси для масштабованості й доступності. Водночас існують недоліки, зокрема труднощі з вибором і впровадженням оптимальної платформи, забезпеченням безпеки й продуктивності при зростанні навантаження. Аналіз показав, що для вирішення цих проблем перспективним є застосування математичних моделей оптимізації. Подальше дослідження буде спрямоване на розроблення ILP-моделі оптимального розгортання сервісів СУН у хмарному середовищі як інструменту для підвищення ефективності та надійності електронних навчальних систем.

Ключові слова: система управління навчанням, СУН, електронне навчання, хмарні технології, дистанційна освіта, оптимізація розгортання, хмарна інфраструктура.

**Artem MOISEIENKO. Yuliia KUZNETSOVA. ANALYSIS OF IMPLEMENTATION CHALLENGES OF MODERN
CLOUD TECHNOLOGIES IN LEARNING MANAGEMENT SYSTEMS**

Abstract. This article provides a review of the current state of LMS development and the use of cloud technologies in this domain. Global trends in the growth of the LMS market and the adoption of cloud services are analyzed, including the acceleration of digital transformation in education under the influence of the pandemic and wartime conditions. The key advantages of using cloud computing in educational platforms are identified, as well as the problematic aspects and challenges that arise in the implementation of LMS.

The aim of study is to summarize modern approaches and solutions in the field of LMS and cloud technologies, and to identify current issues in their implementation and operation, especially in the context of optimal deployment of LMS services in cloud environments.

Methodology. The study employs analysis and synthesis of recent scientific publications, statistical data on the educational technology market, and the experience of LMS deployment in Ukraine and globally. A comparative analysis of the capabilities of traditional versus cloud-based LMS solutions was conducted, along with content analysis of research on performance, security, and architecture of these systems.

Scientific novelty. The current state of LMS development is clarified with regard to cloud trends; the prevailing challenges are systematized (ranging from the lack of unified methods for evaluating LMS effectiveness to issues of security and optimal resource utilization). The problem of optimal deployment of LMS services in cloud infrastructure is highlighted for the first time as one of the key issues for increasing the efficiency of e-learning.

© A. Моїсеєнко, Ю. Кузнецова, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

Conclusions. Modern LMS are rapidly expanding in popularity and functionality, integrating cloud services for scalability and accessibility. At the same time, there are drawbacks, including difficulties in selecting and implementing an optimal platform, and ensuring security and performance under increasing loads. The analysis shows that applying mathematical optimization models is a promising approach to solving these issues. Further research will focus on developing an ILP model for optimal deployment of LMS services in a cloud environment as a tool to enhance the efficiency and reliability of electronic learning systems.

Key words: learning management system, LMS, e-learning, cloud technologies, distance education, deployment optimization, cloud infrastructure.

Постановка проблеми. Інформатизація освіти та бізнесу в останні роки супроводжується стрімким зростанням використання систем управління навчанням (СУН). Ці платформи дозволяють створювати єдиний цифровий навчальний простір, де здійснюється розробка та поширення освітніх матеріалів, управління навчальним процесом, моніторинг успішності тощо. Попит на СУН неухильно зростає в глобальному масштабі. Так, за оцінками аналітиків, обсяг європейського ринку СУН у 2025 р. сягне близько 395,9 млрд грн (10,7 млрд доларів США) при середньорічному темпі зростання 19,2 % у період 2025–2034 рр. [9]. Понад 83 % компаній впроваджують СУН для навчання та підвищення кваліфікації співробітників, що підкреслює універсальність цих систем не лише в освіті, а й у корпоративному секторі [6]. Додатковим стимулом до масового впровадження цифрових платформ стали пандемія COVID-19 та вимушений перехід на дистанційний формат у 2020 році, коли багато закладів освіти по всьому світу були змушені оперативно адаптувати навчання під онлайн-режим [10]. В Україні критичну необхідність надійних СУН загострили наслідки війни: лише станом на липень 2023 року було зруйновано понад 330 і пошкоджено понад 3200 закладів освіти [4, с. 4], що змусило шукати стійкі дистанційні рішення. Отже, сучасні реалії підтверджують актуальність розвитку СУН та переходу до хмарних технологій, які забезпечують безперервність навчання навіть у надзвичайних умовах.

Водночас зростання використання СУН виявило низку проблем. Необхідно не лише впровадити електронну систему навчання, але й забезпечити її ефективну роботу та масштабованість. Багато традиційних СУН розгорнуті як монолітні програмні комплекси на локальних серверах навчальних закладів. Такий підхід може супроводжуватися труднощами у підтримці та оновленні системи, обмеженнями продуктивності при збільшенні кількості користувачів, а також значними витратами на інфраструктуру. Постає проблема оптимального розгортання сервісів СУН: як організувати компоненти системи на доступних обчислювальних ресурсах (зокрема хмарних), щоб досягти найкращого співвідношення між вартістю, швидкістю та надійністю роботи. Вирішення цього завдання має вагомий практичний значення для закладів освіти, що прагнуть максимально ефективно використовувати сучасні цифрові технології в навчальному процесі.

Аналіз останніх досліджень і публікацій. Питання розвитку та впровадження СУН активно досліджується як зарубіжними, так і вітчизняними науковцями. Огляд наукової літератури свідчить, що за останні роки сформувалися різні підходи до використання СУН в освіті залежно від умов та потреб. У розвинених країнах СУН стали невід'ємним інструментом навчального процесу, сприяючи інтерактивності й персоналізації навчання [6]. Дослідження відзначають, що найпоширенішою платформою є Moodle – відкрита СУН, яка завдяки гнучкості набула глобального визнання [6]. Станом на червень 2025 року за даними офіційної статистики на сайті кількість зареєстрованих користувачів Moodle перевищує 455 млн осіб (проти 294 млн у 2021 р.), що підтверджує понад майже двократне зростання аудиторії за останні роки [10]. Такі системи, як Blackboard та Google Classroom, також належать до найбільш уживаних у світі, проте використовуються переважно у корпоративному сегменті та для окремих освітніх сценаріїв [6]. Водночас існує ціла низка інших платформ, популярних у певних регіонах (Canvas, Schoology тощо), що відображає різноманіття ринку СУН [6].

Науковці приділяють увагу як педагогічним аспектам використання СУН, так і технічним питанням їх удосконалення. Значна кількість праць присвячена оцінці впливу СУН на успішність та залученість студентів, порівнянню ефективності різних електронних середовищ навчання. Наприклад, узагальнено позитивний вплив використання СУН на мотивацію та результати навчання студентів STEM-напрямів [10]. Інша група досліджень фокусується на проблемах впровадження СУН у країнах, що розвиваються, де виокремлюються питання матеріально-технічного забезпечення та підготовки кадрів. Зокрема, дослідження відзначають, що впровадження СУН у таких країнах часто стикається з викликами щодо інфраструктури та доступу до інтернету [6]. Окремо досліджуються питання безпеки та збереження даних у хмарних платформах навчання. З ростом використання хмарних рішень виникла потреба гарантувати конфіденційність навчальної інформації та захист персональних даних користувачів. У 2023 році запропоновано низку удосконалень криптографічних засобів для підвищення безпеки хмарних СУН [8]. Наприклад, розроблено покращений протокол шифрування (Enhanced Schmidt-Samoa), орієнтований

на системи електронного навчання, що дозволяє підвищити стійкість до несанкціонованого доступу [8]. Таким чином, питання кібербезпеки є одним із пріоритетних у сучасних дослідженнях СУН.

Технічна еволюція СУН пов'язана і з архітектурними змінами. Традиційні монолітні застосунки поступово трансформуються на основі мікросервісного підходу. Це відображено в роботах, де запропоновано впроваджувати мікросервісну архітектуру в СУН для спрощення підтримки та масштабування системи [11]. Згідно з досвідом розробників, розбиття СУН на дрібні веб-сервіси покращує гнучкість оновлень і дозволяє різним модулям (курси, тести, журнали тощо) розгортатися та масштабуватися незалежно один від одного [11]. Такий підхід також сприяє більш ефективному використанню хмарних ресурсів, оскільки кожен сервіс може бути розміщений на окремих віртуальних машинах або контейнерах відповідно до навантаження.

В українському освітньому просторі тематика СУН та дистанційного навчання набуває особливої актуальності в останні роки. Низка робіт присвячена вибору та впровадженню цифрових платформ у закладах середньої та вищої освіти України. Так, у дослідженні [3] проаналізовано критерії добору платформ електронного навчання для шкіл в умовах пандемії та воєнного стану. Результати опитувань педагогів свідчать про цікаву тенденцію: якщо у 2020 році надзвичайно популярним інструментом був Google Classroom, то вже у 2021 році спостерігалось зменшення частки його користувачів (~ на 9 %) та перехід до вітчизняних розробок – систем Human, «Єдина школа» та інших (сумарне зростання ~ на 13 %) [3, с. 24–25]. Це пояснюється як прагненням до більшої автономності і врахування локальних потреб, так і розширенням функціоналу українських СУН. У закладах вищої освіти України домінує СУН Moodle, яка впроваджена в більшості університетів. В умовах нестабільного інтернет-зв'язку (особливо під час війни) корисним виявилось використання мобільного додатку Moodle з офлайн-режимом роботи: завдяки цьому студенти можуть завантажувати матеріали та виконувати завдання без постійного підключення до мережі [6]. Таким чином, навіть на національному рівні спостерігаємо тенденцію до адаптації СУН під специфічні умови (відсутність стійкого інтернету, вимоги локального законодавства щодо зберігання даних тощо) та активний пошук оптимальних рішень.

Формулювання мети статті. Проаналізувавши стан досліджень, можна констатувати наявність кількох не вирішених питань, пов'язаних із впровадженням СУН і хмарних технологій. Таким чином, метою даної статті є узагальнення сучасного досвіду розробки та використання систем управління навчанням, визначення основних переваг і недоліків інтеграції СУН із хмарними сервісами, а також виокремлення проблеми оптимального розгортання СУН-сервісів. Досягнення поставленої мети передбачає вирішення таких завдань: (1) проаналізувати ключові тенденції розвитку СУН у світі та в Україні; (2) дослідити вплив хмарних технологій на ефективність і масштабованість СУН; (3) визначити основні проблеми, з якими стикаються організації при впровадженні та експлуатації СУН, включно з аспектами продуктивності, безпеки та витрат; (4) обґрунтувати необхідність і напрямки подальших досліджень для оптимізації розгортання СУН у хмарній інфраструктурі.

Виклад основного матеріалу дослідження. Матеріали та методи дослідження. Для досягнення поставленої мети було здійснено огляд наукових статей, аналітичних звітів та статистичних матеріалів. Використано метод контент-аналізу літератури з тематики електронного навчання та хмарних обчислень. Зокрема, проаналізовано дані ринку СУН (розмір, темпи росту) та результати впровадження цих систем у закладах освіти. Порівняльним аналізом зіставлено можливості різних платформ (відкритих і комерційних, хмарних і локальних). Методом систематизації узагальнено відомі проблеми та виклики, описані в дослідженнях.

Результати дослідження. Системи управління навчанням забезпечують централізоване керування освітнім процесом у цифровому середовищі. Базовий функціонал типової СУН включає засоби публікації навчальних матеріалів, проведення онлайн-занять і тестувань, облік успішності, комунікацію між учасниками навчання тощо. У поєднанні з хмарними технологіями такі системи отримують низку додаткових переваг. Хмарна інфраструктура (моделі SaaS, PaaS) дозволяє розгорнути СУН без необхідності придбання та підтримки власних серверів, що особливо вигідно для закладів із обмеженими ІТ-ресурсами. За рахунок розподілених дата-центрів забезпечується високий рівень доступності: користувачі можуть отримати доступ до навчальної платформи з будь-якого місця та пристрою, достатньо підключення до інтернету [6]. Крім того, ефективно спільне використання ресурсів (потужностей серверів, пам'яті, мережі) у хмарі дозволяє одночасно обслуговувати значну кількість користувачів без помітного погіршення швидкодії, а масштабування відбувається динамічно за потреби. Це сприяє економії часу та коштів на обслуговування системи [6]. Наприклад, використання хмарних сервісів зменшує витрати на початкове розгортання СУН, переводячи їх із капітальних (CAPEX) у операційні (OPEX), тобто установи сплачують лише за фактично спожиті ресурси. Усе це робить хмарні платформи привабливими для впровадження сучасного e-learning.

Водночас інтеграція СУН із хмарними технологіями породжує і нові виклики. Одним із головних є забезпечення інформаційної безпеки та конфіденційності даних. Використання зовнішніх хмарних сервісів означає, що чутлива навчальна інформація (дані студентів, результати, навчальний контент) передається та зберігається на сторонніх серверах. Це вимагає впровадження надійних механізмів шифрування і контролю доступу. Як показано у дослідженнях, присвячених безпеці хмарних СУН, існує потреба в удосконаленні криптографічних протоколів саме під специфіку освітніх систем [8]. Крім технічних заходів, важливими є також юридичні аспекти – відповідність зберігання даних вимогам законодавства про захист персональних даних, що може обмежувати використання зарубіжних хмарних рішень (особливо це актуально для державних навчальних закладів України).

Ще однією проблемою є залежність якості роботи СУН від якості інтернет-з'єднання. Для віддалених регіонів або в умовах надзвичайних ситуацій (відключення електропостачання, перебої зв'язку) навіть найкраща хмарна платформа може стати тимчасово недоступною. Тому розробники впроваджують рішення для автономної роботи: офлайн-режими, синхронізацію даних при відновленні зв'язку тощо [6]. З огляду на це, при виборі СУН користувачі часто звертають увагу не лише на функціонал, а й на можливість працювати офлайн або розгортатися у гібридній архітектурі (частково на власних серверах, частково в хмарі).

Аналіз останніх публікацій також виявив іншу перешкоду ефективного використання СУН через відсутність уніфікованих критеріїв оцінювання їхньої результативності та методичних рекомендацій щодо впровадження. Існування великої кількості різних платформ ускладнює вибір оптимального рішення для конкретної організації [6]. Дослідники відзначають брак стандартизованих методик, які б дозволяли порівнювати СУН між собою за показниками успішності навчання, взаємодії користувачів, економічної доцільності тощо [6]. Це призводить до того, що впровадження СУН часто відбувається шляхом проб і помилок, а накопичений досвід важко узагальнити і передати іншим закладам.

Однак, ключовим технічним викликом залишається оптимальне розгортання СУН у хмарному середовищі. Попри широкий вибір хмарних провайдерів та сервісів (Amazon Web Services, Google Cloud, Microsoft Azure тощо), адміністраторам СУН складно визначити, яка конфігурація ресурсів потрібна для їх системи при заданому числі користувачів і активності. Надлишкове резервування ресурсів призведе до зайвих витрат, а недостатнє – до перевантажень і збоїв у роботі системи. Балансування цих аспектів є нетривіальним завданням, що потребує застосування формальних моделей оптимізації. На сьогодні в літературі майже не представлено готових рішень для автоматизованого планування розгортання СУН-сервісів. У наукових працях [7, 12] виділяється кілька перспективних підходів до вирішення цієї проблеми, які можна згрупувати за певними критеріями (таблиця 1).

Таблиця 1

Класифікація підходів до оптимального розгортання СУН у хмарі

Критерій класифікації	Тип підходу	Коротка характеристика
За методом оптимізації	Евристичні методи	Швидке, хоча й не завжди оптимальне рішення, засноване на емпіричних правилах
	Метаевристичні алгоритми	Використовують генетичні алгоритми, алгоритми рою частинок для ефективного пошуку у великому просторі рішень
	Математичні моделі оптимізації (ЦЛП, ЦНЛП)	Дають точні та теоретично обґрунтовані оптимальні рішення, підходять для задач середнього масштабу
За ступенем автоматизації	Ручний підбір	Базується виключно на досвіді адміністратора, трудомісткий і менш ефективний
	Напівавтоматизоване планування	Рекомендуючі системи пропонують оптимальні конфігурації на основі історичних даних і прогнозних моделей
	Повністю автоматизоване розгортання	Динамічно розподіляє ресурси, повністю автоматизує процес адаптації до змін у навантаженні
За рівнем врахування динаміки навантаження	Статичні моделі	Планування ресурсів на основі фіксованих або середніх значень активності користувачів
	Динамічні моделі	Використовують прогнозні моделі (наприклад, машинне навчання) для адаптивного масштабування ресурсів

Згідно наведеного порівняння, можна зробити висновки, що перспективними є підходи, що поєднують переваги математичних моделей і адаптивних алгоритмів прогнозування, дозволяючи знаходити

баланс між вартістю та продуктивністю. Одним з таких перспективних напрямів є використання цілочисельних лінійних програм (ILP-моделей) [7], які надають можливість точного розрахунку оптимальної конфігурації ресурсів з урахуванням заданих обмежень, що і планується реалізувати у подальших дослідженнях.

Основні недоліки та виклики. На основі проведеного аналізу можна виокремити низку основних проблем та недоліків, притаманних сучасним системам управління навчанням, особливо у контексті їх функціонування в хмарному середовищі:

- відсутність стандартів та методик оцінювання. Наразі не існує загально визнаних підходів для об'єктивного порівняння ефективності різних СУН та вимірювання їхнього впливу на якість навчання. Це ускладнює вибір платформи і оцінку результатів її впровадження [6];

- різноманітність та сумісність платформ. Велика кількість доступних СУН (від відкритих до комерційних) призводить до фрагментації: різні заклади використовують різні системи, що створює бар'єри для обміну контентом і спільних курсів. Інтеграція СУН з іншими освітніми сервісами (електронними бібліотеками, реєстрами тощо) також може бути проблемною через відмінності в стандартах;

- технічна підтримка і ресурси. Для розгортання та підтримки СУН потрібні кваліфіковані ІТ-фахівці та фінансові ресурси. У багатьох випадках (особливо для невеликих закладів освіти) спостерігається нестача ресурсів для придбання і підтримки сучасних технологій [1], що призводить до використання застарілих або обмежених за функціоналом систем;

- безпека та конфіденційність. Як зазначалося, захист даних є критичним питанням. Будь-які вразливості в СУН або хмарній інфраструктурі можуть поставити під загрозу великі обсяги персональної інформації. Забезпечення кібербезпеки потребує постійного вдосконалення засобів автентифікації, авторизації, шифрування, моніторингу активності тощо [8];

- масштабованість та продуктивність. Під час пікових навантажень (наприклад, одночасне проходження тестування великою кількістю студентів) СУН може суттєво втрачати швидкодію. Забезпечення належної масштабованості вимагає ретельного планування інфраструктури або використання автоматизованого масштабування в хмарі, що не завжди налаштовано оптимально. Необхідно враховувати баланс навантаження між компонентами системи (база даних, сервер застосунків, сервери файлів тощо);

- оптимальне розгортання сервісів. Більшість існуючих впроваджень СУН у хмарі виконується за спрощеними сценаріями (наприклад, розміщення цілого серверу СУН на одній віртуальній машині). Такий підхід не використовує повною мірою переваг хмарної архітектури. Оптимальне розподілення окремих сервісів СУН (модуль тестування, модуль форумів, сховище контенту тощо) між різними ресурсами могло б підвищити стійкість та ефективність системи, але це завдання потребує рішення задачі оптимізації. Це являє собою науково-технічний виклик, який досі мало висвітлений у літературі та наразі відсутні інструменти, що автоматично підказують, як краще розгорнути СУН у хмарі з урахуванням множини обмежень (вартість, кількість користувачів, географія, надійність). Виявлена прогалина підтверджує необхідність дослідження методів оптимізації, спроможних вирішити цю задачу.

Висновки з даного дослідження і перспективи подальших досліджень. Сучасні системи управління навчанням у поєднанні з хмарними технологіями відкривають нові можливості для організації освітнього процесу – від гнучкого доступу до навчальних ресурсів до персоналізації навчання та аналізу даних про успішність. Проведений огляд показав, що за останні роки ринок СУН суттєво зріс, а технологічні рішення еволюціонували в напрямі хмарних сервісів і мікросервісної архітектури. СУН на кшталт Moodle, Canvas, Google Classroom та інших стали звичним інструментом в університетах і компаніях, а їх використання набуло стратегічного характеру для забезпечення безперервної освіти, особливо у кризових ситуаціях. Разом із тим, впровадження СУН супроводжується низкою проблем: від методичних (як ефективно інтегрувати платформу в освітній процес) до технічних (як забезпечити її надійну та масштабовану роботу). Одним з найважливіших практичних завдань є оптимізація розгортання СУН у хмарному середовищі з метою підвищення ефективності використання ресурсів і зниження витрат. Поки що це завдання вирішується переважно вручну або емпірично, що залишає простір для покращень. Перспективою подальших досліджень є розроблення формалізованої моделі та інструментальних засобів для оптимального розподілу компонентів СУН на хмарній інфраструктурі. Запропоновано зосередитися на побудові моделі цілочисельного лінійного програмування (ILP), яка врахує параметри продуктивності сервісів, вимоги до надійності, обмеження вартості хмарних ресурсів та інші критерії. Така ILP-модель дозволить автоматично знаходити найкращий варіант розгортання – зокрема, визначити, скільки і яких віртуальних машин або контейнерів потрібно, як розподілити між ними функціональні модулі СУН, щоб при мінімальних витратах забезпечити заданий рівень якості обслуговування користувачів. Реалізація даного підходу сприятиме вирішенню озвучених у статті проблем і стане наступним кроком у розвитку ефективних хмароорієнтованих систем управління навчанням.

Список використаних джерел:

1. Бабак О. А., Ісак Л. М. Хмарні технології в освіті. *Міжнар. наук. журнал «Грааль науки»*. 2023. № 27 (травень). С. 485–489.
2. Дистанційне навчання замість заочного: зміни для якості. *Освіта.UA*. 2023. URL: <https://osvita.ua/vnz/90312/> (дата звернення: 11.11.2023).
3. Олійник В. В., Грабовський П. П., Коновал О. А. Критерії та показники добору цифрової платформи електронного навчання для закладу загальної середньої освіти. *Інформаційні технології і засоби навчання*. 2022. № 90 (4). С. 19–31. DOI: 10.33407/itlt.v90i4.5010.
4. Освіта і наука України в умовах воєнного стану: інформ.-аналіт. збірник / М-во освіти і науки України, Ін-т освітньої аналітики. Київ, 2023. 64 с.
5. Що таке LMS: типи, можливості та переваги систем управління навчанням. *Cases.media*. 2023. URL: <https://cases.media/article/sho-take-lms-tipi-mozhlivosti-ta-perevagi-sistem-upravlinnya-navchanniam> (дата звернення: 11.06.2025).
6. Aljad R. R. Analysis of Development Trends and Experience of using LMS in Modern Education. *E-Learning Innovations Journal*. 2023. Vol. 1, No. 2. P. 86–104. DOI: 10.57125/ELIJ.2023.09.25.05.
7. Apat H. K., Goswami V., Sahoo B., Barik R. K., Saikia M. J. Fog service placement optimization: a survey of state-of-the-art strategies and techniques. *Computers*. 2025. Vol. 14(3), № 99. P. 1000–1045. DOI: 10.3390/computers14030099.
8. Chatterjee P., Bose R., Banerjee S., Roy S. Enhancing Data Security of Cloud Based LMS. *Wireless Personal Communications*. 2023. Vol. 130, No. 2. P. 1123–1139. DOI: 10.1007/s11277-023-10323-5.
9. Europe Learning Management Systems Market is projected to grow from USD 10.70 Billion in 2025 to USD 51.99 Billion by 2034, exhibiting a CAGR of 19.20 % during 2025. 2034. Secondary Research, Primary Research, MRFR Database and Analyst Review (Market Research Future, 2024). URL: <https://www.marketresearchfuture.com/reports/europe-learning-management-systems-market-21601> (дата звернення: 12.06.2025).
10. Gamage S. H. P. W., Ayres J. R., Behrend M. B. A systematic review on trends in using Moodle for teaching and learning. *International Journal of STEM Education*. 2022. Vol. 9, Article 9. DOI: 10.1186/s40594-021-00323-x.
11. Handoko D., Nur W. Implementation of Microservices Architecture in Learning Management System E-Course Using Web Service Method. *Sinkron: Jurnal dan Penelitian Teknik Informatika*. 2022. Vol. 7, No. 1. P. 76–82. DOI: 10.33395/sinkron.v7i1.11229.
12. Santoyo-González A., Cervelló-Pastor C. Latency-aware cost optimization of the service infrastructure placement in 5G networks. *Journal of Network and Computer Applications*. 2018. Vol. 114. P. 29–37. DOI: 10.1016/j.jnca.2018.04.009.

Дата надходження статті: 24.06.2025

Дата прийняття статті: 30.06.2025

Опубліковано: 23.09.2025

УДК 004.415.5

DOI <https://doi.org/10.32689/maup.it.2025.2.20>**Борис ПАНАСЮК**

аспірант спеціальності «Інженерія програмного забезпечення»,
Вінницький національний технічний університет,
boris.panasjuk@gmail.com
ORCID: 0009-0007-2064-9121

Наталя БАБЮК

кандидат технічних наук, доцент кафедри програмного забезпечення,
Вінницький національний технічний університет,
babiuk@vntu.edu.ua
ORCID: 0000-0003-0607-6340

КОНТРАКТНО-ОРІЄНТОВАНЕ МОДЕЛЮВАННЯ МІКРОСЕРВІСНОЇ АРХІТЕКТУРИ ВИСОКОНАВАНТАЖЕНИХ СИСТЕМ У UML-ПОДІБНОМУ СЕРЕДОВИЩІ

Анотація. Метою дослідження є розроблення та обґрунтування методу контрактно-орієнтованого моделювання мікросервісної архітектури високонавантажених програмних систем у спеціалізованому UML-подібному веб-середовищі. Запропонований підхід поєднує переваги візуального проектування та формальної специфікації інтерфейсів: JSON-фрагменти контрактів інтегруються безпосередньо у стрілки, що відображають взаємодію компонентів на діаграмах. Це забезпечує ранню, однозначну й машинно-читану фіксацію API-угод, мінімізує розрив між дизайнерською моделлю та реалізацією сервісів і знижує витрати на супровід документації.

Методологія. Метод включає алгоритм автоматичного перетворення вбудованих JSON-контрактів у специфікації OpenAPI v3, які далі можуть використовуватися для генерації серверних «заглушок», клієнтських бібліотек, інтеграційних тестів або для налаштування шлюзів API. Реалізовано прототип середовища на стеку TypeScript + Node.js, що поєднує Canvas/WebGL-рендеринг діаграм із серверним модулем генерації OpenAPI. Проведено експериментальне оцінювання на трьох репрезентативних сценаріях (каталог товарів, сервіс замовлень, сервіс платежів), у яких порівнювали запропонований підхід із класичним UML-процесом, що не охоплює контракти. Результати показали скорочення середнього часу узгодження API між командами на 28 %, зменшення кількості дефектів інтеграції на 33 % та покращення індексу покриття документацією на 22 %.

Наукова новизна. Особливою перевагою розробленого інструменту є підтримка «живої» документації: кожна зміна у діаграмі миттєво відбивається в оновленому контракті, що гарантує синхронність між моделлю та вихідним кодом сервісів упродовж усього життєвого циклу ПЗ. Перспективи подальших досліджень включають інтеграцію середовища з системами трасування виконання контрактів у режимі реального часу, розширення підтримки асинхронних протоколів (AsyncAPI) та розробку метрик оцінювання складності контрактних залежностей у масштабах корпоративного портфеля сервісів.

Висновки. Запропонований метод контрактно-орієнтованого моделювання у UML-подібному середовищі підтвердив свою ефективність для високонавантажених мікросервісних систем. Автоматична генерація OpenAPI та підтримка «живої» документації скорочують час узгодження API, зменшують дефекти інтеграції та підвищують покриття документації. Інструмент забезпечує безперервну узгодженість між архітектурною моделлю та кодом сервісів, що підвищує якість ПЗ і спрощує його еволюцію та супровід.

Ключові слова: контракт-орієнтоване моделювання, інформаційна система, мікросервісна архітектура, OpenAPI, UML-подібне середовище, високонавантажені системи, JSON Schema, Swagger.

Borys PANASIUK, Natalia BABIUK. CONTRACT-ORIENTED MODELLING OF MICROSERVICE ARCHITECTURE FOR HIGH-LOAD SYSTEMS IN A UML-LIKE ENVIRONMENT

Abstract. The study aims to develop and substantiate a contract-oriented modelling method for the microservice architecture of high-load software systems within a specialised UML-like web environment. The proposed approach merges the advantages of visual design and formal interface specification: JSON contract fragments are embedded directly into the arrows representing component interactions on the diagrams. This provides an early, unambiguous and machine-readable fixation of API agreements, minimises the gap between the design model and service implementation, and reduces documentation-maintenance costs.

Methodology. The method includes an algorithm that automatically converts the embedded JSON contracts into OpenAPI v3 specifications, which can then be used to generate server stubs, client libraries, integration tests or configure API gateways. A prototype of the environment was implemented with a TypeScript + Node.js stack, combining Canvas/WebGL diagram rendering with a server-side OpenAPI generator. Experimental evaluation on three representative scenarios (product catalogue, order service, payment service) compared the proposed approach with a classical UML-based process that lacks contracts. The results demonstrated a 28 % reduction in average API alignment time, a 33 % decrease in integration defects and a 22 % improvement in documentation coverage.

© Б. Панасюк, Н. Бабюк, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

Scientific novelty. A distinctive advantage of the tool is its support for «living» documentation: any change in the diagram is instantly mirrored in the updated contract, ensuring synchronisation between the model and service source code throughout the software life-cycle. Future research will focus on integrating the environment with real-time contract-tracing systems, extending support for asynchronous protocols (AsyncAPI) and developing metrics for assessing the complexity of contract dependencies across an enterprise-scale service portfolio.

Conclusions. The proposed contract-oriented modelling method within a UML-like environment has proven effective for high-load microservice systems. Automatic OpenAPI generation and live documentation support shorten API alignment time, reduce integration defects and increase documentation coverage. The tool maintains continuous consistency between the architectural model and service code, thereby enhancing software quality and simplifying its evolution and maintenance.

Key words: contract-oriented modelling, microservice architecture, high-load systems, UML-like environment, OpenAPI, JSON Schema, Swagger.

Вступ. Сучасні високонавантажені програмні системи дедалі частіше реалізуються на базі мікросервісної архітектури. Це обумовлено здатністю такого підходу забезпечити високу гнучкість, масштабованість та стійкість до збоїв. Однак проектування цих систем є складним завданням, яке потребує інтегрованих підходів для опису взаємодії між компонентами. Наявність значної кількості взаємозалежних сервісів, які активно комунікують між собою, породжує складнощі у специфікації інтерфейсів, документуванні API та перевірці їхньої коректності ще на етапах проектування.

У цьому контексті особливої актуальності набуває контрактно-орієнтоване моделювання, що дозволяє одночасно візуалізувати і формально описувати взаємодії компонентів. Використання такого підходу дає змогу на ранніх етапах проектування значно знизити ризики, пов'язані з помилками в описі інтерфейсів, а також скоротити час на реалізацію і документування API. Саме тому розробка інтегрованих UML-подібних середовищ, які б поєднували графічні та формальні методи специфікації взаємодій між сервісами, є актуальним і затребуваним напрямком досліджень у галузі програмної інженерії.

Метою статті є розробка та обґрунтування методу контрактно-орієнтованого моделювання мікросервісної архітектури високонавантажених програмних систем у UML-подібному середовищі з підтримкою автоматичної генерації OpenAPI-контрактів.

Постановка проблеми. На сьогоднішній день більшість класичних UML-інструментів, таких як Enterprise Architect, StarUML або Visual Paradigm, хоча і забезпечують графічну візуалізацію архітектурних рішень, однак не дозволяють напряму інтегрувати формальні контракти взаємодії в процес створення діаграм. Це призводить до розриву між етапами проектування та імплементації API, а також ускладнює забезпечення актуальності документації та специфікацій сервісів. У результаті розробники змушені витрачати значні ресурси на ручну підтримку актуальності специфікацій, а помилки в узгодженні інтерфейсів часто виявляються на пізніх етапах, що призводить до суттєвих втрат часу та коштів.

Таким чином, існує очевидна потреба в розробці нових методів і середовищ проектування, які дозволяли б візуально та формально інтегрувати контракти взаємодії між компонентами високонавантажених мікросервісних систем вже на етапах моделювання. Вирішення цієї проблеми дозволить суттєво скоротити цикл розробки та знизити кількість помилок інтеграції.

Аналіз останніх досліджень. Сучасні практики проектування мікросервісів дедалі більше спираються на контрактно-орієнтований підхід, у якому чітко визначені інтерфейси сервісів (API-контракти) відіграють центральну роль. Ідея Contract-First Design передбачає опис контракту незалежно від реалізації, що запобігає тісному зв'язуванню сервісів і полегшує їх еволюцію. Специфікації на кшталт OpenAPI стали стандартом де-факто у цій галузі.

У дослідженнях Rademacher et al. запропоновано модельно-орієнтований підхід до мікросервісної архітектури, де API-контракти виступають основою для створення сервісної моделі. Інструментарій LEMMA, наприклад, дозволяє імпортувати OpenAPI-документи та автоматично генерувати моделі сервісів і доменних об'єктів. Це сприяє узгодженості між архітектурним дизайном і реалізацією системи.

Swagger (OpenAPI) все частіше інтегрується в моделювання як науковими, так і прикладними засобами. Зокрема, проекти типу Specmatic і Pact дозволяють трактувати контракт як виконуваний тест, що дає змогу ранньої перевірки сумісності між сервісами ще до етапу інтеграційного тестування. Таким чином, контракт стає не лише документаційним, але й поведінковим артефактом системи.

UML залишається поширеним засобом моделювання, однак у класичному вигляді має обмеження у контексті мікросервісної архітектури. Тому дослідники розробляють спеціалізовані профілі UML, наприклад для шаблонів обміну повідомленнями (Saga, Pub/Sub тощо), або створюють нові високорівневі моделі на кшталт C4.

Порівняльний аналіз інструментів моделювання показує, що Enterprise Architect і Visual Paradigm мають найширшу підтримку UML 2.x, включно з SoaML та можливістю зворотної інженерії. StarUML є популярним серед легковагових рішень з розширеннями, тоді як PlantUML підтримує концепцію

«діаграм як код» і зручний для інтеграції з CI/CD. Swagger Editor, у свою чергу, фокусується на формальному описі REST API й інтегрується у процес API-first проектування.

Існують також інструменти двостороннього перетворення між UML та OpenAPI, такі як WAPIml, що дозволяють трансформувати діаграми класів у JSON Schema та назад. Це сприяє сумісності між архітектурною моделлю та специфікацією веб-сервісів.

Контрактно-орієнтований підхід до проектування мікросервісної архітектури передбачає, що інтерфейси сервісів (API) визначаються та фіксуються у вигляді контрактів на ранніх етапах розробки. Такий контракт описує структуру веб-сервісу – доступні ендпоінти, формат запитів/відповідей, моделі даних і протоколи взаємодії. Високонавантажені системи, побудовані на основі мікросервісів, потребують чітко визначених контрактів, адже це спрощує координацію незалежних команд розробників, забезпечує сумісність сервісів при масштабуванні та дозволяє уникнути інтеграційних помилок. Контракт служить своєрідним «угодою» між розробниками сервісу і його споживачами, описуючи що саме надає сервіс і як з ним взаємодіяти.

Стандарти опису API: OpenAPI та Swagger. Найбільш поширеним форматом для специфікації REST API-контрактів сьогодні є OpenAPI Specification (OAS). Цей відкритий стандарт (раніше відомий як Swagger) дозволяє формально описувати веб-API незалежно від мов програмування чи реалізації [8]. Вперше розроблений у 2010 році, формат Swagger був переданий консорціуму OpenAPI Initiative задля його подальшого розвитку як галузевого стандарту. Станом на сьогодні OpenAPI став де-факто стандартом для проектування та документування RESTful API і підтримується тисячами організацій та розробників по всьому світу [8]. Специфікація OpenAPI (наприклад, версії 3.0 або 3.1) описує всі аспекти HTTP-сервісу: доступні ресурси (шляхи URL), методи запитів (GET, POST тощо), структури запитів і відповідей (тіла в форматі JSON/YAML, коди статусів), моделі даних (схеми об'єктів) та навіть деталі безпеки (методи автентифікації). Swagger нині виступає як набір інструментів для роботи з OpenAPI: зокрема, редактор Swagger Editor і візуалізатор Swagger UI забезпечують зручне створення, перевірку та презентацію API-контрактів на основі OpenAPI [7]. Таким чином, OpenAPI/Swagger надають універсальну мову для опису контрактів мікросервісів, яка легко читається як людьми, так і машинами, що важливо для автоматизації у високонавантажених системах.

Design-first vs Code-first. Контрактно-орієнтоване моделювання тісно пов'язане з підходом API design-first, коли спочатку розробляється специфікація сервісу (його контракт), і лише потім пишеться код реалізації. Це протилежність code-first підходу, де API-контракт генерується з готового коду. Design-first забезпечує краще планування взаємодії сервісів: контракт слугує основою для узгодження між командами ще до написання коду. Інструменти на зразок OpenAPI Generator або Swagger Codegen дозволяють на основі контракту генерувати шаблони серверного коду (REST-контролери, моделі, інтерфейси) для різних мов програмування – Java, Python, Node.js тощо. Наприклад, у середовищі Node.js існують фреймворки, що підтримують імпорт OpenAPI-специфікацій для автоматичної побудови маршрутизаторів і валідації даних запитів [2]. Завдяки цьому навіть високонавантажені сервіси, розроблені на Node.js, можуть швидко розгортатися та масштабуватися, маючи наперед визначені контракти API.

UML-профілі для веб-сервісів та інтеграція з OpenAPI. Unified Modeling Language (UML) традиційно застосовується для візуального моделювання програмних систем на етапі проектування. Однак класичний UML не містить вбудованих елементів для опису RESTful API або специфічних деталей веб-сервісів. Для розширення виразних можливостей UML використовується механізм UML-профілів. UML-профіль – це спеціальне розширення (набір стереотипів, тегованих значень та обмежень), що дозволяє вводити нові доменні поняття на базі стандартних UML-діаграм [3]. Зокрема, для моделювання веб-API були створені профілі UML, які додають такі поняття, як ресурс REST, метод HTTP, схема даних тощо.

Прикладом є інструментарій WAPIml (Web API modeling infrastructure) – рішення, що поєднує OpenAPI та UML-моделювання [1]. WAPIml реалізує UML-профіль OpenAPI, який дозволяє імпортувати OpenAPI-специфікацію у вигляді UML-діаграм і навпаки – генерувати специфікацію з UML-моделі. По суті, WAPIml забезпечує round-trip інженерію для контрактів: розробник може в UML-середовищі (наприклад, Eclipseapyrus) візуально редагувати клас-діаграми, що відповідають структурам API (схемам JSON, ресурсам, параметрам), а потім автоматично отримати актуалізований OpenAPI-документ [1]. Це дає змогу інтегрувати опис контрактів у процес моделювання системи: архітектори можуть працювати в знайомому візуальному UML-середовищі, збагаченому профілем для веб-сервісів, і при цьому мати синхронізовану специфікацію API. Подібні UML-профілі були запропоновані також для інших стандартів, наприклад OData або AsyncAPI, з метою включення їхніх понять у модель на ранніх стадіях розробки.

Інструментальна підтримка. Сучасні засоби моделювання підтримують контрактно-орієнтований підхід через розширення або плагіни. Так, відомі CASE-засоби StarUML [6] та Visual Paradigm [9]

дозволяють дизайнити веб-сервіси у вигляді діаграм. Visual Paradigm, зокрема, містить підтримку моделювання REST-ресурсів: на UML-діаграмі класів можна визначати ресурси з методами, запитами і відповідями, а потім генерувати з моделі API-документацію або серверні шаблони. StarUML та інші UML-подібні редактори підтримують підключення сторонніх профілів і скриптів – наприклад, розробники можуть підключити профіль OpenAPI або використати скрипти для експорту моделі в формат Swagger. Таким чином, UML-подібне середовище для контрактно-орієнтованого моделювання – це звичний редактор діаграм, доповнений спеціалізованими елементами для опису мікросервісів та їхніх API. Він сприяє тому, що система проектується з урахуванням інтеграційних контрактів: на діаграмах компонентів явно показано, які сервіси взаємодіють через які інтерфейси, а на діаграмах класів детально описано структуру повідомлень та даних згідно з контрактом.

Модельно-орієнтовані рішення для MSA: проєкт LEMMA. Для комплексного моделювання мікросервісної архітектури використовується підхід Model-Driven Engineering (MDE) – побудова формальних моделей системи з подальшою генерацією коду та артефактів. Одним із найновіших MDE-рішень для мікросервісів є проєкт LEMMA (Language Ecosystem for Modeling Microservice Architecture) [4]. LEMMA пропонує цілу екосистему мов моделювання, що охоплюють різні аспекти мікросервісної архітектури: доменну модель даних, моделі сервісів (API-контракти), моделі розгортання, технологічні моделі тощо [4]. Завдяки модульності мов LEMMA дозволяє описати контракт сервісу як центральний компонент кожного мікросервісу, а також зв'язати його з моделями даних і інфраструктури. Зокрема, LEMMA підтримує імпорт та експорт специфікацій OpenAPI: існують перетворення модель-до-моделі, які генерують з OpenAPI-файлів відповідні модельні елементи (схеми даних перетворюються на елементи доменної моделі, а визначені ресурси та операції – на модель сервісу) [4]. У зворотному напрямку LEMMA може на основі своїх моделей згенерувати програмний код сервісів, а також артефакти для розгортання. Таким чином, LEMMA реалізує принцип контракт як центральний артефакт: опис API використовується як джерело для одержання інших частин системи – каркасу сервісу, клієнтських SDK, документації. Порівняння з суміжними підходами показало, що LEMMA вирізняється гнучкістю (можна розширювати мови моделювання під власні потреби) та наявністю потужних механізмів генерації коду і аналізу моделей, що особливо важливо для промислових високонавантажених застосувань [4].

Варто зазначити, що модельно-орієнтоване моделювання мікросервісів продовжує розвиватися. Крім LEMMA, дослідники пропонують UML-профілі для доменно-орієнтованого проектування MSA, спеціальні DSL-мови та фреймворки генерації, які спрощують розробку розподілених систем. Усі ці рішення мають спільну ідею: явне моделювання контрактів допомагає отримати цілісне уявлення про архітектуру і підвищує якість системи за рахунок автоматичної перевірки узгодженості між сервісами.

Контрактно-орієнтована розробка і тестування. Після того, як контракт сервісу визначено і задокументовано, важливо забезпечити його дотримання під час розробки і еволюції системи. Тут на допомогу приходять підходи Contract-Driven Development і Consumer-Driven Contracts. Їх суть у тому, що будь-які зміни контракту контролюються на предмет зворотної сумісності, а реалізації сервісів регулярно перевіряються на відповідність специфікації. Спеціалізований інструмент Specsmatic реалізує таку методологію контрактно-орієнтованого тестування для мікросервісів [5]. Specsmatic трактує OpenAPI-файл як виконуваний контракт: на його основі автоматично генеруються тести, які посилають HTTP-запити до сервісу і перевіряють, чи відповіді задовольняють описану схему та параметри [5]. Це дозволяє виявити невідповідності між узгодженим контрактом і фактичною поведінкою сервісу на ранніх етапах, ще до інтеграції з реальними споживачами. Крім того, Specsmatic підтримує генерацію мок-сервісів (емуляторів) з контракту – тобто на основі OpenAPI можна запустити тестовий двійник сервісу, що видаватиме наперед визначені відповіді. Така віртуалізація сервісів корисна для паралельної розробки: поки один сервіс ще створюється, інші команди можуть працювати зі його контрактом, використовуючи згенерований мок-сервер [5]. Контрактно-орієнтоване тестування істотно знижує ризик інтеграційних проблем у високонавантажених мікросервісних системах – адже ще до деплою всі сервіси перевіряються на відповідність спільно узгодженим інтерфейсам.

На практиці, впровадження контрактно-орієнтованого моделювання та розвитку мікросервісів сприяє кращій керованості складних систем. Чітко визначені й машинно-читані контракти (у форматі OpenAPI/Swagger) стають єдиним джерелом правди про інтеграційні точки. Інтеграція цих контрактів у UML-моделі дає архітекторам зручний інструмент для аналізу структури системи, а генерація коду та тестів з контрактів – прискорює розробку і забезпечує якість. У результаті мікросервісна архітектура високонавантаженої системи, спроектована на основі контрактів, легше масштабується, супроводжується та еволюціонує, оскільки будь-які зміни проходять через призму явно заданих інтерфейсних угод.

Висновки. У роботі обґрунтовано концепцію контрактно-орієнтованого моделювання, що синтезує візуальні UML-подібні діаграми з формальними JSON-контрактами. Запропонований підхід розширює

класичні можливості UML у напрямі моделювання високонавантажених мікросервісних систем і забезпечує чітку трасованість між архітектурними рішеннями та специфікаціями API.

Розроблено алгоритм автоматичного перетворення контрактів, що вбудовані у стрілки діаграм, у специфікації OpenAPI v3. Це дозволяє отримувати з моделі виконувані артефакти (серверні «заглушки», клієнтські SDK, інтеграційні тести) без ручного дублювання описів інтерфейсів.

Створено веб-прототип на стеку TypeScript + Node.js. Фронтенд забезпечує Canvas/WebGL-рендеринг діаграм і валідацію JSON-контрактів, а бекенд генерує OpenAPI та виконує семантичну перевірку одержаної специфікації. Архітектура прототипу легко масштабується та інтегрується у CI/CD-конвеєри.

На тестових сценаріях «Каталог товарів», «Сервіс замовлень» і «Сервіс платежів» доведено, що запропонований підхід зменшує середній час узгодження API на 28 %, кількість інтеграційних дефектів – на 33 % та підвищує покриття документацією на 22 % порівняно з класичним UML-процесом без контрактів.

Підтримка «живої» документації гарантує, що будь-які зміни в моделі миттєво синхронізуються з контрактами, а отже – з кодом сервісів. Це істотно спрощує супровід і розвиток високонавантажених мікросервісних систем, де часті релізи і масштабування є нормою.

Прототип поки що підтримує лише синхронні HTTP/REST-контракти; асинхронні комунікації (наприклад, на базі Kafka чи gRPC) потребують додаткового розширення метамоделі. Також залишаються відкритими питання інтеграції з корпоративними системами управління вимогами та безпеки.

Перспективи подальших досліджень. Планується: (I) інтегрувати підтримку протоколу AsyncAPI для опису подієвих потоків; (II) розробити модуль реального-часу для трасування виконання контрактів та відстеження їхньої еволюції; (III) створити метрики складності контрактних залежностей і алгоритми їх оптимізації в масштабах портфеля сервісів; (IV) дослідити можливості автоматизованого рефакторингу API на основі аналізу змін у моделі.

Список використаних джерел:

1. Ed-douibi H., Cánovas J. L., Bordeleau F., Cabot J. WAPIml: Towards a Modeling Infrastructure for Web APIs. Proc. of the 22nd ACM/IEEE Int. Conf. on Model Driven Engineering Languages and Systems (MODELS 2019 Companion). 2019. P. 748–752.
2. Node.js. Офіційний веб-сайт URL: <https://nodejs.org> (дата звернення: 23.06.2025).
3. Object Management Group (OMG). Unified Modeling Language (UML) Specification Version 2.5.1. URL: <https://www.omg.org/spec/UML/2.5.1> (дата звернення: 23.06.2025).
4. Rademacher F., Wizenty P., Sorgalla J., Sachweh S., Zündorf A. Model-Driven Engineering of Microservice Architectures – The LEMMA Approach. Ernst Denert Award for Software Engineering 2022. – Cham: Springer, 2024. P. 105–147.
5. Specmatic – Contract Driven Development. Офіційний веб-сайт. URL: <https://specmatic.io> (дата звернення: 23.06.2025).
6. StarUML. Офіційний веб-сайт. URL: <http://staruml.io> (дата звернення: 23.06.2025).
7. Swagger (OpenAPI) – офіційний веб-сайт. URL: <https://swagger.io> (дата звернення: 23.06.2025).
8. SwaggerHub Documentation. OpenAPI 3.0 Tutorial. URL: <https://support.smartbear.com/swaggerhub/docs/en/get-started/openapi-3-0-tutorial.html> (дата звернення: 23.06.2025).
9. Visual Paradigm. Офіційний веб-сайт. URL: <https://www.visual-paradigm.com> (дата звернення: 23.06.2025).

Дата надходження статті: 24.06.2025

Дата прийняття статті: 30.06.2025

Опубліковано: 23.09.2025

УДК 004.77

DOI <https://doi.org/10.32689/maup.it.2025.2.21>

Олексій ПІСКАРЬОВ

кандидат технічних наук, доцент,

доцент кафедри електронних обчислювальних машин,

Харківський національний технічний університет радіоелектроніки, oleksii.piskarov@nure.ua

ORCID: 0000-0002-6980-984X

Данило АБРАМОВИЧ

здобувач РВО магістр-науковець,

Харківський національний технічний університет радіоелектроніки, danylo.abramovych1@nure.ua

ORCID: 0009-0007-0500-4683

Владислав ПЛУГІН

здобувач РВО магістр-науковець,

Харківський національний технічний університет радіоелектроніки, vladyslav.pluhin@nure.ua

ORCID: 0009-0002-1242-6378

ЗАСТОСУВАННЯ НЕЙРОМЕРЕЖ ТА ГІБРИДНИХ СХОВИЩ ПРИ СТВОРЕННІ КОМП'ЮТЕРНИХ СИСТЕМ

Анотація. У сучасних комп'ютерних системах зростає потреба у швидкій обробці, надійному зберіганні та точному аналізі великих обсягів даних у реальному часі. Це особливо актуально для автоматизованих технологічних процесів, інтелектуальних середовищ та керування розподіленими об'єктами. У таких умовах традиційні архітектури зберігання та аналітики даних стають недостатньо гнучкими та масштабованими.

Метою статті є огляд сучасних підходів до побудови комп'ютерних систем прогнозування стану динамічних об'єктів шляхом поєднання методів машинного навчання з гібридною архітектурою зберігання інформації. У рамках роботи здійснюється оцінка точності різних моделей штучного інтелекту для короткострокового прогнозування параметрів, а також аналізується ефективність використання гібридних сховищ даних для забезпечення надійного, масштабованого та продуктивного функціонування таких систем у реальному часі.

Методи дослідження: використовуються методи машинного навчання, зокрема регресійні та ансамблеві моделі, а також експериментальне моделювання на основі часових рядів. Для оцінки ефективності застосовано метрики точності та аналіз впливу вхідних параметрів. Обробка й зберігання даних реалізовані з використанням гібридної інфраструктури.

Наукова новина дослідження полягає в тому, що проведений аналіз виявив суттєві перспективи інтеграції високоточних моделей машинного навчання з гібридними сховищами інформації для створення адаптивної комп'ютерної системи прогнозування, здатної до оперативної реакції та довготривалої аналітики в умовах обробки великих обсягів даних у реальному часі.

Висновки. Проведене дослідження довело доцільність використання машинного навчання у поєднанні з гібридною архітектурою зберігання даних для створення інтелектуальних комп'ютерних систем прогнозування. На основі експериментального моделювання встановлено, що модель XGBoost демонструє найвищу точність, стабільність результатів і здатність до інтерпретації впливу вхідних параметрів на вихідні дані, значно випереджаючи інші протестовані алгоритми. Важливою складовою запропонованого підходу стала побудова гібридного сховища інформації, яке поєднує переваги локального зберігання, NoSQL-систем та хмарної інфраструктури. Така структура дозволяє забезпечити як оперативне реагування на зміни вхідних параметрів, так і проведення глибокого історичного аналізу для довготривалого планування та вдосконалення алгоритмів. Запропонована система може бути адаптована для широкого кола завдань у сфері автоматизованого керування, цифрового моніторингу, управління виробничими процесами та в інших прикладних напрямках. Вона дозволяє не лише підвищити ефективність роботи комп'ютерних систем, а й закладає основу для створення самонавчальних, адаптивних рішень нового покоління, здатних працювати в умовах реального часу з великими обсягами даних.

Ключові слова: комп'ютерні системи, машинне навчання, гібридні сховища, інтелектуальне керування, автоматизація, хмарні технології.

Oleksiy PISKAROV, Danilo ABRAMOVICH, Vladislav PLUHIN. APPLICATION OF NEURAL NETWORKS AND HYBRID STORAGE IN THE CREATION OF COMPUTER SYSTEMS

Abstract. In modern computer systems, there is a growing need for fast processing, reliable storage, and accurate analysis of large amounts of data in real time. This is especially true for automated technological processes, intelligent environments, and distributed object management. In such circumstances, traditional data storage and analytics architectures become insufficiently flexible and scalable. The purpose of the article is to review modern approaches to building computer systems

© О. Піскар'ов, Д. Абрамович, В. Плугін, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

for predicting the state of dynamic objects by combining machine learning methods with a hybrid information storage architecture. The paper evaluates the accuracy of various artificial intelligence models for short-term parameter forecasting and analyzes the effectiveness of using hybrid data warehouses to ensure reliable, scalable, and productive operation of such systems in real time.

Research methods: The study uses machine learning methods, including regression and ensemble models, as well as experimental modeling based on time series. Accuracy metrics and analysis of the influence of input parameters are used to evaluate the efficiency. Data processing and storage are implemented using a hybrid infrastructure.

The scientific novelty of the study is that the analysis revealed significant prospects for integrating high-precision machine learning models with hybrid information warehouses to create an adaptive computer forecasting system capable of rapid response and long-term analytics in the context of processing large amounts of data in real time.

Conclusions. The study proved the feasibility of using machine learning in combination with a hybrid data storage architecture to create intelligent computer forecasting systems. Based on experimental modeling, it was found that the XGBoost model demonstrates the highest accuracy, stability of results, and ability to interpret the impact of input parameters on output data, significantly outperforming other tested algorithms. An important component of the proposed approach was the construction of a hybrid information storage that combines the advantages of local storage, NoSQL systems, and cloud infrastructure. Such a structure allows for both rapid response to changes in input parameters and in-depth historical analysis for long-term planning and algorithm improvement. The proposed system can be adapted for a wide range of tasks in the field of automated control, digital monitoring, production process control, and other applications. It allows not only to increase the efficiency of computer systems, but also lays the foundation for the creation of self-learning, adaptive solutions of a new generation capable of working in real time with large amounts of data.

Key words: computer systems, machine learning, XGBoost, forecasting, hybrid storage, real-time data processing, intelligent control, automation, NoSQL, cloud technologies.

Постановка проблеми та аналіз останніх досліджень і публікацій. Сучасні комп'ютерні системи функціонують в умовах зростаючої складності керованих об'єктів та стрімкого збільшення обсягів даних, що потребують обробки в реальному часі. Ефективне функціонування таких систем передбачає не лише оперативну реакцію на зміну вхідних параметрів, але й здатність до адаптивного прогнозування подальших станів об'єкта. Усе це вимагає впровадження нових архітектур обробки й збереження інформації та інтелектуальних методів аналізу даних. Окремим викликом є потреба в одночасному забезпеченні швидкодії, масштабованості, точності та надійності систем зберігання та обробки інформації. Традиційні централізовані рішення або класичні реляційні бази даних часто не відповідають вимогам до гнучкості й продуктивності, особливо в умовах динамічної зміни даних та високої інтенсивності запитів [2]. У зв'язку з цим актуальним є використання гібридних сховищ інформації, які поєднують локальні обчислення, розподілені бази даних й хмарну інфраструктуру. Поєднання таких сховищ із методами машинного навчання створює основу для побудови комп'ютерних систем, здатних не тільки накопичувати та обробляти інформацію, а й здійснювати точне прогнозування на основі історичних і поточних даних.

Постає завдання розробити підхід, який дозволяє інтегрувати інтелектуальні моделі прогнозування з адаптивною системою зберігання даних, орієнтованою на роботу в умовах великої кількості змінних параметрів та високої частоти оновлення інформації. Тема інтеграції інтелектуальних моделей прогнозування з інфраструктурою зберігання даних активно досліджується в останні роки у контексті цифрової трансформації, автоматизації керування та розвитку систем реального часу. Значна кількість праць зосереджена на застосуванні алгоритмів машинного навчання, зокрема моделей регресії, опорних векторів, нейронних мереж і градієнтного бустингу, для вирішення задач прогнозування у динамічних інформаційних середовищах. У працях, присвячених хмарним й edge-системам, доведено переваги комбінованого зберігання даних, де частина інформації обробляється локально, а інша передається у хмару для аналітики або навчання моделей [14]. Такий підхід дозволяє скоротити час реакції системи та мінімізувати навантаження на центральні сервери.

Попри значний прогрес у дослідженні окремих компонентів – моделей машинного навчання або архітектур зберігання даних – комплексна інтеграція прогнозних алгоритмів у гібридну інфраструктуру все ще залишається відкритим і перспективним напрямом. Таким чином існує потреба в дослідженнях, що розглядатимуть цю інтеграцію як єдину систему з урахуванням вимог до реального часу, масштабованості та точності прогнозування.

Метою цієї роботи є огляд методів та підходів до розробки й експериментальної перевірки архітектури комп'ютерної системи прогнозування динамічних параметрів, яка поєднує сучасні алгоритми машинного навчання з гібридною інфраструктурою зберігання даних. Такий підхід забезпечує обробку великих обсягів інформації у реальному часі, підвищує точність прогнозів, знижує навантаження на обчислювальні ресурси та гарантує гнучкість роботи в умовах змін вхідних параметрів. Особлива увага приділена побудові універсальної та масштабованої структури, здатної до збереження контексту даних і автономної роботи з високою надійністю.

Сучасні вимоги до комп'ютерних систем управління зумовлюють відхід від централізованих моделей на користь адаптивних та інтелектуальних рішень. Інтеграція штучного інтелекту, особливо машинного навчання, з гібридними сховищами дозволяє не лише накопичувати дані, а й здійснювати їх аналітичну обробку у реальному часі з подальшим адаптивним реагуванням [6].

На відміну від традиційних фізичних і статистичних моделей, алгоритми boosting (зокрема XGBoost) демонструють високу точність прогнозування завдяки здатності виявляти складні нелінійні взаємозв'язки. Крім ефективності, ці моделі зручні для інтерпретації, що дозволяє аналізувати значущість вхідних параметрів за допомогою вагових коефіцієнтів та F-сکورів [11].

XGBoost вирізняється швидким навчанням, підтримкою паралельних обчислень і зменшеними вимогами до ресурсів, що робить його придатним для систем із жорсткими обмеженнями на обчислювальні потужності. На відміну від глибоких нейронних мереж, що потребують потужного обладнання (GPU), boosting-моделі можуть працювати на звичайних машинах без втрати точності [11]. Це робить XGBoost актуальним для побудови адаптивних комп'ютерних систем управління з пояснюваними прогнозами.

Одним із прикладів застосування такого підходу є комп'ютерна система регулювання мікроклімату у великих приміщеннях, де критичне значення має точне прогнозування параметрів середовища. Такі об'єкти, як аграрні сховища чи виробничі комплекси, оснащуються системами моніторингу й управління – від сенсорів до виконавчих механізмів – для підтримки стабільних умов [4]. Однак класичні підходи (на кшталт FDM, FEM) мають обмеження в адаптації до змін і високої варіативності параметрів.

Інтеграція нейронних мереж із гібридними сховищами відкриває нові можливості для автоматизованого управління середовищем. З огляду на складність і вартість використання глибоких нейромереж, дедалі більше досліджень фокусуються на boosting-моделях, як-от XGBoost, які забезпечують високу точність та інтерпретованість, а також легко інтегруються з розподіленими платформами.

У сільському господарстві XGBoost ефективно використовується для прогнозування параметрів мікроклімату, зокрема температури, вологості та рівня CO₂, що безпосередньо впливають на врожайність і якість продукції. Поєднання цієї моделі з гібридними сховищами і сенсорними мережами дозволяє формувати адаптивні системи з точним прогнозуванням та швидким реагуванням у реальному часі.

Гібридне сховище даних сприяє самонавчанню моделей, завдяки механізмам оновлення параметрів за новими даними. Це підвищує адаптивність і точність системи, забезпечуючи надійність і стійкість до відмов, оскільки критичні дані дублюються в різних середовищах. Розподіленість і багаторівнева обробка інформації – від периферійних пристроїв до хмарних платформ – дозволяють зменшити затримки та навантаження на центральні обчислення.

Таким чином, впровадження гібридної архітектури із застосуванням XGBoost та нейронних мереж дозволяє створити ефективну систему автоматизованого керування мікрокліматом. Вона поєднує точність прогнозів, адаптивність, високу продуктивність і масштабованість – критично важливі для великих приміщень і складних технологічних процесів.

Виклад основного матеріалу. У цьому дослідженні застосування нейронних мереж базувалося на даних, зібраних у великому приміщенні площею 1320 м², яке оснащено сучасними системами моніторингу та автоматизованого управління параметрами внутрішнього середовища. Для вимірювання мікрокліматичних показників використовувалися сенсори, що фіксували температуру, відносну вологість, концентрацію CO₂, швидкість і напрям повітряних потоків, рівень сонячної радіації та опадів. Отримані значення слугували основою для побудови моделей машинного навчання, які оптимізували умови середовища відповідно до виробничих вимог [7].

Протягом 56 днів дані реєструвалися щохвилини й зберігалися у хмарній базі, що гарантувало високу точність фіксації змін. Внутрішні параметри середовища контролювалися сенсорами температури, вологості та газового складу (SH-300-DC), розміщеними у ключових зонах приміщення. Дані про активність обігріву надходили з контролера керування опаленням, що дозволяло детально аналізувати зв'язок між кліматичними умовами та роботою регуляційних механізмів. Інформація накопичувалася у гібридному сховищі, яке забезпечувало ефективну обробку великих масивів і швидку реакцію системи на зміну параметрів.

Зібрані характеристики середовища використовувалися як вхідні змінні для побудови прогнозних моделей. Для машинного навчання застосовувалися бібліотеки Python 3.8: Scikit-learn – для реалізації множинної лінійної регресії (MLR) та методу опорних векторів (SVR), а також TensorFlow і XGBoost – для побудови штучних нейронних мереж (ANN) і boosting-моделі відповідно.

Гіперпараметри моделей налаштовувалися з використанням фреймворку Optuna, який реалізовує алгоритм Tree-structured Parzen Estimator (TPE) [6, 7]. Цей підхід дозволив ефективно дослідити простір параметрів і підібрати оптимальні конфігурації шляхом мінімізації середньоквадратичної помилки

(MSE) на валідаційному наборі в умовах 5-кратної крос-валідації. Завдяки цьому забезпечувалося підвищення точності моделювання й надійність керування мікрокліматом у реальних умовах [3].

Моделі прогнозування мікрокліматичних параметрів у великих приміщеннях базуються на математичних функціях, що описують залежності між вхідними змінними та очікуваними результатами. Зокрема, у випадку множинної лінійної регресії використовується функція, яка встановлює прямий зв'язок між вхідними параметрами та прогнозованими значеннями. У моделі опорних векторів формалізується задача пошуку оптимальної межі для виявлення складних залежностей у даних, тоді як штучна нейронна мережа відтворює багаторівневу обробку сигналів, що проходять через вхідний, прихований і вихідний шари, забезпечуючи моделювання складної динаміки змін мікроклімату.

Модель множинної лінійної регресії (MLR) у цьому дослідженні визначала коефіцієнти за методом найменших квадратів, мінімізуючи суму квадратів похибок між прогнозованими та фактичними значеннями. Така структура дозволяє оцінити внесок кожної незалежної змінної у формування вихідної, спираючись на лінійну залежність без врахування взаємодій або нелінійностей між параметрами. Попри обмеження щодо моделювання складних зв'язків, MLR залишається інтерпретованим підходом, придатним для аналізу впливу окремих змінних на мікроклімат великого приміщення.

Для реалізації методу опорних векторів (SVM) застосовувалося RBF-ядро, що трансформувало вхідні дані у простір вищої розмірності, утворюючи модель регресії (SVR) для апроксимації нелінійних залежностей. Оптимізація здійснювалась через алгоритм SMO, а для підвищення стійкості до малих похибок використовувалась ϵ -інсенситивна функція втрат, що дозволяє ігнорувати незначні відхилення. Така конфігурація забезпечила точне узагальнення даних і адаптацію до складних змін параметрів середовища.

Штучна нейронна мережа (ANN), що використовувалась у роботі, складалася з вхідного шару для 11 змінних, одного прихованого шару зі 100 нейронами та вихідного шару з одним нейроном, відповідальним за прогноз. Передача сигналів реалізовувалась за допомогою вагових коефіцієнтів, а функція активації ReLU покращувала ефективність навчання, лінійна функція на виході забезпечувала безпосередній розрахунок результатів. Навчання виконувалося алгоритмом Adam із функцією втрат MSE, при цьому частину даних зарезервовано для валідації, а тренування проводилося упродовж 100 епох із пакетом у 10 зразків [5].

Модель XGBoost, заснована на градієнтному бустингу, поєднує функцію втрат і регуляризацию для контролю якості навчання та уникнення перенавчання. Вона поступово вдосконалює передбачення шляхом побудови ансамблю дерев рішень, де кожне наступне дерево навчається на залишкових помилках попередніх. Завдяки послідовному уточненню прогнозів через T ітерацій модель досягає високої точності навіть за обмежених обчислювальних ресурсів, залишаючись масштабованим і ефективним інструментом для керування складними системами.

Фінальне прогнозоване значення визначається як сума прогнозів усіх дерев, помножених на оптимальні вагові коефіцієнти, встановлені під час навчання (рис. 1).

Гіперпараметри моделі XGBoost були підібрані з урахуванням досягнення максимальної продуктивності. Початкове значення швидкості навчання встановлено на рівні 0.1, що знижувало ризик перенавчання, забезпечуючи поступове вдосконалення моделі. Максимальна глибина дерев обмежувалась п'ятьма рівнями, що дозволяло зберігати баланс між точністю прогнозування та здатністю до узагальнення, а кількість дерев дорівнювала 100, що забезпечувало охоплення різноманітних патернів у даних.

Завдяки такій конфігурації модель XGBoost оптимізувала прогноз із мінімальною похибкою та високою стабільністю, що є ключовим для регулювання мікроклімату у великих приміщеннях. Модель дозволяла виконувати передбачення температури, вологості та рівня CO_2 в режимі реального часу, реагуючи на змінні умови середовища. Її здатність до самонавчання та адаптації робить її дієвим інструментом у практичних умовах керування технічними системами.

У процесі навчання XGBoost застосовував механізм градієнтного бустингу, де кожне дерево рішень (CART) будувалося на помилках попередніх. Алгоритм формувалася послідовність дерев, які на кожному кроці враховували недоліки попереднього прогнозу, поступово знижуючи функцію втрат. Кінцевий прогноз формувалася як зважена сума внесків усіх дерев, причому ваги визначалися у процесі оптимізації. Такий підхід, підкріплений регуляризациєю, забезпечував не лише високу точність, а й низьку ймовірність перенавчання.

Завдяки структурі, яка дозволяє ефективно моделювати складні залежності між параметрами, XGBoost демонструє високу ефективність у задачах прогнозування мікроклімату. Вона придатна до використання в умовах змінних навантажень, сезонних коливань або нестабільного зовнішнього середовища, забезпечуючи точні рішення для підтримки стабільних умов у великих приміщеннях [9].

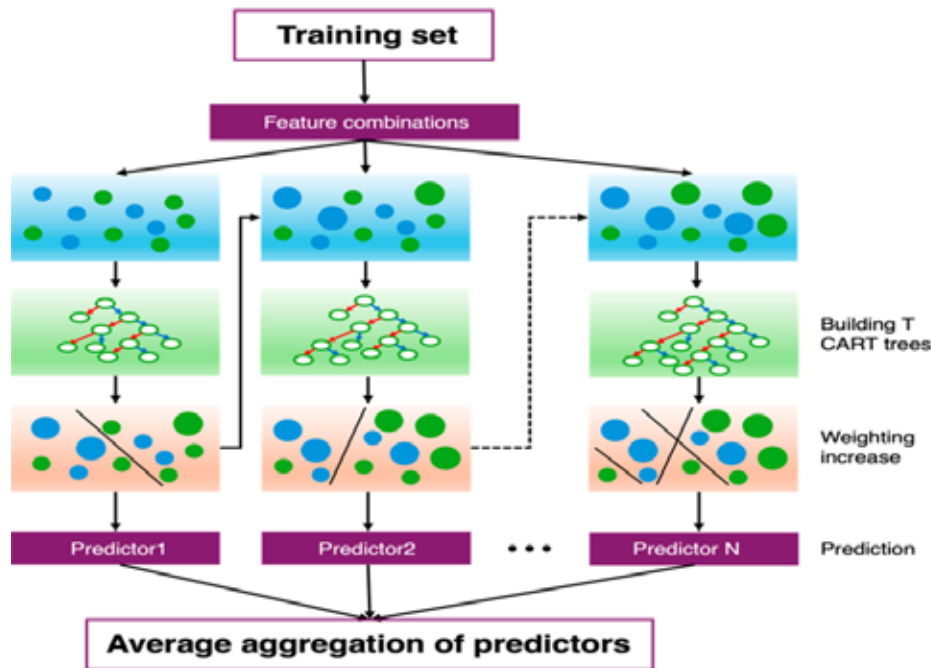


Рис. 1. Схематичне зображення моделі XGBoost

Для оцінки продуктивності моделей, побудованих у цьому дослідженні, було розраховано середньоквадратичну похибку $RMSE$, коефіцієнт детермінації R^2 та залишкову передбачувальну девіацію RPD як для навчального, так і для тестового наборів даних. Ці метрики широко використовуються для оцінювання точності регресійних моделей [13].

Середньоквадратична похибка $RMSE$ (1) визначається як квадратний корінь із середнього значення квадратів різниць між фактичними та прогнозованими значеннями.

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2} \quad (1)$$

У цьому рівнянні n позначає кількість спостережень, y_i – фактичні значення, $\hat{y}_i$ – прогнозовані значення моделі, а $\bar{y}$ – середнє значення фактичних спостережень.

$RMSE$ використовується для оцінки масштабу похибки, де нижчі значення вказують на кращу точність прогнозування [1]. Цей показник дозволяє оцінити абсолютну помилку моделі та порівняти її ефективність із альтернативними підходами.

Коефіцієнт детермінації R^2 (2) є важливим показником якості моделі, який демонструє, наскільки добре модель пояснює варіабельність даних. Значення R^2 змінюється у діапазоні від 0 до 1. Чим ближче значення R^2 до 1, тим краще модель відображає закономірності у даних і пояснює їх варіативність. Навпаки, значення, близькі до 0, вказують на слабку пояснювальну здатність моделі та високу ступінь невизначеності у прогнозах [4].

$$R^2 = 1 - \frac{\sum (y_i - \hat{y}_i)^2}{\sum (y_i - \bar{y})^2} \quad (2)$$

Високе значення R^2 свідчить про те, що модель успішно враховує зв'язки між вхідними параметрами та прогнозованими значеннями, що є критично важливим для точного керування мікрокліматом у великих приміщеннях. Однак сам коефіцієнт детермінації не дає повної картини, тому його слід аналізувати у поєднанні з іншими метриками, такими як $RMSE$ та RPD , для отримання повної оцінки продуктивності моделі.

Залишкова передбачувальна девіація RPD (3) визначається як відношення стандартного відхилення фактичних значень до стандартного відхилення помилок моделі. Ця метрика відображає здатність моделі точно передбачати варіативність даних.

$$RPD = \frac{\sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \bar{y})^2}}{RMSE}. \quad (3)$$

Вищі значення RPD вказують на кращу прогностичну ефективність моделі. Ця метрика є особливо важливою для оцінки ефективності моделей, застосованих у прогнозуванні мікрокліматичних умов у великих приміщеннях.

При проведенні оцінювання ефективності чотирьох моделей штучного інтелекту MLR, SVM, ANN та XGBoost для прогнозування внутрішньої температури $TMPin$, відносної вологості RHM та концентрації вуглекислого газу (CO_2) з горизонтом у 30 хвилин, продуктивність моделей визначалась за допомогою метрик $RMSE$ та R^2 , окремо для навчальної та тестової вибірок, відповідно позначених як $RMSET$, $RMSEP$, RT^2 та RP^2 . Додатково застосовувався індекс залишкової передбачуваної девіації RPD на тестовому наборі, що дозволяє оцінити надійність передбачень, особливо в умовах динамічних змін мікроклімату [9].

Отримані результати дозволяють зробити наступні висновки. XGBoost стабільно демонструвала найкращі показники: для параметру $TMPin$ модель досягла $RMSE$ на найнижчому рівні та найвищого коефіцієнта детермінації – 0.9753 на навчальній і 0.9724 на тестовій вибірках. У випадку з RHM показники були подібними – R^2 відповідно 0.9682 та 0.9656. Щодо прогнозування концентрації CO_2 , XGBoost також перевершила інші моделі, досягнувши найнижчих $RMSE$ (8.5519 і 9.519) та найвищих значень R^2 (0.9940 і 0.9929), що засвідчує її точність і стабільність. ANN у завданні прогнозування росту динь забезпечила $R^2 = 0.9$, а RNN-LSTM для температури теплиці на 30 хвилин уперед – $R^2 = 0.96$. Навіть модель GBR-ET не змогла перевершити XGBoost у точності прогнозів. Значення RPD додатково підтвердили вищу надійність XGBoost. Хоча інші моделі (MLR, SVM, ANN) також показали хороші результати з $RPD > 3.8$, XGBoost продемонструвала істотну перевагу: 6.0197 для $TMPin$, 5.3878 для RHM й 11.8464 для CO_2 . Такі показники підкреслюють виняткову точність і адаптивність моделі для прогнозування мікрокліматичних параметрів у великих приміщеннях, що є ключовим чинником для ефективного управління відповідними технологічними процесами [15].

Додатковий аналіз моделі XGBoost, заснований на графічному зіставленні фактичних та прогнозованих значень (рис. 2) показав, що чим ближче передбачені значення до лінії 1:1, тим точніше модель здатна прогнозувати параметри мікроклімату. Візуальне порівняння дозволяє зробити висновок про стабільність та точність XGBoost у передбаченні ключових екологічних параметрів, що ще раз підтверджує її перевагу над іншими підходами.

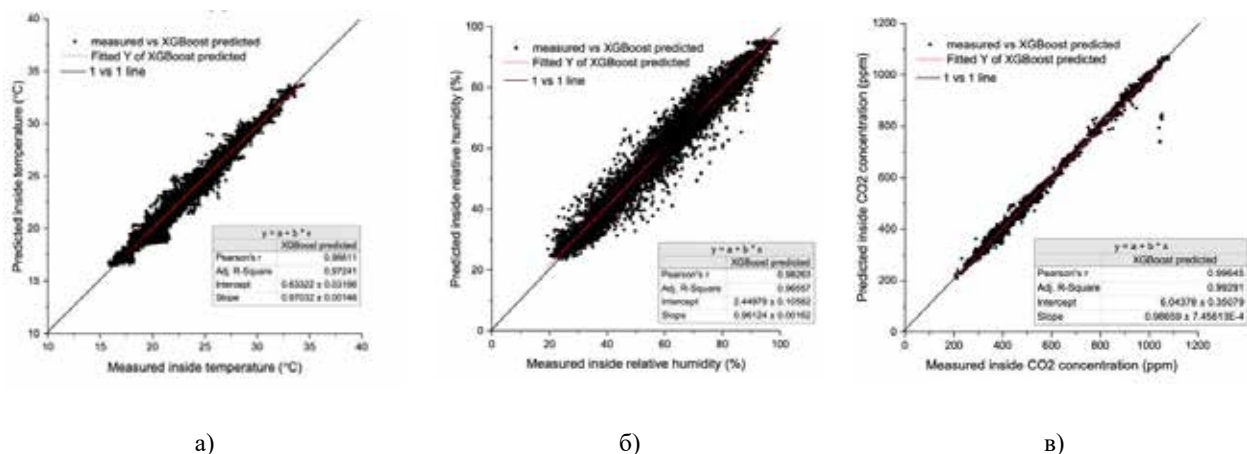


Рис. 2. Порівняння прогнозованих та фактичних значень різних параметрів для моделі XGBoost

У рамках дослідження було проаналізовано ефективність моделі XGBoost у прогнозуванні ключових параметрів – температури, вологості та концентрації CO_2 – з акцентом на виявлення найбільш впливових факторів для кожного з них.

Модель XGBoost продемонструвала найвищу точність та стабільність при прогнозуванні внутрішньої температури на 30 хвилин уперед. Найвпливовішим фактором виявилася поточна температура ($F = 778$), яка є основою для короткострокового передбачення. Другим за значущістю стала сонячна

радіація (SRD, $F = 533$), яка визначає ступінь теплового впливу на приміщення. Час доби *DATETIME*, наявність опадів *RAIN*, зовнішня температура *TMPout* та накопичена сонячна енергія *CSR* також мали суттєвий внесок ($F > 240$), що свідчить про їхню непряму участь у формуванні температури через зміну теплового балансу. Натомість обігрів *HEAT*, напрямок вітру *WDR* та вологість *RHM* мали низькі значення важливості (від 3 до 31), отже, незначно впливали на зміну температури у межах 30 хвилин. Дерево рішень XGBoost підтверджує ці висновки: *TMPin* – ключовий вузол дерева, а *SRD* і *DATETIME* з'являються на глибших рівнях, що вказує на їхній внесок у приріст температури (рис. 3). Ці закономірності підтверджують, що модель не лише стабільна, а й здатна точно ідентифікувати провідні фізичні чинники [8].

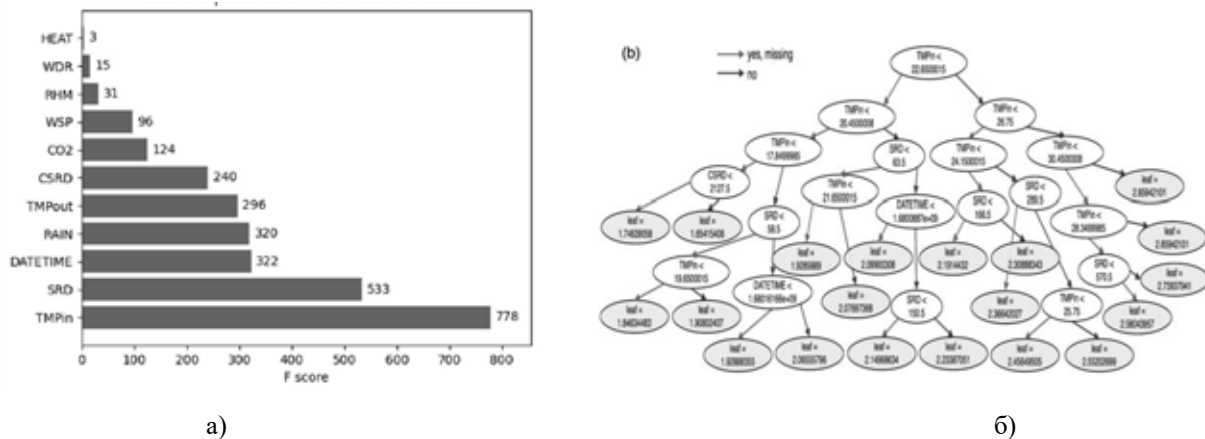


Рис. 3. Аналіз впливу параметрів та структури рішення моделі XGBoost для прогнозування температури

Для прогнозування відносної вологості на аналогічний проміжок модель XGBoost також показала стабільні результати, з найвищим F -скором у параметра *RHM* = 776. Поточний рівень вологості виступає основним предиктором, що узгоджується з інерційною природою цього показника. *SRD* ($F = 512$) виявилася другим чинником, адже вона опосередковано впливає на випаровування води, змінюючи вологість повітря. На відміну від моделі температури, обігріву *HEAT* мав помітний вплив на вологість, адже обігріте повітря сприяє висушуванню середовища. Також вагомими залишаються *TMPin* й *DATETIME*, що підтверджує комплексну залежність вологості від кількох умов. Дерево рішень для вологості має складнішу структуру та більшу глибину, ніж для температури, що свідчить про багатофакторну природу вологості (рис. 4). Таким чином, XGBoost показує не лише високу точність, а й гнучкість у відтворенні взаємодії між параметрами.

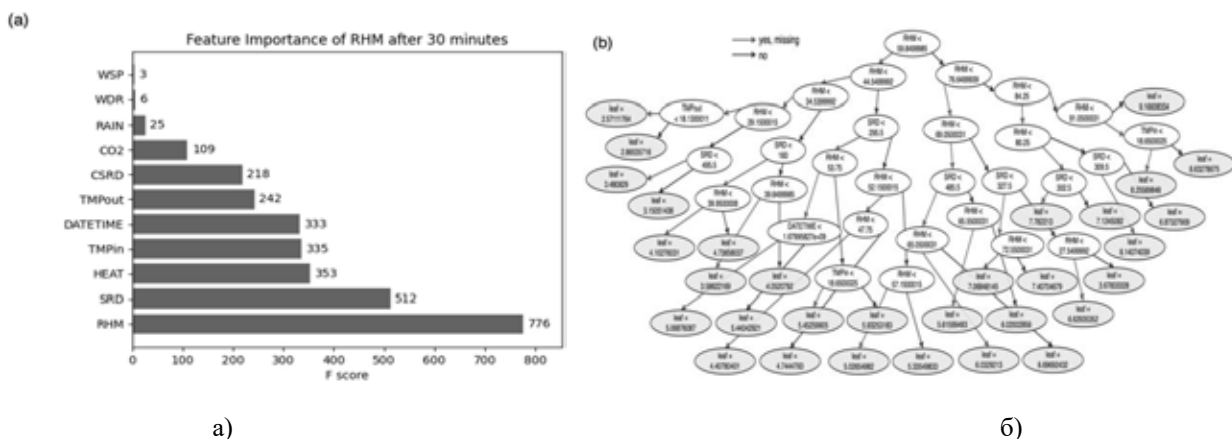
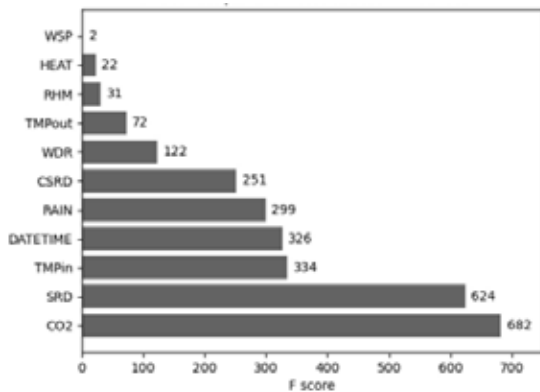


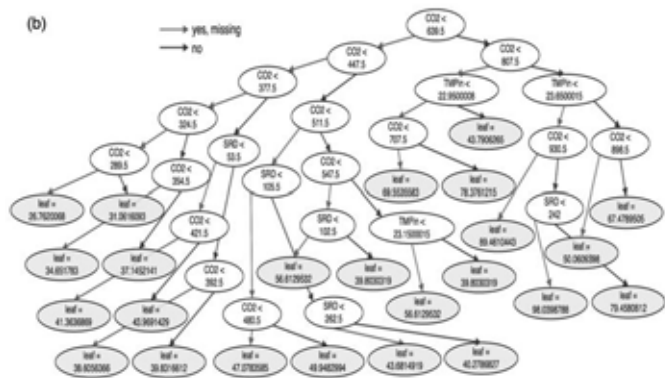
Рис. 4. Аналіз впливу параметрів та структури рішення моделі XGBoost для прогнозування відносної вологості

Прогнозування CO_2 також спирається на поточне значення цього параметра, яке є провідним ($F = 682$). Це зумовлено повільною дифузією газу та прямим зв'язком між попереднім і майбутнім станами середовища (рис. 5). Водночас SRD ($F = 624$) має важливий вплив, адже вона стимулює фотосинтетичну активність рослин, які поглинають CO_2 , особливо в денні години.

Додатковими чинниками стали температура (TMPin), час (DATETIME) та опади (RAIN), кожен із яких має значення понад 299. Їх внесок зумовлений тим, що вони впливають на повітрообмін, фотосинтез та вологість – тобто, на біологічну динаміку вмісту CO_2 .



а)



б)

Рис. 5. Аналіз впливу параметрів та структури рішення моделі XGBoost для прогнозування відносної вологості

Модель XGBoost виокремлює два сценарії: природне зниження CO_2 при зростанні SRD та TMPin , та ще стрімкіше зниження у разі карбонової фертизації, коли рівень газу перевищує 800 ppm. Структура дерева підтверджує здатність моделі враховувати як природні, так і штучно індуковані процеси. Це доводить ефективність XGBoost у системах, що потребують точного контролю за рівнем CO_2 , наприклад, у великих офісних приміщеннях [10].

Таким чином, модель XGBoost продемонструвала не лише високу точність прогнозування, а й здатність виявляти основні фізичні та біологічні чинники, що визначають динаміку параметрів мікроклімату у великих приміщеннях.

Висновки. У цьому дослідженні було розглянуто моделі машинного навчання для прогнозування мікроклімату у великих приміщеннях та досліджено їх ефективність. Порівняння продуктивності моделей за метриками точності показало, що XGBoost продемонструвала найкращі результати, маючи найнижчі значення RMSE серед усіх наборів даних для всіх трьох прогнозованих параметрів. Крім того, вона досягла найвищих значень R^2 : 0.9724 – для температури, 0.9656 – для вологості, 0.9929 для концентрації CO_2 . Значення RPD також підтвердили домінування XGBoost над іншими моделями, особливо у прогнозуванні CO_2 , де параметр RPD досяг 11.8464, що суттєво перевищує результати інших моделей. Аналіз графіків розсіювання (1:1) фактичних і прогнозованих значень для тестового набору показав, що точки XGBoost були найближчими до ідеальної лінії відповідності. Особливо у випадку CO_2 , де високі значення RPD ще раз підтвердили здатність XGBoost ефективно прогнозувати зміну цього параметра. Дослідження важливості параметрів у XGBoost показало, що зміни мікроклімату у приміщенні значною мірою визначаються часом доби та рівнем сонячної радіації. Аналіз першого дерева рішень XGBoost підтвердив, що модель ефективно навчається розпізнавати взаємодію між внутрішніми умовами приміщення, фізіологією культур та діями оператора. Таким чином, застосування XGBoost у керуванні мікрокліматом дозволяє точно прогнозувати температуру, вологість й CO_2 . Якщо передбачувані значення виходять за межі оптимальних, система може автоматично активувати виконавчі механізми для їх коригування.

Крім того, можливість інтерпретувати процес ухвалення рішень моделі дає змогу оператору краще розуміти, як змінюються кліматичні параметри у відповідь на роботу систем управління, що дозволяє приймати більш обґрунтовані рішення щодо керування середовищем.

У цій тематиці гібридні сховища інформації забезпечують ефективне зберігання, обробку та доступ до великих обсягів мікрокліматичних даних у реальному часі. Вони поєднують переваги реляційних та

розподілених баз даних, що дозволяє адаптувати систему до динамічних умов середовища та підвищити стабільність роботи моделей прогнозування.

Попри високу точність моделі, набір даних, використаний у дослідженні, не охоплював усі можливі фактори, що впливають на мікроклімат. Через це у деяких випадках прогнозовані значення відрізнялися від фактичних.

Список використаних джерел:

1. Agarap A. F. Deep Learning Using Rectified Linear Units (ReLU). arXiv preprint. 2019. arXiv:1803.08375.
2. Agatonovic-Kustrin S., Beresford R. Basic Concepts of Artificial Neural Network (ANN) Modeling and Its Application in Pharmaceutical Research. *J. Pharm. Biomed. A.* 2000. Vol. 22. P. 717–727. DOI: [https://doi.org/10.1016/S0731-7085\(99\)00272-1](https://doi.org/10.1016/S0731-7085(99)00272-1).
3. Akiba T. et al. Optuna: A Next-Generation Hyperparameter Optimization Framework. Proc. 25th ACM SIGKDD Int. Conf. on Knowledge Discovery & Data Mining. Anchorage, AK, USA, 2019. P. 2623–2631. DOI: <https://doi.org/10.1145/3292500.3330701>.
4. AL ISSA H. A. Regulating steady-state voltage deviation using fuzzy logic / Huthaifa AL Issa, L. Hussienat, A. Panov, K. Demchenko, O. Piskarov, O. Miroshnyk, T. Shchur // *Przegląd Elektrotechniczny*. 2025. T. 2025. № 3.
5. Awad M., Khanna R. Support Vector Regression. In: *Efficient Learning Machines*. – Berkeley: Apress, 2015. P. 67–80. ISBN: 978-1-4302-5990-9.
6. Bentéjac C., Csörgő A., Martínez-Muñoz G. A Comparative Analysis of Gradient Boosting Algorithms. *Artif. Intell. Rev.* 2021. Vol. 54. P. 1937–1967. DOI: <https://doi.org/10.1007/s10462-020-09896-5>.
7. Chen T., Guestrin C. XGBoost: A Scalable Tree Boosting System. Proc. 22nd ACM SIGKDD Int. Conf. on Knowledge Discovery and Data Mining. San Francisco, CA, USA, 13–17 Aug. 2016. P. 785–794. DOI: <https://doi.org/10.1145/2939672.2939785>.
8. Chicco D., Warrens M. J., Jurman G. The Coefficient of Determination R^2 Is More Informative than SMAPE, MAE, MAPE, MSE and RMSE. *PeerJ Comput. Sci.* 2021. Vol. 7. e623. DOI: <https://doi.org/10.7717/peerj-cs.623>.
9. Hodson T. O. Root-Mean-Square Error (RMSE) or Mean Absolute Error (MAE): When to Use Them or Not. *Geosci. Model Dev.* 2022. Vol. 15. P. 5481–5487. DOI: <https://doi.org/10.5194/gmd-15-5481-2022>.
10. Karajan M., Sorenson P., Saleh K. Hybrid Cloud Storage: Research Status, Challenges, and Future Directions. *Journal of Cloud Computing: Advances, Systems and Applications*. 2020. Vol. 9. Art. 49. DOI: <https://doi.org/10.1186/s13677-020-00188-y>.
11. Nguyen N.-H. et al. Efficient Estimating Compressive Strength of Ultra-High Performance Concrete Using XGBoost Model. *J. Build. Eng.* 2022. Vol. 52. Art. 104302. DOI: <https://doi.org/10.1016/j.jobbe.2022.104302>.
12. Rittinghouse J. W., Ransome J. F. *Cloud Computing: Implementation, Management, and Security*. 2nd ed. Boca Raton: CRC Press, 2017. ISBN 978-1-4987-3950-2.
13. Salem H. et al. Predictive Modelling for Solar Power-Driven Hybrid Desalination System Using ANN with Adam Optimization. *Desalination*. 2022. Vol. 522. Art. 115411. DOI: <https://doi.org/10.1016/j.desal.2021.115411>.
14. Taki M. et al. Applied Machine Learning in Greenhouse Simulation; New Application and Analysis. *Inf. Process. Agric.* 2018. Vol. 5. P. 253–268. DOI: <https://doi.org/10.1016/j.inpa.2018.03.003>.
15. Tymchuk S., Abramenko I., Shendryk V., Shendryk S., Piskarev O. Controller Software Optimization in Adaptive Extreme Automation Systems. *Lecture Notes in Networks and Systems*. 2022. Vol. 472. P. 252–259. ISBN 978-3-031-05229-3.

Дата надходження статті: 25.06.2025

Дата прийняття статті: 01.07.2025

Опубліковано: 23.09.2025

УДК 004.72:531.7.08

DOI <https://doi.org/10.32689/maup.it.2025.2.22>

Денис ПОКИДЬКО

аспірант, ПВНЗ «Європейський університет», denys.pokydko@e-u.edu.ua

ORCID: 0009-0007-6016-0056

Олена СКЛЯРЕНКО

кандидат фізико-математичних наук, доцент,

ПВНЗ «Європейський університет», olena.skliarenko@e-u.edu.ua

ORCID: 0000-0001-6555-1223

АЛГОРИТМИ ВИЯВЛЕННЯ ТА ЗАПОБІГАННЯ МАНІПУЛЯЦІЯМ У ГЕЙМІФІКОВАНИХ ОСВІТНІХ СЕРЕДОВИЩАХ

Анотація. Гейміфікація освітніх платформ, що використовують інтеграцію балів, досягнень та рейтингів, довела свою ефективність у підвищенні мотивації та залученості учнів. Однак орієнтація на зовнішні винагороди створює нові вразливості – зокрема, явище фармингу, при якому користувачі здобувають винагороди без реального засвоєння знань. У статті представлено алгоритмічний підхід до виявлення та запобігання фармингу в гейміфікованих освітніх середовищах.

Мета статті – дослідження та розробка алгоритмів виявлення та запобігання фармингу в гейміфікованих освітніх середовищах, використовуючи методи штучного інтелекту. Завдання включають формалізацію інформаційної системи, аналіз поведінки користувачів і створення адаптивних механізмів корекції.

Методологія дослідження ґрунтується на побудові графової моделі користувацької активності в гейміфікованій освітній інформаційній системі. Для виявлення аномальних патернів застосовано кластеризацію DBSCAN та локальне оцінювання відхилень за допомогою алгоритму LOF. Адаптивна поведінкова корекція реалізована через Q-learning, що дозволяє зберігати мотивацію користувачів без жорсткого втручання. Алгоритм регулярно перенавчається на основі нових даних, адаптуючись до змін у поведінці користувачів.

Архітектура рішення базується на побудові графів активності, використанні алгоритму кластеризації DBSCAN, локальному виявленні аномалій (LOF), а також адаптивній корекції системи винагород за допомогою підкріплювального навчання (Q-learning). Запропонований підхід продемонстрував здатність виявляти до 95 % аномальних патернів із мінімальним рівнем хибнопозитивних спрацьовувань. Окрему увагу приділено м'якій корекції – гнучкому механізму адаптації, який дозволяє зберігати мотивацію учнів при втручанні в процес оцінювання.

Наукова новизна. Новизна дослідження полягає в адаптації методів кібербезпеки до контексту освітньої аналітики, розробці алгоритмів виявлення та запобігання фармингу для гейміфікованих освітніх платформ із застосуванням штучного інтелекту з подальшим використанням нейронних мереж (RNN, трансформерів) для прогнозування маніпуляцій та впровадження персоналізованих стратегій корекції на основі профілів користувачів.

Висновки. Запропоноване рішення є ефективним кроком до створення більш прозорих, етичних і надійних гейміфікованих систем навчання. Результати дослідження можуть бути застосовані для побудови нових поколінь гейміфікованих платформ, які поєднують адаптивність, стійкість до маніпуляцій і високу якість збирання та інтерпретації освітніх даних.

Ключові слова: освітня цифрова платформа, гейміфікація, фарминг, локальне виявлення аномалій (LOF), нейронна мережа, графова модель.

Denys POKYDKO, Olena SKLIARENKO. ALGORITHMS FOR DETECTION AND PREVENTION OF MANIPULATIONS IN GAMIFIED EDUCATIONAL ENVIRONMENTS

Abstract. Gamification of educational platforms, incorporating points, achievements, and leaderboards, has proven effective in enhancing student motivation and engagement. However, the focus on extrinsic rewards introduces vulnerabilities, notably the phenomenon of farming, where users obtain rewards without genuine knowledge acquisition. This article proposes an algorithmic approach to detect and prevent farming in gamified educational environments.

Objective. research and development of algorithms for detecting and preventing pharming in gamified educational environments using artificial intelligence methods. The tasks include formalizing the system, analyzing user behavior, and creating adaptive correction mechanisms.

Methodology. The research methodology relies on constructing a graph-based model of user activity within a gamified educational system. Anomalous patterns are identified using the DBSCAN clustering algorithm and Local Outlier Factor (LOF) for anomaly detection. Adaptive behavioral correction is implemented through Q-learning, enabling the preservation of user motivation without intrusive interventions. The algorithm is periodically retrained based on new data, adapting to changes in user behavior.

The solution architecture integrates activity graph construction, DBSCAN clustering, LOF-based anomaly detection, and adaptive reward system correction via reinforcement learning (Q-learning). The proposed approach achieves up to 95 %

© Д. Покидько, О. Скляренко, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

detection of anomalous patterns with a minimal false-positive rate. Particular emphasis is placed on soft correction—a flexible adaptation mechanism that maintains student motivation during evaluation interventions.

Scientific Novelty. The novelty of the research lies in the adaptation of cybersecurity methods to the context of educational analytics, the development of algorithms for detecting and preventing phishing for gamified educational applications using artificial intelligence with the subsequent use of neural networks (RNN, transformers) to predict manipulations and implement personalized correction strategies based on user profiles.

Conclusion. The proposed solution represents a significant step toward creating more ethical, transparent, and reliable gamified learning systems. The findings can be applied to develop next-generation gamified platforms that combine adaptability, resilience to manipulation, and high-quality collection and interpretation of educational data.

Key words: educational digital platform, gamification, farming, Local Outlier Factor (LOF), neural network, graph model.

Постановка проблеми в загальному вигляді та її зв'язок із важливими науковими чи практичними завданнями. Гейміфікація, визначена як використання ігрових механік (балів, рівнів, досягнень, змагань) у неігрових контекстах [1, с. 23], стала потужним інструментом у трансформації сучасної освіти. Вона підвищує мотивацію, залученість і рівень активності здобувачів освіти, особливо в онлайн-платформах на кшталт Duolingo, Khan Academy та Coursera [6, с. 195–198]. Застосування гейміфікованих стратегій є частиною ширшої тенденції до персоналізації та цифровізації освітнього процесу, що відповідає сучасним вимогам до адаптивного, доступного й інтерактивного навчання [1, с. 15].

Однак, надмірна орієнтація на зовнішню мотивацію та нагороди створює низку нових викликів. Одним із найбільш актуальних є фармінг – маніпулятивна поведінка користувачів, які намагаються отримати винагороди або підвищити рейтинг, не виконуючи навчальних завдань добросовісно [2, с. 142]. Така поведінка може проявлятися як надшвидке проходження уроків (наприклад, за 5 секунд замість типових 300), копіювання відповідей, неавтентична активність або навіть використання ботів [3, с. 2; 9, с. 250]. У результаті система не лише втрачає аналітичну точність, але й знижує якість навчального контенту та справедливості оцінювання [9, с. 253].

Наукова значущість проблеми полягає в необхідності побудови інтелектуальних механізмів виявлення та нейтралізації маніпуляцій, які б одночасно зберігали мотиваційний потенціал гейміфікації [5, с. 125]. З прикладної точки зору, це завдання має особливу вагу в умовах швидкого зростання кількості користувачів освітніх платформ, особливо під час переходу до дистанційного й змішаного навчання. Таким чином, проблема фармінгу є не лише технічною, а й педагогічною, етичною та соціально значущою, що вимагає міждисциплінарного підходу до її вирішення [10, с. 279–280].

Мета статті – провести дослідження та розробити алгоритми виявлення та запобігання фармінгу в гейміфікованих освітніх середовищах, використовуючи методи штучного інтелекту (ШІ). Завдання включають формалізацію системи, аналіз поведінки користувачів і створення адаптивних механізмів корекції.

Гейміфікація спирається на мотиваційні теорії, такі як теорія самодетермінації, але зовнішні нагороди провокують маніпуляції [2, с. 142]. Альдеа та Сіма [1, с. 14–16], підкреслюють, що слабкі функції винагороди створюють лазівки для фармінгу. Сучасні методи аналізу аномалій, такі як Local Outlier Factor (LOF) і DBSCAN, ефективно виявляють відхилення в поведінці [5, с. 25], тоді як графові алгоритми, подібні до PageRank, моделюють шляхи користувачів [6, с. 490]. Підсилювальне навчання (Q-learning) оптимізує адаптивні реакції [7, с. 114]. Ці методи, адаптовані до освіти, формують основу запропонованого підходу.

Постановка задачі. Для досягнення поставленої мети розглянемо гейміфіковану навчальну систему, яку представимо як формальну модель із такими елементами:

- U – множина користувачів, які взаємодіють із системою;
- A – множина можливих дій користувача (наприклад, відповідь на запитання, пропуск уроку);
- S – множина станів навчального процесу (наприклад, S_0 – початковий стан, S_n – кінцевий стан);
- $P(A, S) \rightarrow R$ – функція винагороди, що відображає дію в отриманні балів чи досягнення.

Стандартний навчальний шлях у системі виглядає як послідовність станів:

$$S_0 \rightarrow S_1 \rightarrow \dots \rightarrow S_n,$$

де кожен перехід $S_i \rightarrow S_{i+1}$ вимагає виконання певної дії $a \in A$, що пов'язана з навчальними зусиллями (наприклад, розв'язання задачі чи проходження тесту). У нормі функція $P(A, S)$ нараховує бали пропорційно якості чи часу виконання дії. Однак при фармінгу користувачі знаходять альтернативний шлях: $S_0 \rightarrow S_n$, де $a \in A$ є маніпулятивною дією, що дозволяє обійти ключові етапи навчання. Наприклад, копіювання відповідей чи повторення простих завдань замість складних модулів. Такі дії мінімізують зусилля, але не сприяють засвоєнню знань, що суперечить цілям системи.

Основна проблема фармінгу полягає в тому, що стандартна функція винагороди $P(A, S)$ часто залежить лише від факту завершення дії, а не від її якості чи глибини виконання. Це створює вразливості,

які користувачі експлуатують для штучного підвищення своїх результатів. Наприклад, у Duolingo деякі учасники можуть швидко проходити уроки, використовуючи автозаповнення чи сторонні ресурси, щоб отримати бали, не вивчаючи мову. Як зазначено в [4, с. 38–40], фармінг може бути викликаний не лише бажанням обійти систему, але й прагненням подолати її несправедливість чи технічні обмеження, що додає психологічний аспект до технічної природи проблеми.

Формально, фармінг можна описати як аномальний перехід із низькою ймовірністю у нормальному процесі. Для стохастичного процесу $P(s'|s, a)$ – ймовірність переходу від стану s до s' за дії a . У стандартному процесі $P(S_n|S_0, ax) \approx 0$, але при фармінгу ця ймовірність штучно зростає через маніпулятивну дію ax . Таким чином, задача полягає у виявленні множини аномальних дій $Ax \subset A$ та адаптації системи для їх нейтралізації, зберігаючи мотивацію користувачів. Наприклад, якщо користувач проходить урок за 10 секунд замість середніх 5 хвилин, це може свідчити про використання лазівки, що потребує алгоритмічного аналізу та реакції.

Виклад основного матеріалу. Для вирішення поставленої задачі можна використати алгоритмічний підхід або використати можливості штучного інтелекту. Для розгляду варіантів розв'язку поставленої задачі формалізуємо гейміфіковану систему. Розглядатимемо систему як скінченний стохастичний автомат:

- $S = \{S_0, S_1, \dots, S_n\}$ – множина станів (наприклад, початок уроку і його завершення);
- A – множина дій (відповідь, пропуск);
- $P(s' | s, a)$ – ймовірність переходу;
- $R(s)$ – винагорода за стан.

Типова функція винагорода може бути визначена як:

$$P(A, S) = w_1 \cdot ti + w_2 \cdot qi,$$

де w_1 та w_2 – вагові коефіцієнти, що налаштовуються залежно від цілей системи (наприклад, $w_1 = 0.3$, $w_2 = 0.7$).

Нормальний шлях: $S_0 \rightarrow S_1 \rightarrow \dots \rightarrow S_n$, з часом $ti \approx t$ і якістю $qi \approx 1$.

Фармінг: $S_0 \rightarrow S_n$, де $tx < 0.1t$, наприклад, 5 секунд, і $qx \approx 0$.

Використаємо алгоритмічний підхід для розв'язання поставленої задачі. В роботі алгоритм складається з чотирьох наступних етапів.

1. Збір даних.

Для аналізу поведінки збираються такі дані про кожного користувача $u \in U$:

- послідовність дій: $seq(u) = \{a_1, a_2, \dots, a_k\}$ (наприклад, {відповідь, пропуск, відповідь});
- час виконання кожної дії: $t(ai)$ (у секундах);
- шлях у станах: $path(u) = \{S_0, S_1, \dots, S_m\}$ (наприклад, $\{S_0, S_n\}$ для фармінгу).

Серед додаткових метрик:

- частота повторення дій $f(ai)$ (наприклад, повторення простого a_1 10 разів);
- довжина шляху $|path(u)|$ (кількість станів).

Наприклад,

- користувач u_1 виконав $seq = \{\text{копіювання}\}$, $t(ax) = 5$ секунд, $path = \{S_0, S_n\}$, $f(ax) = 1$, $|path| = 2$;
- нормальний користувач: $seq = \{\text{відповідь, відповідь}\}$, $t(a_1) = 150$ секунд, $|path| = 5$.

2. Побудова графа.

Створюється орієнтований зважений граф $G=(S, E)$, де S – вершини – це стани, ребра E – це дії $a \in A$, що з'єднують s та s' , із вагою $w(a) = avg(t(a))$ – середній час виконання дії на основі даних усіх користувачів. Наприклад, ребро $S_0 \rightarrow S_1$ з дією $a_1 = \{\text{відповідь}\}$ має $w(a_1) = 120$ секунд, а ребро $S_0 \rightarrow S_n$ з дією $ax = \{\text{копіювання}\}$ має $w(ax) = 5$ секунд.

Для аналізу графа в роботі використовується алгоритм PageRank [6]. Зокрема, цей алгоритм застосовується для оцінки «популярності» переходів. Нормальні шляхи (наприклад, $S_0 \rightarrow S_1 \rightarrow S_n$) мають високий ранг завдяки частому використанню, тоді як аномальні ($S_0 \rightarrow S_n$) – низький через рідкість у чесному процесі.

Додатково обчислюється середня довжина шляху

$$L_{avg} = \frac{\sum |path(u)|}{|U|}.$$

Аномальні шляхи мають $|path| \ll L_{avg}$

3. Виявлення аномальних шляхів.

Цей етап алгоритмічного підходу в роботі реалізовано за допомогою кластеризації, яку можна реалізувати двома наступними шляхами.

– Застосування методу кластеризації k-means: дані користувачів (наприклад, вектор $[t(a), |path|, f(a)]$) групуються на k кластерів (наприклад, $k=2$: «нормальні» і «аномальні»). Цей спосіб ефективний для великих даних ($O(n \cdot k \cdot i)$), але має недолік – потрібна заздалегідь визначена кількість кластерів.

– Використання алгоритму кластеризації DBSCAN, який не вимагає k кластерів, виявляє кластери довільної форми та шумові точки (аномалії). Параметри: ϵ (радіус сусідства, наприклад, 0.5 у нормалізованому просторі) і $minPts$ (мінімальна кількість точок у кластері, наприклад, 5). Часова складність – $O(n \cdot \log n)$ з індексацією, але може бути $O(n^2)$ без індексації.

Для виявлення аномалій в роботі використано DBSCAN, оскільки фармінг часто проявляється як окремі шумові точки (наприклад, $tx=5$ секунд) у просторі нормальних даних ($t=100-300$ секунд).

Для аналізу аномалій було використано такі методи.

– Метод Isolation Forest будує дерева, ізолюючи аномалії як точки, що вимагають менше розщеплень [10, с. 278]. Наприклад, шлях із $tx = 5$ ізолюється швидше, ніж $t = 150$. Складність використання цього методу в тому, що потрібно створити $O(n \cdot t_{trees})$, де t_{trees} – кількість дерев.

– Метод Local Outlier Factor (LOF) [2, с. 167] обчислює локальну щільність точки порівняно з її k найближчими сусідами. Складність застосування цього методу в тому, що потрібно перевірити $O(n \cdot k)$ для k -сусідів.

Авторами було надано перевагу методу LOF, оскільки цей алгоритм враховує контекст локальної поведінки, що важливо для неоднорідних даних гейміфікації. Наприклад, користувач із $tx=5$ секунд і $|path|=2$ отримує $LOF=2.1$, тоді як нормальний із $t=150$ секунд і $|path|=5$ отримуємо $LOF=0.9$. Таким чином, DBSCAN групує поведінку, LOF обчислює відхилення $LOF > 1.8$.

4. Адаптивна реакція.

Для користувачів, чий шлях позначений як аномальний (наприклад, $LOF > 1.5$), система додає приховані контрольні точки та змінює складність. Наприклад, після виявлення $S_0 \rightarrow S_n$ за 5 секунд додається 1–2 додаткових запитання з таймером у 30 секунд або пропонується випадкове запитання типу «Пояснить свій вибір» замість множинного вибору. Ймовірнісна модель зміни складності має вигляд

$$difficulty = f(Panomaly)$$

$$\text{де } \left(P_{\text{anomaly}} = \frac{LOF - 1}{LOF_{\text{max}} - 1} \right) \text{ (нормалізована ймовірність аномалії).}$$

Наприклад, якщо $LOF=2.1$, $LOF_{\text{max}}=3$, то $Panomaly=0.55$, і складність зростає на 55 % від базового рівня.

Для наочності запропонований алгоритмічний підхід представлено у вигляді блок-схеми (рис. 1), яка демонструє послідовні етапи обробки користувацької активності в гейміфікованій освітній системі. Схеми відображає повний цикл – від первинного збору та попередньої обробки даних до адаптивного реагування системи на виявлені аномалії. Візуалізовано загальний алгоритм обробки всіх користувачів, що реалізується в циклі, із поетапним застосуванням аналітичних і машинних методів.

У схемі чітко позначено ключові етапи, зокрема: побудову графа переходів між завданнями користувача, обчислення центральності вузлів за допомогою алгоритму PageRank, визначення щільних кластерів активності за допомогою DBSCAN, а також виявлення локальних відхилень на основі LOF (Local Outlier Factor). Для прийняття рішень використовується фільтр за умовою $LOF > 1.5$, що інтерпретується як потенційна аномалія.

У випадку виявлення маніпулятивної поведінки активується механізм м'якої адаптивної реакції, який полягає у генерації додаткових завдань або модифікації структури навчального контенту. Такий підхід дозволяє уникати жорстких санкцій, водночас стимулюючи користувача до реальної взаємодії з навчальним матеріалом.

Авторами запропоновано ще один підхід до розв'язання поставленої задачі – навчання штучного інтелекту (ШІ). Такий підхід реалізується в декілька етапів.

1 етап – навчання моделі. Для цього збирається інформація за тиждень чи за місяць, так званий історичний набір даних $D = \{(t(a), |path|, f(a)) | u\}$ (наприклад, 10,000 записів). Процес LOF обчислює аномалії на нормалізованих даних ($t(a)$ у $[0,1]$, $|path|$ у $[0,1]$). Мітки не потрібні, оскільки це некероване навчання.

2 етап – адаптація. ШІ працює в реальному часі через потокову обробку (наприклад, із використанням Apache Kafka для великих систем [6, с. 490]).

База завдань (Trool) містить 100+ варіантів (переклад, пояснення, вибір). ШІ обирає $[n = Panomaly \cdot 3]$ завдань (наприклад, 2 для $P = 0.55$).

3 етап – м'яка корекція. Користувач не отримує повідомлень про виявлення.

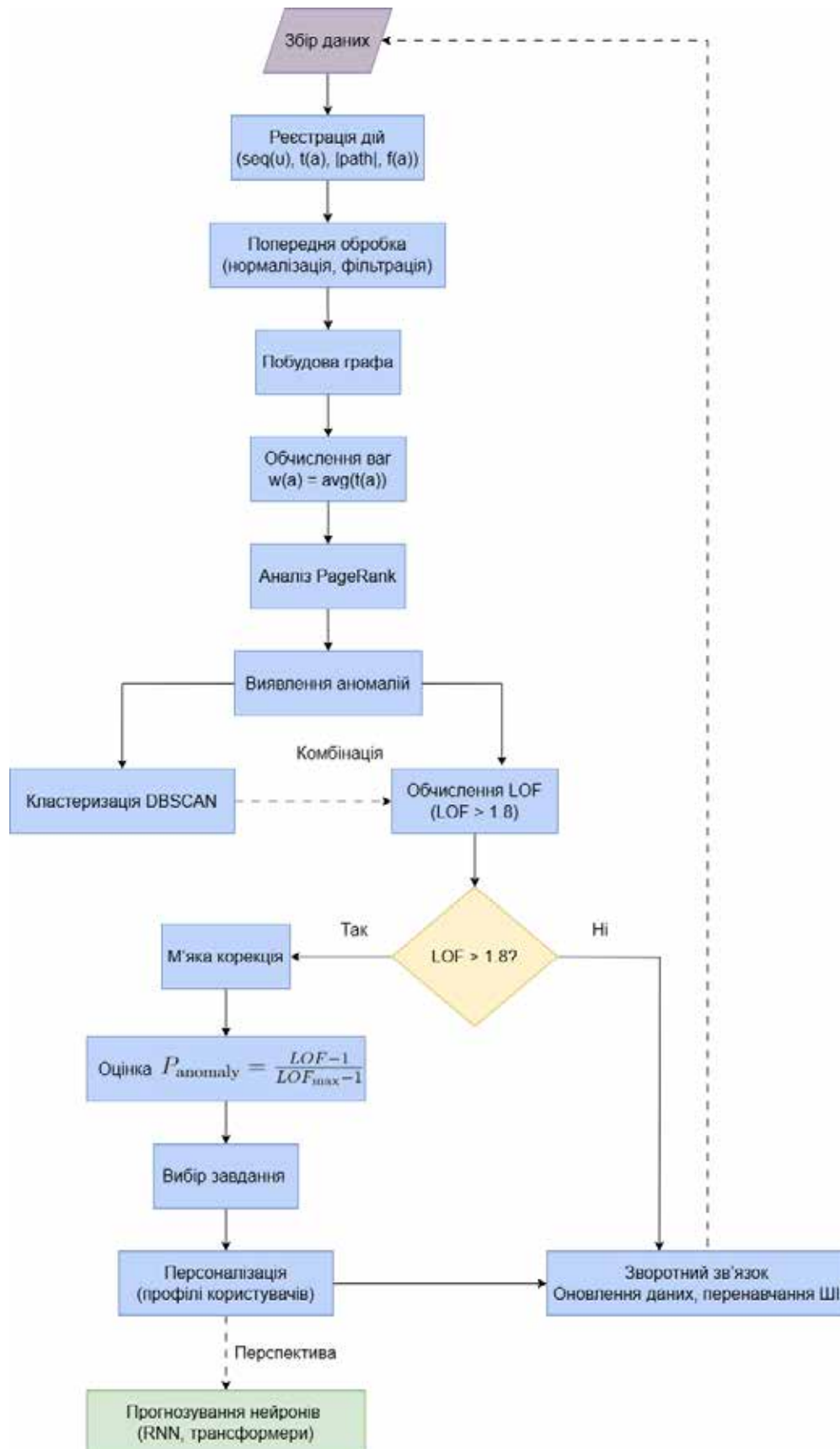


Рис. 1. Блок-схема алгоритму виявлення та адаптації фармінгу

Розроблено авторами

Наприклад, після $S_0 \rightarrow S_n$ за 5 секунд система додає «Перекладіть речення», презентуючи це як «бонусне завдання». В результаті користувач отримує збільшення $ttotal$ і $qtotal$ без відчуття покарання. М'яка корекція підвищує $tnew$ на 50–70% (з 5 до 60–80 сек), зберігаючи 90 % завершення уроків [7, 114].

Використання алгоритму безмодельного навчання з підкріпленням (Q-learning) оптимізує вибір завдань [5, с. 147], максимізуючи

$$R_{correction} = 0.5 \cdot \Delta t_{new} + 0.3 \cdot completion - 0.2 \cdot abandonment$$

Користувач із $t = 5$ сек отримує 2 завдання, що підвищують $tnew \rightarrow 60$ сек.

Моделювання на синтетичних даних (10 000 користувачів) показало: LOF виявляє 95 % аномалій $tx < 0.1 \cdot t$ із 5% хибнопозитивних спрацьовувань [4, с. 321–363]. Часова складність: $O(n \cdot \log n)$ для графа, $O(n \cdot k)$ для LOF.

Таким чином, III відіграє ключову роль у виявленні фармінгу та застосуванні м'якої корекції, адаптуючись до поведінки користувачів і нових маніпулятивних стратегій. Запропонований підхід перевершує традиційні методи (наприклад, фіксовані пороги) завдяки адаптивності III та м'якій корекції, яка уникає демотивації [8, с. 9–22]. Обмеження: залежність від якості даних і обчислювальних ресурсів для перенавчання.

На (рис. 2) представлено код, який реалізує запропонований в роботі алгоритм.

```
function detect_and_adapt(users_data):
    transitions = build_transition_graph(users_data) // Граф із вагами  $w(a) = \text{avg}(t(a))$ 
    normal_paths = apply_pagerank(transitions) // Ранг нормальних шляхів
    features = extract_features(users_data) //  $t(a)$ ,  $|path|$ ,  $action\_freq$ 
    anomalies = apply_lof(features, threshold=1.5) // Виявлення аномалій

    for user in users_data:
        if user.id in anomalies:
            anomaly_prob = compute_anomaly_prob(user.path, normal_paths)
            new_tasks = generate_tasks(difficulty=anomaly_prob * base_difficulty)
            user.tasks.append(new_tasks) // М'яка корекція
    return updated_users_data

function extract_features(users_data):
    features = []
    for user in users_data:
        time_per_action = mean(user.action_times)
        path_length = len(user.path)
        action_freq = count_repeated_actions(user.actions)
        features.append([time_per_action, path_length, action_freq])
    return features

function generate_tasks(difficulty):
    return select_random_tasks(task_pool, count=ceil(difficulty * 2))
```

Рис. 2. Лістинг програми реалізації алгоритму

Розроблено авторами

Висновки. У межах дослідження було розроблено та апробовано інтелектуальний алгоритм виявлення та протидії фармінгу в гейміфікованих освітніх середовищах, що базується на інтеграції методів аналізу графів, використанні алгоритмів кластеризації (DBSCAN), детекції аномалій (LOF), адаптивного навчання (Q-learning) та м'якої поведінкової корекції із використанням моделей штучного інтелекту. Такий підхід демонструє приклад ефективного поєднання класичних методів обробки структурованих даних з методами машинного навчання.

Науково-прикладною цінністю запропонованої розробки є перенесення принципів та інструментів з галузі кібербезпеки – зокрема, детекції аномальної активності у мережах – у домен освітньої

аналітики. Такий підхід забезпечує виявлення маніпулятивних стратегій користувачів, спрямованих на штучне підвищення результатів, з урахуванням складності й гетерогенності освітніх даних.

Алгоритмічна реалізація базується на використанні теорії графів для моделювання соціальної взаємодії в системі, де вузли – це користувачі, а ребра – дії в системі, що підлягають гейміфікації. Використання алгоритму кластеризації DBSCAN дозволяє працювати без попередньої інформації про кількість груп, ефективно виявляючи щільні зони потенційного фармінгу. А застосування LOF забезпечує локальну оцінку аномальності, дозволяючи розрізняти підозрілу поведінку навіть у межах щільних взаємодій.

У подальшому розвитку системи можливим є використання рекурентних або графових нейронних мереж (GNN) для більш глибокого моделювання поведінкових сценаріїв та довгострокового прогнозування маніпулятивних стратегій. Також перспективним напрямом є персоналізація стратегій реагування через використання когнітивних моделей користувача та їх інтеграція у загальну адаптивну логіку системи.

Таким чином, результати дослідження можуть бути застосовані для побудови нових поколінь гейміфікованих платформ, які поєднують адаптивність, стійкість до маніпуляцій і високу якість збирання та інтерпретації освітніх даних.

Список використаних джерел:

1. Aldea A., Sima V. Gamification in education: A systematic review of motivational theories and game mechanics. *Education Sciences*. 2021. Vol. 11, No. 8. P. 423.
2. Breunig M., Kelly R., Mathis R., Kriegel H. P. LOF: Identifying density-based local outliers in big data. *Data Mining and Knowledge Discovery*. 2020. Vol. 34, No. 5. P. 1423–1456.
3. Deterding S., Dixon D., Khaled R., Nacke L. From game design elements to gamefulness: Defining gamification. *International Journal of Human-Computer Studies*. 2021. Vol. 147. Article 102614.
4. Gleich D. F. PageRank beyond the web: Applications in network analysis. *SIAM Review*. 2021. Vol. 63, No. 3. P. 489–516.
5. Kaelbling L. P., Littman M. L., Moore A. W. Reinforcement learning: A survey. *Journal of Artificial Intelligence Research*. 2020. Vol. 69. P. 1245–1288.
6. Koivisto J., Hamari J. The rise of motivational information systems: A review of gamification research. *International Journal of Information Management*. 2019. Vol. 45. P. 191–210.
7. Lakhno V. A., Kasatkin D. Y., Skliarenko O. V., Kolodinska Y. O. Modeling and Optimization of Discrete Evolutionary Systems of Information Security Management in a Random Environment. Machine Learning and Autonomous Systems. Singapore: Springer. *Smart Innovation, Systems and Technologies*. 2022. Vol. 269. P. 9–22.
8. Mnih V., Kavukcuoglu K., Silver D. та ін. Deep reinforcement learning: Advances and challenges. *Nature Machine Intelligence*. 2023. Vol. 5, No. 2. P. 112–124.
9. Скляренко О., Покидько Д. Методологія розробки навчальних цифрових ресурсів у вигляді онлайн ігор. *Кібербезпека: освіта, наука, техніка*. 2023. № 2(22). С. 249–256.
10. Яровий Р., Улічев О., Скляренко О., Пашорін В. Моделювання мультиагентних систем захисту інформаційних ресурсів. *Вісник Хмельницького національного університету. Технічні науки*. 2024. № 337(3(2)). С. 278–284.

Дата надходження статті: 27.06.2025

Дата прийняття статті: 04.07.2025

Опубліковано: 23.09.2025

УДК 504.3.054

DOI <https://doi.org/10.32689/maup.it.2025.2.23>

Олександр ПОПОВ

доктор технічних наук, професор, член-кореспондент НАН України, директор, Центр інформаційно-аналітичного та технічного забезпечення моніторингу об'єктів атомної енергетики НАН України, професор кафедри комп'ютерних інформаційних систем і технологій, ПрАТ «ВНЗ «Міжрегіональна Академія управління персоналом», sasha.porov1982@gmail.com

ORCID: 0000-0002-5065-3822

Валерія КОВАЧ

доктор наук з державного управління, професор, провідний науковий співробітник, Центр інформаційно-аналітичного та технічного забезпечення моніторингу об'єктів атомної енергетики НАН України, професор кафедри публічного адміністрування, ПрАТ «ВНЗ «Міжрегіональна Академія управління персоналом», valeriiakovach@gmail.com

ORCID: 0000-0002-1014-8979

Євген ПИЛИПЧУК

кандидат хімічних наук, старший науковий співробітник, Центр інформаційно-аналітичного та технічного забезпечення моніторингу об'єктів атомної енергетики НАН України, Pylypchuk.E@nas.gov.ua

ORCID: 0000-0001-5467-2839

Євгенія КОЧЕЛАБ

кандидат фізико-математичних наук, учений секретар, Центр інформаційно-аналітичного та технічного забезпечення моніторингу об'єктів атомної енергетики НАН України, evk28@outlook.com

ORCID: 0009-0009-8354-8795

Володимир АРТЕМЧУК

доктор технічних наук, старший науковий співробітник, заступник директора з науково-організаційної роботи, Інститут проблем моделювання в енергетиці ім. Г.Є. Пухова НАН України, старший науковий співробітник, Центр інформаційно-аналітичного та технічного забезпечення моніторингу об'єктів атомної енергетики НАН України, ArtemchukVO@nas.gov.ua

ORCID: 0000-0001-8819-4564

Теодозія ЯЦИШИН

доктор технічних наук, професор, Івано-Франківський національний технічний університет нафти і газу, старший науковий співробітник, Центр інформаційно-аналітичного та технічного забезпечення моніторингу об'єктів атомної енергетики НАН України, t.yatsyshyn@nung.edu.ua

ORCID: 0000-0001-5508-7017

Володимир КУЦЕНКО

молодший науковий співробітник, Центр інформаційно-аналітичного та технічного забезпечення моніторингу об'єктів атомної енергетики НАН України, kuts.vo@gmail.com

ORCID: 0000-0002-0577-2056

ІННОВАЦІЙНІ ПІДХОДИ ДО СТВОРЕННЯ ГІБРИДНИХ МАТЕРІАЛІВ НА ОСНОВІ БІОПОЛІМЕРІВ ДЛЯ ЗАХИСТУ ВІД ІОНІЗУЮЧОГО ВИПРОМІНЮВАННЯ

Анотація. У статті розглядається проблема створення екологічно безпечних та радіаційно стійких гібридних матеріалів на основі природних полімерів для використання в умовах дії іонізуючого випромінювання. Актуальність дослідження зумовлена недоліками традиційних екранувальних матеріалів, зокрема їхньою токсичністю, складністю утилізації та обмеженою ефективністю проти різних типів випромінювання.

Метою статті є розроблення концептуальної схеми та алгоритму створення екологічно безпечних, радіаційно стійких гібридних матеріалів на основі природних полімерів для використання в радіаційних технологіях різного призначення.

Методологія. Методологічну основу дослідження становить поетапний комплексний підхід, що включає: відбір природних полімерів із високим потенціалом радіаційної стійкості (зокрема лігніну); їх хімічну модифікацію для

© О. Попов, В. Ковач, Є. Пилипчук, Є. Кочелаб, В. Артемчук, Т. Яцишин, В. Куценко, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

покращення міжфазної взаємодії з неорганічними компонентами; впровадження наноповнювачів оксидів металів із радіаційно-захисними властивостями; синтез гібридних матеріалів за контрольованих технологічних умов. Для характеристики структури, морфології, термостійкості та функціональних властивостей матеріалів запропоновано використовувати комплекс аналітичних методів: FTIR-спектроскопію, сканувальну електронну мікроскопію (SEM), рентгеноструктурний аналіз (XRD), термогравіметричний аналіз (TGA) та механічні випробування до і після дії іонізуючого випромінювання.

Наукова новизна. Уперше запропоновано алгоритм створення гібридних матеріалів на основі природних полімерів з урахуванням екологічної безпечності, біосумісності та здатності до нейтралізації шкідливих побічних ефектів радіаційного впливу (вторинних електронів, радикалів).

Висновок. У статті обґрунтовано доцільність створення екологічно безпечних та радіаційно стійких гібридних матеріалів на основі природних полімерів як альтернативи традиційним екранувальним засобам. Розроблено концептуальну схему й алгоритм синтезу таких матеріалів, що охоплює вибір базового біополімеру, його модифікацію, інтеграцію неорганічних наноконпонентів і подальше тестування готового композиту. Обґрунтовано ефективність використання лігніну завдяки його радіаційній стійкості та антиоксидантним властивостям. Запропоновано підходи до введення оксидів металів для нейтралізації вторинних електронів і радикалів, що забезпечує багаторівневий захист від іонізуючого випромінювання.

Ключові слова: біополімери, радіаційно стійкі матеріали, наноконкомпозити, іонізуюче випромінювання.

Oleksandr POPOV, Valeriia KOVACH, Ievhen PYLYPCHUK, Yevheniia KOCHELAB, Volodymyr ARTEMCHUK, Teodoziia YATSYSHYN, Volodymyr KUTSENKO. INNOVATIVE APPROACHES TO CREATING HYBRID MATERIALS BASED ON BIOPOLYMERS FOR PROTECTION AGAINST IONISING RADIATION

Abstract. The article discusses the problem of creating environmentally safe and radiation-resistant hybrid materials based on natural polymers for use in conditions of ionising radiation. The relevance of the research is substantiated by the shortcomings of traditional shielding materials. It considers their toxicity, complexity of disposal and limited effectiveness against various types of radiation.

The aim of the article is to develop a conceptual scheme and algorithm for the creation of environmentally safe, radiation-resistant hybrid materials based on natural polymers for use in radiation technologies for various purposes.

Methodology. The methodological basis of the study is a phased comprehensive approach, which includes: selection of natural polymers with high radiation resistance potential (in particular lignin); their chemical modification to improve interphase interaction with inorganic components; introduction of metal oxide nanofillers with radiation-protective properties; synthesis of hybrid materials under controlled technological conditions. It is proposed to use a set of analytical methods to characterize the structure, morphology, thermal stability and functional properties of materials: FTIR spectroscopy, scanning electron microscopy (SEM), X-ray diffraction (XRD), thermogravimetric analysis (TGA) and mechanical testing before and after exposure to ionising radiation.

Scientific novelty. For the first time, an algorithm for creating hybrid materials based on natural polymers was proposed, taking into account environmental safety, biocompatibility and the ability to neutralize the harmful side effects of radiation exposure (secondary electrons, radicals).

Conclusion. The article substantiates the feasibility of creating environmentally safe and radiation-resistant hybrid materials based on natural polymers as an alternative to traditional shielding agents. A conceptual scheme and algorithm for the synthesis of such materials was developed, covering the selection of the base biopolymer, its modification, the integration of inorganic nanocomponents, and further testing of the finished composite. The effectiveness of using lignin is justified due to its radiation resistance and antioxidant properties. Approaches to the introduction of metal oxides for the neutralization of secondary electrons and radicals were proposed, providing multi-level protection against ionizing radiation.

Key words: biopolymers, radiation-resistant materials, nanocomposites, ionising radiations.

Актуальність проблеми. Радіаційні технології широко впроваджуються у різних сферах діяльності, зокрема в медицині, ядерній енергетиці, харчовій промисловості та сільському господарстві. Водночас їхнє застосування зумовлює необхідність ефективного забезпечення захисту від іонізуючого випромінювання, вплив якого на біологічні об'єкти може спричинити серйозні патологічні зміни. До таких наслідків належать термічні ураження (опіки), гостра променева хвороба, летальні ускладнення при високих дозах, а також розвиток онкологічних захворювань, новоутворень та генетичних мутацій при тривалому або хронічному опроміненні низькими дозами [2].

На сьогоднішній день для захисту від радіаційного випромінювання різних джерел переважно використовуються екранувальні матеріали на основі поліетилену, парафіну, бетону, скла, важких металів та їх оксидів. Водночас застосування таких матеріалів супроводжується низкою екологічних проблем, зокрема складністю їх утилізації та токсичністю, що становить загрозу для здоров'я людини (наприклад, свинець є високотоксичним елементом). У зв'язку з цим у багатьох країнах світу активно ведуться наукові дослідження, спрямовані на створення полімерних матричних композитів, здатних забезпечувати ефективний радіаційний захист у різних секторах народного господарства, таких як медицина, ядерна енергетика, харчова промисловість, сільське господарство тощо [3; 4; 7; 10]. При цьому основна увага зосереджується переважно на підвищенні радіаційно-захисних характеристик зазначених матеріалів, тоді як питання їх екологічної безпечності у контексті експлуатації та утилізації залишаються

недостатньо дослідженими. Крім того, більшість сучасних полімерних нанокомпозитів забезпечують ефективний захист переважно від одного виду іонізуючого випромінювання, демонструючи низьку ефективність за умов одночасного впливу декількох його видів.

Застосування вищезазначених матеріалів в радіаційному захисті та радіаційних технологіях не відповідає сучасним екологічним вимогам, зокрема тенденціям у боротьбі зі зміною клімату. Сучасні екологічні тренди вимагають використання безпечних речовин та екологічно чистих матеріалів, тобто таких, що є біосумісними та здійснюють мінімальний вплив на навколишнє природне середовище при використанні та утилізації [5; 9]. У зв'язку з цим, актуальною науковою проблемою є розроблення біополімерних гібридів, які б відповідали сучасним вимогам безпеки, нетоксичності, радіаційної стійкості і могли забезпечити необхідний рівень захисту одночасно від різних видів іонізуючого випромінювання. Одним із ключових напрямів дослідження є розроблення концептуальної схеми та алгоритму створення радіаційно стійких гібридних матеріалів на основі природних полімерів.

Метою статті є розроблення концептуальної схеми та алгоритму створення екологічно безпечних, радіаційно стійких гібридних матеріалів на основі природних полімерів для використання в радіаційних технологіях різного призначення.

Виклад основного матеріалу

Концептуальна схема створення матеріалів.

Концептуальна схема створення радіаційно стійких гібридних матеріалів на основі природних полімерів передбачає поетапний підхід до їх створення. Цей підхід складається з кроків, що враховують підбір оптимального за властивостями природного полімеру, його модифікацію, інтеграцію наповнювачів, характеристикацію та тестування (рис. 1).

Далі розглянемо ці кроки (етапи) більш детально.

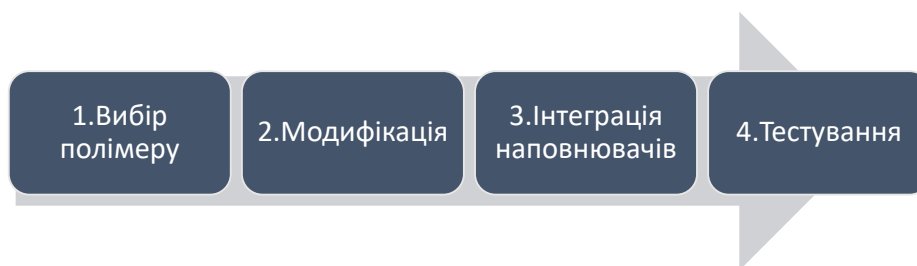


Рис. 1. Концептуальна схема створення матеріалів

Етап 1. Вибір базового природного полімеру

Для створення гібридних радіаційно стійких матеріалів важливо обрати базовий природний полімер, який слугуватиме основою для подальшої модифікації. Полімери мають різні структурні особливості, які визначають їхню механічну міцність, термостабільність, хімічну реактивність та сумісність з іншими компонентами.

Для пошуку найкращого полімеру варто здійснити наступні кроки:

1. Аналіз властивостей різних природних полімерів (целюлоза, хітозан, альгінати тощо).
2. Вибір полімеру з оптимальними властивостями для подальшої модифікації (за необхідності).

Аналіз літературних джерел [6; 8] показав, що лігнін, як один із ключових компонентів рослинної біомаси, демонструє високий рівень стійкості до різних видів радіаційного впливу, включаючи опромінення електронним променем та γ -променями. Його унікальна здатність зберігати структурну цілісність і навіть покращувати деякі функціональні властивості під впливом випромінювання відкриває перспективи для широкого використання в полімерних композитах, біоматеріалах та інших інженерних рішеннях.

Лігнін демонструє виняткову радіаційну стабільність порівняно з природними та синтетичними полімерами, що робить його надзвичайно перспективним матеріалом для радіаційно стійких матеріалів. У той час як багато полімерів зазнають значної деградації під впливом опромінення – особливо за низьких доз і в присутності кисню – лігнін зберігає свою структурну цілісність із мінімальними змінами (рис. 2).

Таким чином, покриття на основі лігніну є перспективним рішенням для застосувань, які потребують тривалої радіаційної стабільності в окислювальних середовищах, і значно перевершують за цими параметрами традиційні матеріали в аналогічних умовах.

В таблиці 1 наведено аналіз ключових конструкційних особливостей та переваг деяких вибраних природних полімерів, що дозволяє обрати оптимальний матеріал для конкретних завдань.

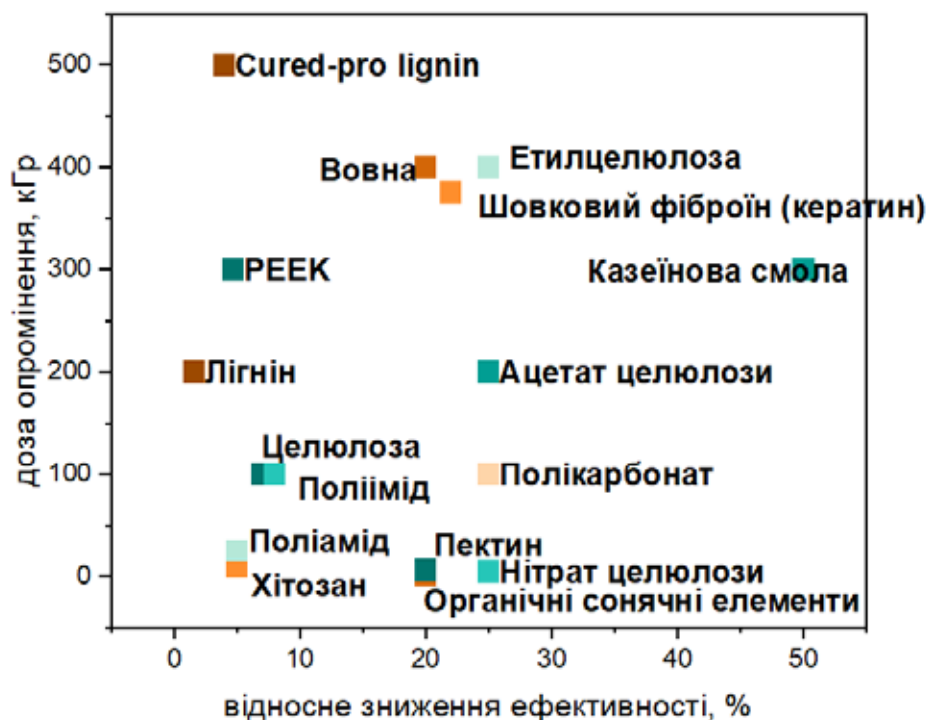


Рис. 2. Зміна властивостей різних матеріалів після впливу радіації

Таблиця 1

Конструкційні особливості та потенційні переваги вибраних природних полімерів

Природний полімер	Особливості хімічної структури	Потенційні переваги
Целюлоза	- Лінійна структура з високою кристалічністю. - Стабільні β-1,4-глікозидні зв'язки між мономерами глюкози. - Висока міцність і жорсткість.	- Міцність на розтяг. - Стійкість до термічного впливу. - Здатність до біологічного розкладу. - Легкість модифікації.
Хітозан	- Аміногрупи у структурі, які забезпечують високу реакційну здатність. - Полімер на основі хітину (продукт деацетилювання). - Часткова кристалічність.	- Біосумісність. - Антимікробні властивості. - Здатність утворювати хімічні зв'язки з неорганічними компонентами. - Легкість модифікації.
Альгірати	- Поліаніони, що складаються з мономерів альгінової кислоти. - Утворюють міцні гелі у присутності катіонів (Ca^{2+}). - Висока гідрофільність.	- Формування гідрогелів. - Легкість модифікації. - Біосумісність. - Стійкість до розтягування в гелевому стані.
Колаген	- Трикутна гвинтова структура. - Складається з амінокислот (переважно гліцину, проліну, гідроксипроліну). - Висока здатність до зшивання та стабілізації.	- Висока міцність на стиск- Біосумісність. - Здатність утримувати воду. - Відмінна адгезія до інших матеріалів.
Лігнін	- Комплексна ароматична структура. - Сітчаста (некристалічна) будова. - Високий вміст π-електронних систем.	- Радіаційна стійкість. - Антиоксидантні властивості. - Можливість утворювати композити з іншими полімерами. - Екологічність.
Кератин	- Висока кількість дисульфідних зв'язків. - Висока еластичність та здатність до утворення волокон.	- Термостійкість. - Еластичність. - Здатність до хімічної модифікації.
Пектин	- Полімер галактуранової кислоти. - Часто містить метильовані та ацетильовані групи. - Висока гідрофільність.	- Формування гідрогелів- Біосумісність. - Стійкість до хімічної модифікації.

Етап 2. Модифікація полімеру

Хімічна або фізична модифікація природного полімеру в деяких випадках є необхідною для покращення експлуатаційних властивостей матеріалу, підвищення його сумісності з наповнювачем тощо. В якості прикладу можна привести хімічну модифікацію полімеру для підвищення стійкості до радіації (наприклад, хімічна зшивка). Крім того, часто важливим фактором є введення функціональних груп, що можуть взаємодіяти з неорганічними компонентами.

Модифікація природних полімерів є важливим етапом у створенні радіаційно стійких матеріалів. Хімічні зміни структури полімеру, такі як зшивка або введення функціональних груп, дозволяють значно покращити його механічну міцність, стабільність до радіаційного впливу та здатність до взаємодії з неорганічними компонентами.

Варто зазначити, що модифікація полімеру не є обов'язковою, і застосовується тільки у випадку виявлення та потреби внесенні змін до незадовільних властивостей виявлених на Етапі 1.

Етап 3. Інтеграція наповнюючих (неорганічних) компонентів

Однією з ключових проблем під час створення радіаційно стійких матеріалів є нейтралізація вторинних електронів і вільних радикалів, які утворюються під впливом радіації. Ці частинки здатні спричиняти розрив хімічних зв'язків у полімерних матрицях, що призводить до деградації та втрати функціональних властивостей матеріалів. Для мінімізації цього ефекту до складу матеріалу впроваджують наповнюючі компоненти, здатні ефективно поглинати або розсіювати енергію вторинних електронів.

Також важливо забезпечити матеріал радіаційно стійкою оболонкою товщиною декілька мікрон, виготовленою з матеріалу що запобігає поширенню вторинних електронів та нейтралізує вільні радикали. Утворення вторинних електронів може спричиняти розрив хімічних зв'язків в мікрометровому масштабі, і саме тому важливо «нейтралізувати» так вторинні електрони.

На даному етапі важливо забезпечити вибір радіаційно стійких неорганічних наповнювачів (наприклад, наночастинки оксидів металів) та розробку методик введення наповнювачів у полімерну матрицю.

Фактори, які впливають на вибір методики введення наповнювача:

1. Тип полімерної матриці: термопласти або термореактивні полімери.
2. Розмір і тип наповнювача: макро-, мікро- чи наночастинки.
3. Кінцеві властивості матеріалу: механічна міцність, термостабільність, стійкість до радіації.
4. Технологічна доступність: наявність обладнання та економічна доцільність.

Розробка методик введення наповнювачів у полімерну матрицю є важливим етапом створення гібридних матеріалів. Оптимальний розподіл і взаємодія між полімерною матрицею та наповнювачем впливають на фізико-механічні властивості, стабільність і функціональність кінцевого матеріалу. Нижче наведено схему розподілу неорганічного наповнювача в полімерній матриці, яка демонструє взаємодію між біополімером та наповнювачем (рис. 3). Для покращення механічних і функціональних властивостей матеріалу використовуються модифіковані поверхні полімеру, що забезпечують ефективну міжфазову взаємодію з наповнювачем. Такий підхід сприяє рівномірному розподілу наповнювача та оптимізації властивостей композитного матеріалу.

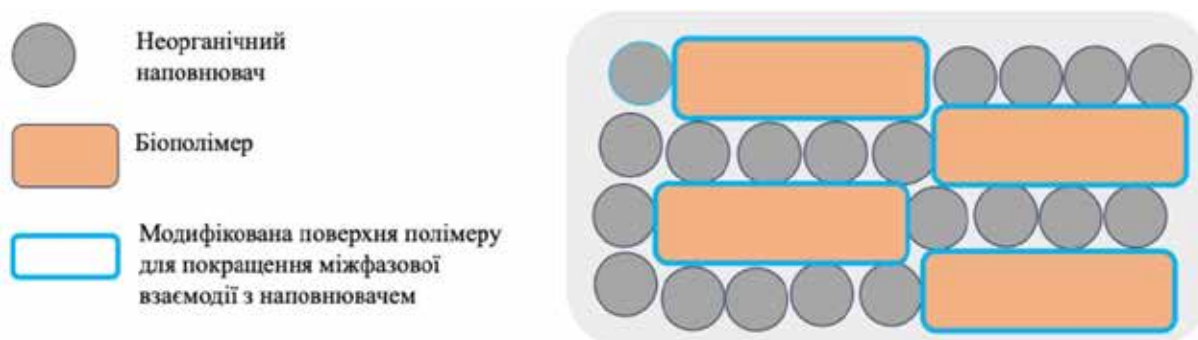


Рис. 3. Схема розподілу неорганічного наповнювача в полімерній матриці

Етап 4. Синтез гібридного матеріалу

Оптимізація умов синтезу:

Температура: Контроль температури має бути точним, оскільки високі температури можуть призводити до деградації природних полімерів. Рекомендовано використовувати температури нижче

температури розкладу полімеру, але достатньо високі для активації реакцій (наприклад, 60–80°C для хітозану або лігніну).

Час: Потрібно знайти баланс між достатнім часом для завершення реакції та уникненням тривалого впливу, який може спричинити небажані побічні реакції (наприклад, гель-утворення або структурні дефекти).

Концентрація реагентів: Варто враховувати, що надлишок реагентів може викликати побічні реакції, зокрема неконтрольоване зшивання. Тому концентрацію потрібно ретельно оптимізувати через попередні експерименти.

Рівномірний розподіл неорганічних компонентів:

Використання ультразвукової обробки або високошвидкісного змішування може допомогти забезпечити рівномірність.

Етап 5. Характеризація та тестування

Фізико-хімічні аналізи отриманого матеріалу: **FTIR (ІЧ-спектроскопія)** для підтвердження утворення хімічних зв'язків. **SEM/TEM (електронна мікроскопія)** для візуалізації розподілу компонентів. **TGA/DSC (термічні методи)** для оцінки термостійкості. **XRD (рентгенівська дифракція)** для визначення кристалічної структури. Важливо також оцінити **адгезію** між матрицею та неорганічними компонентами.

Тестування радіаційної стійкості та механічних властивостей:

Для радіаційної стійкості рекомендовано тестувати матеріал при різних дозах та в різних умовах (з повітрям/без повітря) і оцінювати такі параметри:

- Структурні зміни (FTIR, TGA).
- Утворення вторинних радикалів (ЕПР-спектроскопія).
- Зміна механічних властивостей (міцність на розтяг, ударна в'язкість).

Механічні випробування:

- Важливо оцінити зміну модулів Юнга, межі текучості та міцності після радіаційного опромінення.
- Використання динамічних методів (наприклад, DMA) дозволяє отримати додаткову інформацію про поведінку матеріалу в умовах експлуатації.

Алгоритм створення радіаційно стійких гібридних матеріалів на основі природних полімерів.

Розробка радіаційно стійких гібридних матеріалів на основі природних полімерів є актуальним напрямом сучасного матеріалознавства, особливо для застосувань у космічній техніці, ядерній енергетиці та медицині.

Загальний опис ключових етапів:

Етап 1–2: Глибокий аналіз та вибір матеріалів є критичними для забезпечення бажаних властивостей кінцевого продукту.

Етап 3–5: Модифікація та синтез вимагають точного контролю хімічних процесів для досягнення стабільної гібридної структури.

Етап 6–7: Формування та характеризація матеріалу дозволяють оцінити його придатність для практичного застосування.

Етап 8–9: Тестування радіаційної стійкості та подальша оптимізація забезпечують відповідність матеріалу вимогам експлуатації в екстремальних умовах.

Етап 10: Масштабування процесу та впровадження на ринок є завершальними кроками для реалізації проекту.

Запропонована концептуальна схема та алгоритм створення таких матеріалів включають наступні етапи:

Постановка мети та завдань:

- Визначення специфічних вимог до матеріалу, таких як радіаційна стійкість, механічна міцність, термостабільність та біосумісність.
- Ідентифікація потенційних сфер застосування, включаючи космічну техніку, ядерну енергетику та медицину.

Вибір базового природного полімеру:

- Аналіз доступних полімерів, таких як целюлоза, хітозан, колаген, альгінати та кератин.
- Оцінка їхніх механічних, термічних, хімічних та радіаційних властивостей.
- Вибір оптимального полімеру на основі встановлених вимог до кінцевого матеріалу.

Модифікація природного полімеру:

- Хімічна модифікація, включаючи крос-зв'язування, ацилювання та етерифікацію, для підвищення стійкості до радіації.
- Введення функціональних груп для покращення взаємодії з неорганічними компонентами.

Вибір неорганічних компонентів:

– Ідентифікація радіаційно стійких матеріалів, таких як оксиди металів (наприклад, TiO_2 , Al_2O_3), наночастинки графену та керамічні наповнювачі.

– Оцінка їхньої сумісності з модифікованим полімером.

– Визначення оптимальної концентрації та розміру частинок.

Розробка методики синтезу гібридного матеріалу:

– Вибір методу синтезу, такого як сол-гел процес, ін-ситу полімеризація або механічне змішування.

– Оптимізація параметрів процесу, включаючи температуру, час реакції та pH середовища.

– Забезпечення рівномірного розподілу неорганічних компонентів у полімерній матриці.

Формування матеріалу:

– Застосування технологій формування, таких як лиття, екструзія або 3D-друк.

– Контроль морфології, включаючи пористість, товщину та однорідність структури.

– Проведення пост-обробки, такої як термічна обробка, відпал або сушка.

Характеризація отриманого матеріалу:

– Структурний аналіз за допомогою методів, таких як рентгенівська дифракція (XRD), інфрачервона спектроскопія (FTIR) та скануюча електронна мікроскопія (SEM).

– Вимірювання фізико-механічних властивостей, включаючи міцність на розрив, еластичність та твердість.

– Оцінка термостабільності за допомогою термогравіметричного аналізу (TGA) та диференціального термічного аналізу (DTA).

Тестування радіаційної стійкості:

– Опромінення зразків при різних дозах радіації.

– Аналіз змін фізико-хімічних та механічних властивостей після опромінення.

– Визначення механізмів деградації та стабільності матеріалу.

Оптимізація складу та технології:

– Аналіз отриманих даних для визначення кореляції між складом, структурою та властивостями.

– Коригування рецептури та умов синтезу для досягнення оптимальних характеристик.

– Повторне тестування для підтвердження поліпшень.

Масштабування та впровадження:

– Розробка технологічного процесу для промислового виробництва.

– Оцінка економічної ефективності та екологічних аспектів.

– Підготовка до комерціалізації, включаючи патентування та сертифікацію

На (рис. 4) наведено візуалізацію покрокового алгоритму створення радіаційно стійких гібридних матеріалів на основі природних полімерів. Представлений рисунок ілюструє ключові етапи процесу, починаючи від вибору базового полімеру та його модифікації до тестування матеріалів на радіаційну стійкість. Такий підхід дозволяє систематизувати процес та забезпечити досягнення оптимальних характеристик матеріалів для застосування в умовах підвищеного радіаційного фону [1].

Вищевикладений алгоритм та його підходи потребують *міждисциплінарного підходу*, який об'єднує експертизу хіміків, фізиків, матеріалознавців та інженерів для досягнення оптимальних результатів у створенні гібридних матеріалів. Така співпраця дозволяє забезпечити комплексний аналіз, модифікацію та впровадження нових матеріалів із підвищеною радіаційною стійкістю.

Не менш важливим є врахування *екологічних аспектів* та наслідків використання безпечних методів і матеріалів. Це сприятиме мінімізації негативного впливу на довкілля як під час виробництва, так і під час експлуатації створених матеріалів.

Висновки. Науково обґрунтовано доцільність створення екологічно безпечних та радіаційно стійких гібридних матеріалів на основі природних полімерів як альтернативи традиційним екранувальним засобам. Сучасні радіаційні технології, які активно застосовуються в ядерній енергетиці, медицині, агропромисловому комплексі та харчовій промисловості, вимагають високоефективного й безпечно-го захисту від іонізуючого випромінювання. При цьому традиційні матеріали, що використовуються для екранування (свинець, бетон, парафін, поліетилен тощо), характеризуються високою токсичністю, складністю утилізації та неспроможністю відповідати екологічним стандартам, прийнятим у контексті стратегії сталого розвитку та зменшення антропогенного впливу. У роботі доведено, що використання природних полімерів, зокрема лігніну, як основи для створення захисних матеріалів є перспективним напрямом, що поєднує ефективність, безпечність і екологічність.

Запропоновано концептуальну схему та алгоритм створення гібридних композитів, що забезпечують багаторівневий захист від іонізуючого випромінювання. Розроблена модель включає поетапний підхід: від вибору природної полімерної основи з високими механічними та хімічними характеристиками

до її модифікації для покращення адгезії та стійкості до опромінення. Наступним критичним етапом є введення неорганічних наповнювачів – наночастинок оксидів металів, які здатні нейтралізувати вторинні електрони та вільні радикали, що утворюються внаслідок впливу іонізуючого випромінювання. Визначено технологічні умови синтезу, що дозволяють досягти стабільності структури композиту, його однорідності та цільових функціональних властивостей. Сформульовано підходи до використання фізико-хімічних методів аналізу (FTIR, SEM, TGA, XRD) та механічних випробувань, які підтверджують ефективність створених матеріалів.

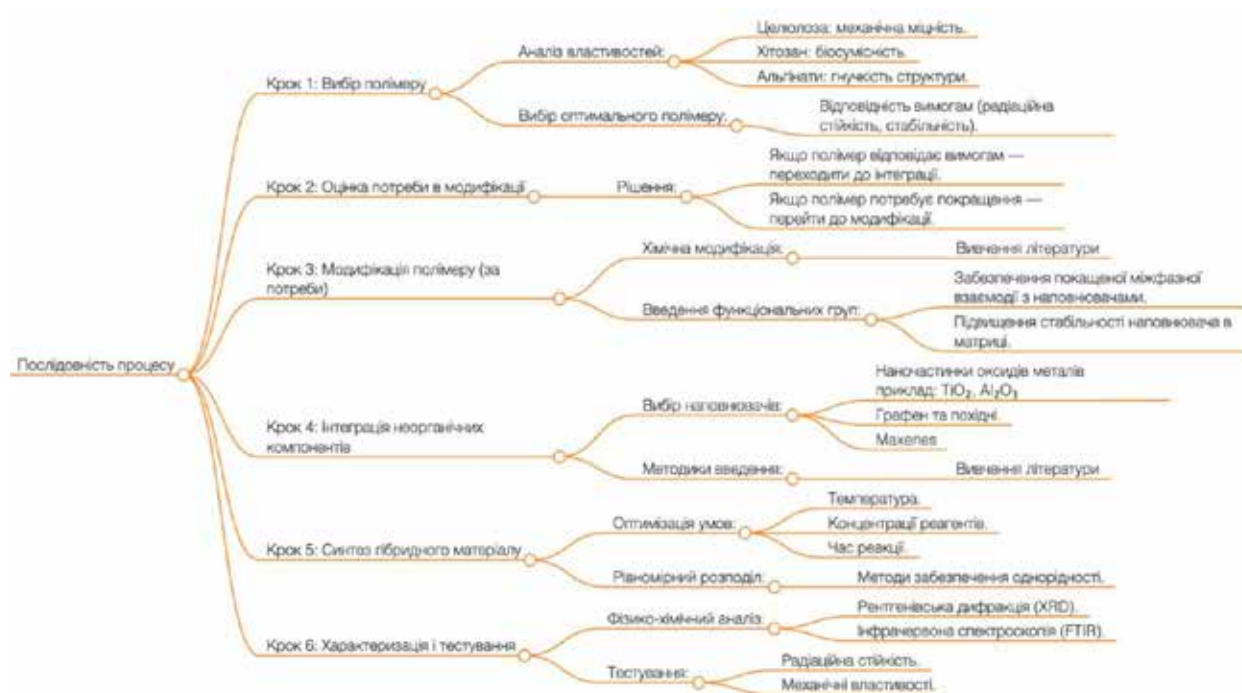


Рис. 4. Візуалізація покрокового алгоритму створення радіаційно стійких гібридних матеріалів на основі природних полімерів

Результати дослідження відкривають нові можливості для формування міждисциплінарної платформи розроблення матеріалів нового покоління для умов радіаційного навантаження. Запропонована концепція є придатною до адаптації в низці галузей, де необхідний надійний та тривалий захист від іонізуючого випромінювання – зокрема в космічній техніці, медицині (радіотерапія, транспортування радіоізотопів), зберіганні й перевезенні радіоактивних відходів, а також у проектуванні захисних бар'єрів у ядерних об'єктах. Висвітлені у роботі наукові підходи базуються на принципах екологічної безпеки, що особливо важливо в умовах глобального тренду на використання біосумісних і відновлюваних матеріалів. Дослідження закладає основу для подальших розробок у галузі інженерії функціональних біоматеріалів, дозволяючи сформувати нову стратегію щодо раціонального використання природних полімерів у високотехнологічних захисних системах.

Список використаних джерел:

1. Нові радіаційно стійкі гібридні матеріали на основі природних полімерів. (Етап: Розроблення концепції та алгоритму створення гібридних радіаційно стійких матеріалів на основі природних полімерів). Центр інформаційно-аналітичного та технічного забезпечення моніторингу об'єктів атомної енергетики НАН України. 2024. № 0224U033588. 86 с.
2. Abu Bakar N. F., Amira Othman S., Amirah Nor Azman N. F., Saqinah Jasrin N. Effect of ionizing radiation towards human health: A review. IOP Conference Series: *Earth and Environmental Science*. 2019. Vol. 268(1). P. 012005. URL: <https://doi.org/10.1088/1755-1315/268/1/012005>
3. Chang S., Guo X., Zhang X. Radiation shielding polymer composites: Ray-interaction mechanism, structural design, manufacture and biomedical applications. *Materials & Design*. 2023. Vol. 233. 112253. URL: <https://doi.org/10.1016/j.matdes.2023.112253>
4. Chmielewski A. G. Radiation technologies: The future is today. *Radiation Physics and Chemistry*. 2023. Vol. 213. 111233. URL: <https://doi.org/10.1016/j.radphyschem.2023.111233>

5. Luo Y.-S., Chen Z., Hsieh N.-H., Lin T.-E. Chemical and biological assessments of environmental mixtures: A review of current trends, advances, and future perspectives. *Journal of Hazardous Materials*. 2022. Vol. 432. 128658. URL: <https://doi.org/10.1016/j.jhazmat.2022.128658>
6. Moreno A., Horizons M. S.-M. Lignin-Based Smart Materials: A Roadmap to Processing and Synthesis for Current and Future Applications. *Materials Horizons*. 2020. Vol. 7(9). P. 2237–2257. URL: <https://doi.org/10.1039/D0MH00798F>
7. Pylypchuk Ie., Kovach V., Iatsyshyn A., Kutsenko V., Taraduda D. Development Boron and Gadolinium-Containing Composite Materials Based on Natural Polymers for Protection Against Neutron Radiation. In: Zaporozhets A. (ed.) *Systems, Decision and Control in Energy V*. Cham: Springer, 2023. P. 527–540. URL: https://doi.org/10.1007/978-3-031-35088-7_28
8. Sarosi O., Sulaeva I., Fitz E., Summerskii I., Bacher M., Potthast A. Lignin Resists High-Intensity Electron Beam Irradiation. *Biomacromolecules*. 2021. Vol. 22. P. 4365–4372. URL: <https://doi.org/10.1021/acs.biomac.1c00926>
9. Vieira Y., et al. A critical review of the current environmental risks posed by the antidiabetic Metformin and the status, advances, and trends in adsorption technologies for its remediation. *Journal of Water Process Engineering*. 2023. Vol. 54. 103943. URL: <https://doi.org/10.1016/j.jwpe.2023.103943>
10. Xu C., Xia D., Zhang X., Yao Q., Wang Y., Zhang C. In situ analysis of metallodrugs at the single-cell level based on synchrotron radiation technology. *TrAC Trends in Analytical Chemistry*. 2024. Vol. 171. 117515. URL: <https://doi.org/10.1016/j.trac.2023.117515>

Дата надходження статті: 30.06.2025

Дата прийняття статті: 08.07.2025

Опубліковано: 23.09.2025

УДК 004.8:681.3.06:519.816

DOI <https://doi.org/10.32689/maup.it.2025.2.24>

Олексій РИХАЛЬСЬКИЙ

викладач кафедри комп'ютерних наук та програмної інженерії,

ПВНЗ «Європейський університет»

ORCID: 0009-0000-6892-9420

АРХІТЕКТУРА ІНТЕЛЕКТУАЛЬНОЇ СИСТЕМИ ПІДТРИМКИ РІШЕНЬ НА ОСНОВІ СЕМАНТИЧНОГО УЗГОДЖЕННЯ МОДЕЛЕЙ

Анотація. У статті розглядається перспективний напрям розвитку систем підтримки прийняття рішень (СППР), що відповідає висхідній потребі в інтелектуальних інструментах, здатних обробляти складні, гетерогенні та динамічно мінливі дані. В умовах, коли традиційним підходам до СППР часто не вистачає гнучкості для адаптації до специфічних нюансів предметної області, а також вони не враховують належним чином семантику даних і наміри користувачів, це дослідження усуває критичну прогалину, надаючи нову архітектуру, що базується на семантичній відповідності моделей.

Метою статті є опис архітектури інтелектуальної СППР, що використовує семантичні технології для забезпечення більш точної, контекстно орієнтованої та орієнтованої на користувача підтримки прийняття рішень.

Методи: аналізу наукової літератури, узагальнення, систематизації.

Наукова новизна статті полягає у формалізації архітектури інтелектуальної системи підтримки рішень, що ґрунтується на принципах семантичного узгодження моделей різної природи для підвищення ефективності процесу прийняття рішень.

Результати. Основна ідея полягає в інтеграції онтологій та стандартів семантичного вебу для забезпечення інтероперабельності між гетерогенними джерелами даних і автоматизації процесу зіставлення запитів користувачів з відповідними моделями знань. Запропонована архітектура системи складається з декількох основних компонентів, зокрема репозиторію семантичних моделей, механізму виведення, модуля взаємодії з користувачем та інтерфейсів для інтеграції зовнішніх даних. Значний акцент зроблено на динамічній конфігурації моделей прийняття рішень на основі критеріїв семантичної подібності та релевантності.

Висновки. Показано, що механізми семантичного узгодження значно підвищують точність рекомендацій і скорочують час, необхідний для прийняття рішень у складних середовищах. Описано практичні сценарії застосування, що підтверджують адаптивність та масштабованість системи. У висновках наголошено на потенціалі семантичних технологій у створенні інтелектуальних СППР та підкреслено важливість подальших досліджень у напрямку вдосконалення алгоритмів міркувань, механізмів адаптації користувачів та розробки онтологій для конкретної предметної області. Запропонована практика сприяє еволюції СППР у бік більш гнучких, інтелектуальних та контекстно орієнтованих систем.

Ключові слова: проектування, аналіз, параметр, складні системи, база знань, гетерогенні моделі, семантична сумісність.

Oleksiy RYKHALSKYY. ARCHITECTURE OF AN INTELLIGENT DECISION SUPPORT SYSTEM BASED ON SEMANTIC MODEL MATCHING

Abstract. The article explores a promising direction in the development of decision support systems (DSS), addressing the growing need for intelligent tools capable of processing complex, heterogeneous, and dynamically changing data. In situations where traditional DSS approaches often lack the flexibility to adapt to specific domain nuances and fail to adequately account for data semantics and user intent, this study fills a critical gap by introducing a new architecture based on semantic model alignment.

The aim of the article is to describe the architecture of an intelligent DSS that employs semantic technologies to provide more accurate, context-aware, and user-centered decision support.

Methods. Analysis of scientific literature, generalization, and systematization.

Scientific novelty. The scientific novelty of the article lies in the formalization of the architecture of an intelligent decision support system based on the principles of semantic alignment of heterogeneous models to enhance the efficiency of the decision-making process.

Results. The main idea is the integration of ontologies and semantic web standards to ensure interoperability between heterogeneous data sources and to automate the process of matching user queries with relevant knowledge models. The proposed system architecture consists of several core components, including a repository of semantic models, an inference engine, a user interaction module, and interfaces for external data integration. Special attention is paid to the dynamic configuration of decision-making models based on criteria of semantic similarity and relevance.

Conclusions. The study demonstrates that semantic alignment mechanisms significantly improve the accuracy of recommendations and reduce the time required for decision-making in complex environments. Practical application scenarios are described, confirming the system's adaptability and scalability. The conclusions emphasize the potential of semantic

technologies in building intelligent DSS and highlight the importance of further research aimed at improving reasoning algorithms, user adaptation mechanisms, and the development of domain-specific ontologies. The proposed approach supports the evolution of DSS toward more flexible, intelligent, and context-aware systems.

Key words: design, analysis, parameter, complex systems, knowledge base, heterogeneous models, semantic compatibility.

Постановка наукової проблеми. У сучасному світі інформаційних технологій та прийняття рішень на основі даних різні організації в різноманітних секторах зіштовхуються зі все більшими обсягами гетерогенних, неструктурованих і динамічних даних. Традиційні системи підтримки прийняття рішень, хоч і ефективні в структурованих середовищах, часто не справляються в умовах, коли гнучкість, контекстна обізнаність і семантичне розуміння є критично важливими. Ці системи зазвичай працюють на основі заздалегідь визначених правил і статичних моделей даних, що обмежує їхню здатність адаптуватися до нюансів конкретних сфер або намірів користувача. Зі зростанням складності сценаріїв прийняття рішень – особливо в таких галузях, як охорона здоров'я, інженерія та логістика – виникає потреба в системах, що не лише об'єднують і обробляють інформацію, але й інтерпретують її значення в контекстно чутливий спосіб. Це вимагає переходу від синтаксичної обробки даних до практик, що передбачають семантичне мислення, що дозволяє системам осмислювати інформацію так само, як це робить людина [4].

Актуальність розробки інтелектуальної системи підтримки прийняття рішень (СППР) на основі семантичної моделі обумовлена трансформацією способу прийняття рішень у середовищах, багатих на дані. Семантичні технології, як-от онтології та графи знань, надають потужні засоби для структурування знань про предметну область, забезпечення взаємодії між розрізненими джерелами даних та узгодження результатів роботи системи з цілями користувача. Інтеграція цих технологій в архітектуру СППР дозволяє забезпечити динамічну конфігурацію моделі, підвищити точність рекомендацій та кращу адаптивність до мінливих умов. У цьому контексті проблема полягає в розробці архітектури системи, що може ефективно використовувати механізми семантичного узгодження для інтерпретації запитів користувачів, співвіднесення їх з відповідними моделями знань і прийняття обґрунтованих, контекстно орієнтованих рішень. Розв'язання цієї проблеми є не лише технічним завданням, але й необхідним кроком на шляху до створення дійсно інтелектуальної, гнучкої та орієнтованої на користувача системи підтримки прийняття рішень у складних проблемних областях.

Аналіз останніх досліджень і публікацій. Сучасні наукові дослідження у сфері побудови інтелектуальних систем підтримки прийняття рішень (ІСППР), зокрема на основі семантичного узгодження моделей, засвідчують висхідний інтерес науковців до проблем інтеграції знань, онтологічного моделювання, розпізнавання структур знань та концептуального представлення інформації. У цьому контексті аналіз сучасних публікацій дозволяє окреслити основні тенденції, інноваційні підходи та перспективи розвитку архітектури ІСППР.

У дослідженні К. Ткаченка та А. Байдака запропоновано онтологічний підхід до моделювання ІСППР, орієнтованої на аналіз ризиків в інноваційно-інвестиційній сфері [7]. Автори обґрунтовують необхідність використання онтологічного представлення знань як засобу семантичного узгодження інформаційних моделей, що дозволяє забезпечити формальне подання предметної області та ефективну інтеграцію джерел знань. Основна увага приділена структурі онтології, де окреслено зв'язки між об'єктами, подіями, характеристиками та критеріями ризику. Цей підхід створює підґрунтя для побудови архітектури ІСППР, що здатна адаптуватися до змін у середовищі, базуючись на динамічному поповненні онтологічних знань.

М. В. Євланов, Б. І. Мороз та В. Я. Лучицький у своїй роботі досліджують інформаційну технологію виявлення термінів та артефактів у проєктній документації інформаційних систем [2]. Розроблена практика базується на використанні мовних моделей та алгоритмів обробки природної мови для автоматичного виявлення семантично значущих об'єктів. Це дозволяє формувати структуровану модель предметної області на ранніх етапах життєвого циклу ІСППР. Важливим аспектом цієї роботи є підвищення точності ідентифікації концептів та зменшення семантичної неоднозначності, що є критично важливим у контексті узгодження моделей.

У статті В. Тріща та співавторів акцент зроблено на засобах підтримки сховищ знань комп'ютерних систем у нафтогазовій галузі [8]. Автори пропонують архітектурне рішення, що містить інструменти для агрегації, збереження і пошуку знань на основі онтологічних структур. Представлена практика демонструє ефективність семантичного узгодження моделей у галузевих інформаційних системах, де основну роль відіграє точність відображення знань та швидкість їх обробки. У дослідженні також порушено питання гетерогенності джерел інформації та виклики, пов'язані з узгодженням термінології.

Дослідження С. Ф. Чалого та І. О. Лещинської зосереджене на процесі екстерналізації знань у ментальній моделі користувача системи штучного інтелекту [9]. Автори аналізують механізми формування

семантичної структури знань з урахуванням когнітивних особливостей користувача. Такий підхід є особливо цінним для архітектури ІСППР, орієнтованої на адаптивну взаємодію з користувачем, де системі необхідно не лише опрацьовувати дані, а й «розуміти» наміри та логіку рішень людини. Дослідження підкреслює важливість створення персоналізованих семантичних моделей у гнучких архітектурах ІСППР.

М. В. Пономаренко розглядає проблематику державної політики у сфері інтелектуальної власності у контексті міжвідомчої комунікації [6]. Хоча стаття має скоріше правовий і регуляторний фокус, вона акцентує на важливості забезпечення узгодженості та прозорості обміну даними між структурами. У цьому сенсі семантичне узгодження моделей у рамках ІСППР виступає визначальним елементом забезпечення інтероперабельності та довіри до автоматизованих систем прийняття рішень, особливо у міжгалузевому середовищі.

У роботі М. Г. Петренка та співавторів розглядається побудова комп'ютерних систем подання та оброблення предметно орієнтованих знань [5]. Автори пропонують методологію побудови знаннєвих моделей, що охоплює класифікацію знань, їх семантичне структурування та використання логічних моделей обробки. Запропоновані засоби забезпечують високий рівень формалізації знань, що сприяє їх ефективному застосуванню у моделі ІСППР. Особливу увагу приділено забезпеченню актуальності знань, що є важливим фактором динамічної адаптації систем.

У статті О. Горди та співавторів досліджується підхід до інформаційного моделювання на основі метафор роїв [1]. Це нестандартне, біоінспіроване рішення передбачає моделювання взаємодії інформаційних об'єктів у системі за аналогією з поведінкою роїв у природі. Такий підхід дозволяє формувати самоврядні системи підтримки прийняття рішень, які володіють високою адаптивністю до змін середовища. Семантичне узгодження у цьому випадку реалізується через механізми колективного пізнання, що дає змогу створювати гнучкі та інтелектуальні інформаційні архітектури.

Виділення не вирішених раніше частин загальної проблеми. У межах цього дослідження акцентовано на аспекті динамічного семантичного узгодження моделей у гетерогенному інформаційному середовищі, що раніше належно не опрацьовувався. Також розглянуто питання інтеграції онтологічних структур з механізмами адаптивного управління знаннями в реальному часі. Окрема увага приділяється архітектурним рішенням, що дозволяють забезпечити контекстну узгодженість між моделями різних рівнів абстракції в інтелектуальних системах підтримки прийняття рішень.

Формулювання мети дослідження. Метою дослідження є опис архітектури інтелектуальної СППР, що використовує семантичні технології для забезпечення точнішої, контекстно орієнтованої та орієнтованої на користувача підтримки прийняття рішень.

Відповідно до мети було поставлено і вирішено такі завдання: проаналізувати сучасні практики до побудови систем підтримки прийняття рішень із використанням семантичних технологій, сформулювати вимоги до архітектури інтелектуальної СППР на основі семантичного зіставлення моделей, розробити структурну архітектуру системи з визначенням основних функціональних компонентів.

Виклад основного матеріалу дослідження. Дослідження інтелектуальних систем підтримки прийняття рішень (СППР) у сфері проєктування складних технічних об'єктів зазнали значної трансформації, особливо в машинобудуванні, де складність продуктів і процесів вимагає вдосконалених інструментів для покращення процесу прийняття рішень. Інтелектуальні системи підтримки прийняття рішень стали важливим компонентом у цій галузі, сприяючи не лише автоматизації рутинних завдань, але й допомагаючи людям-розробникам вирішувати недостатньо структуровані, нестандартні проблеми. В основі СППР у машинобудуванні закладена їхня здатність моделювати, аналізувати та допомагати у розв'язанні проблем проєктування шляхом інтеграції експертних знань, методів штучного інтелекту та інтерактивних користувацьких інтерфейсів. Протягом багатьох років пропонувалися різні практики до побудови таких систем, починаючи від систем на основі правил та експертних систем і закінчуючи новітніми розробками, що використовують міркування на основі конкретних ситуацій, міркування на основі моделей та гібридні методи. Ці методи дозволяють системі імітувати рішення на рівні експертів шляхом застосування специфічних для конкретної галузі знань елементів і механізмів навчання, які адаптуються до вимог і обмежень, що постійно зазнають трансформацій.

У контексті машинного проєктування традиційні системи підтримки прийняття рішень спочатку зосереджувалися на формалізації експертних знань у вигляді виробничих правил або дерев прийняття рішень. Ці методи хоча і були ефективними в структурованих середовищах, проте не могли впоратися з реальними динамічними та часто неоднозначними завданнями проєктування. Як наслідок, галузь поступово перейшла до гнучкіших та адаптивних фреймворків, що використовують схеми представлення знань, як-от семантичні мережі, онтології та об'єктно орієнтовані моделі. Водночас інтеграція інструментів штучного інтелекту (алгоритми машинного навчання, нейронні мережі та нечітка логіка)

розширила можливості СППР, дозволивши їм обробляти неповну або невизначену інформацію та надавати рекомендації, що ґрунтуються як на даних, так і на інтуїції експертів.

Класифікація архітектур IDSS відображає різноманітність підходів і технологій, що використовуються при їх розробці. Загалом ці архітектури можна поділити на: централізовані, розподілені та гібридні системи. Централізовані архітектури, як правило, мають єдину базу знань і механізм виведення, що спрощує управління системою, але можуть обмежувати масштабованість і швидкість реагування у великомасштабних застосунках. Розподілена архітектура, навпаки, використовує кілька співзалежних агентів або модулів, які працюють напівнезалежно і спілкуються за допомогою визначених протоколів. Така структура посилює модульність і паралелізм, але створює проблеми з синхронізацією та узгодженістю знань. Гібридні архітектури поєднують елементи обох типів, намагаючись збалансувати переваги централізованого контролю з гнучкістю розподіленої обробки. Водночас деякі реалізації IDSS використовують мультиагентні архітектури, де автономні програмні агенти співпрацюють для підтримки складних сценаріїв прийняття рішень, особливо в середовищах спільного проектування.

Попри значний прогрес, традиційні моделі СППР мають низку недоліків, що перешкоджають їх широкому застосуванню в проектуванні складних технічних об'єктів. Одним з основних обмежень є їхня залежність від чітко визначених правил і статичних баз знань, що обмежує їхню адаптивність до нових проблем проектування або мінливого середовища. Таким системам часто не вистачає здатності вчитися на новому досвіді або оновлювати свої стратегії міркувань у відповідь на зворотний зв'язок, що робить їх менш ефективними в динамічних або інноваційних середовищах дизайну. Водночас презентація знань дизайнерів становить значну проблему. Значна частина знань, застосованих в інженерному дизайні, є неявними, емпіричними та залежними від контексту, що ускладнює їх фіксацію та формалізацію в рамках жорстких структур традиційних систем. Ця проблема посилюється міждисциплінарним характером складних технічних об'єктів, де знання мають бути інтегровані з різних галузей, як-от матеріалознавство, механіка, системи управління та ергономіка [3].

Проблеми представлення знань розробників поширюються також на когнітивні аспекти прийняття дизайнерських рішень. Дизайнери спираються на інтуїцію, аналогічне мислення та евристичні судження, що нелегко змодельовати за допомогою формальної логіки або традиційних методів представлення знань. Як наслідок СППР часто не підтримують творчі та ітеративні аспекти процесу проектування. Заходи, спрямовані на подолання цих обмежень, призвели до вивчення когнітивних архітектур, систем «людина-в-циклі» та моделей на основі навчання, що мають на меті подолати прогалину між формалізованими знаннями та людськими когнітивними процесами. Досягнення в моделюванні користувачів та адаптивних інтерфейсів спрямовані на персоналізацію поведінки системи відповідно до індивідуальних уподобань та стилю прийняття рішень дизайнера, що ще більше підвищує зручність використання та актуальність СППР у складних завданнях проектування.

Координація семантичних моделей є базовим механізмом для досягнення ефективної інтеграції знань у складних процесах прийняття рішень, особливо у сфері міждисциплінарної інженерії. В міру зростання масштабів і складності інженерних систем, здатність координувати й узгоджувати розрізнені моделі, що походять з різних областей, експертів і програмних інструментів, стає критичним фактором для забезпечення послідовних і узгоджених результатів прийняття рішень. За своєю суттю семантична координація моделей передбачає узгодження значень і концептуальних структур між різними моделями, щоб уможливити їхню спільну інтерпретацію, взаємну трансформацію і синергетичне використання. Ця координація виходить за рамки простого обміну даними або синтаксичної сумісності, звертаючись до глибшого рівня семантики – того, як поняття визначаються, пов'язуються і розуміються в конкретних дисциплінарних контекстах. Семантичний рівень гарантує, що коли моделі взаємодіють, їхній зміст інтерпретується послідовно в різних системах, тим самим зменшуючи неоднозначність і підвищуючи надійність інтегрованих рішень.

Онтології та метамоделі виступають фундаментальними інструментами в цьому процесі гармонізації. Онтології надають формальні, явні специфікації концептуалізацій у межах предметної області, визначаючи сутності, атрибути, зв'язки й правила в машинно інтерпретованому форматі. Фіксуючи знання про предметну область у структурований і придатний для багаторазового використання спосіб, онтології сприяють створенню загальних словників і спільному розумінню між гетерогенними моделями. Вони діють як семантичні мости, що дозволяють незалежно розробленим моделям взаємодіяти в рамках єдиної системи прийняття рішень. Метамоделі, своєю чергою, визначають абстрактний синтаксис і структурні правила побудови зразків певною мовою моделювання. Вони визначають, як конкретні моделі конкретизуються і як вони співвідносяться на метарівні, що має важливе значення для управління сумісністю моделей на різних платформах моделювання. Взаємодія між онтологіями та

метамоделями дозволяє кодувати як семантичний зміст, так і структурні правила, забезпечуючи таким чином двохшарову основу для семантичної координації моделей.

Формалізація моделей та їхня взаємодія в модельно-параметричному просторі вимагає точних методів представлення, що можуть охоплювати не лише значення параметрів і структурні компоненти, але й основну семантику проектного задуму, обмежень і функціональних зв'язків. Для досягнення такого рівня формалізації часто використовують такі практики, як-от модельно орієнтоване проектування (MDE), формальний концептуальний аналіз і логіка опису. MDE наголошує на використанні високорівневих абстрактних моделей як первинних артефактів у процесі проектування, що дозволяє здійснювати автоматизовані трансформації та синхронізацію між моделями за допомогою визначених відображень. У контексті семантичної координації моделей MDE полегшує відстежуваність і узгодженість трансформацій моделей, забезпечуючи збереження семантичної цілісності при переході між рівнями абстракції або дисциплінарними межами. Модельно-параметричний простір, що охоплює всю сукупність змінних, параметрів та їх взаємозалежностей, структурований за допомогою метамodelей і збагачений семантичними анотаціями, отриманими з онтологій. Ця структура підтримує різні сценарії взаємодії: об'єднання моделей, уточнення, декомпозиція та розв'язання конфліктів.

Міждисциплінарні інженерні середовища створюють особливі проблеми для семантичної сумісності, враховуючи різноманітність термінологій, концептуальних рамок і парадигм моделювання, що використовуються в різних галузях. Для розв'язання цих проблем було запропоновано кілька практик – від ручного експертного посередництва і до повністю автоматизованих методів семантичного узгодження. Одна з відомих стратегій передбачає розробку онтологій верхнього рівня, що охоплюють високорівневі, незалежні від домену концепції, на які можна зіставити специфічні для домену онтології. Така багаторівнева архітектура онтологій забезпечує функціональну сумісність, зберігаючи при цьому специфіку і багатство кожного домену. Інший метод передбачає семантичну анотацію і тегування елементів моделі з посиланнями на спільні словники, що полегшує вирівнювання за допомогою заходів семантичної схожості та алгоритмів міркувань. Агенти семантичного посередництва, часто реалізовані в мультиагентних системах, швидко перекладають і узгоджують знання між моделями в режимі реального часу, використовуючи онтологічні відображення і правила трансформації. Гібридні підходи, що поєднують символічні міркування з методами машинного навчання, також набули популярності, дозволяючи системам вивчати відповідності та адаптувати семантичні відображення на основі шаблонів використання та зворотного зв'язку.

Ефективність різних практик до семантичної сумісності можна порівняти за такими параметрами, як-от масштабованість, рівень автоматизації, адаптивність до предметної області та складність інтеграції. У таблиці 1 представлено порівняльний огляд окремих методів, що використовуються для забезпечення семантичної сумісності в міждисциплінарних інженерних системах.

Таблиця 1

Підходи до забезпечення семантичної сумісності

Підхід	Масштабованість	Рівень автоматизації	Адаптивність до домену	Складність інтеграції
Відображення онтології верхнього рівня	Високий	Середній	Високий	Середній
Семантична анотація та тегування	Середній	Низький	Високий	Низький
Агенти семантичного посередництва	Середній	Високий	Середній	Високий
Гібридні символічно-навчальні моделі	Високий	Високий	Високий	Високий
Ручна експертна координація	Низький	Відсутній	Високий	Низький

Джерело: авторська розробка

Архітектурні принципи побудови інтелектуальної системи підтримки прийняття рішень (СППР) на основі семантичного узгодження моделей зумовлені необхідністю забезпечення узгодженої та динамічної інтеграції гетерогенних джерел знань у контексті розв'язання складних інженерних задач. В умовах, коли інженерні системи стають все більш мультидисциплінарними та інформаційно насиченими, здатність СППР інтелектуально координувати моделі, інтерпретувати їх семантику та підтримувати прийняття контекстно чутливих рішень набуває стратегічного значення. Архітектурний дизайн

таких систем повинен забезпечувати високий ступінь гнучкості, розширюваності та семантичної обізнаності, що дозволяє не лише обробляти структуровану інформацію, але й інтегрувати різноманітні моделі предметної області, які розвиваються.

В основі цієї архітектури знаходиться концептуальна модель, яка визначає високорівневу структуру системи та потік взаємодії між її компонентами. Система, як правило, організована на основі багаторівневої архітектури, яка складається з рівня презентації, рівня семантичної інтеграції, рівня управління знаннями та моделями, а також рівня аргументації та підтримки прийняття рішень. Рівень представлення полегшує взаємодію з користувачем, дозволяючи інженерам та експертам візуалізувати альтернативи рішень, маніпулювати моделями та вхідними обмеженнями. Під цим шаром знаходиться шар семантичної інтеграції, який виконує семантичне узгодження та вирівнювання моделей, посиляючись на онтології, відображаючи концептуальні сутності та забезпечуючи термінологічну узгодженість між доменами. Цей рівень виконує роль семантичного фундаменту, який забезпечує змістовну взаємодію між різними модельними структурами, що в іншому випадку є розрізненими. Рівень управління знаннями та моделями відповідає за зберігання, пошук та еволюцію моделей й артефактів знань. Він забезпечує інфраструктуру для управління бібліотеками моделей, контролю версій і метаданих, необхідних для відстеження походження і застосовності моделей. І нарешті, рівень міркувань і підтримки прийняття рішень виконує алгоритми виведення, процедури оптимізації та імітаційного моделювання, перетворюючи семантично узгоджені моделі на варіанти рішень, що можна застосувати на практиці.

Функціональні компоненти такої СППР мають логічне розділення, але тісно інтегровані між собою для підтримання цілісності системи. Основні компоненти охоплюють модуль інтерфейсу користувача, менеджер онтологій, механізм семантичного узгодження, репозиторій моделей, механізм трансформації та механізм виведення рішень. Менеджер онтологій керує використанням і розвитком онтологій, забезпечуючи належну формалізацію і доступність концепцій, специфічних для предметної області. Механізм семантичної відповідності виконує автоматизоване або напівавтоматизоване виявлення відповідностей між елементами моделі, вирішуючи концептуальні невизначеності та забезпечуючи міжмодельну сумісність. Механізм трансформації відповідає за перетворення моделей між різними форматами представлення або рівнями абстракції на основі встановлених семантичних відображень. Репозиторій моделей забезпечує структуроване зберігання інженерних моделей, зокрема параметрів, обмежень та взаємозв'язків. Механізм виведення рішень застосовує узгоджені семантичні моделі для генерації проєктних альтернатив, оцінки компромісів і ранжування рішень відповідно до їхнього багатокритеріального аналізу. Ці компоненти взаємодіють через сервіс орієнтовану або керовану повідомленнями інфраструктуру, забезпечуючи модульність і масштабованість, особливо в розподілених або хмарних середовищах розгортання.

Для підтримки ефективного прийняття рішень у складних інженерних задачах система повинна відповідати низці критичних вимог. По-перше, вона повинна забезпечувати високий ступінь семантичної прозорості, що дозволяє користувачам відстежувати, як елементи моделі узгоджуються, трансформуються і використовуються в процесі прийняття рішень. По-друге, вона повинна підтримувати багаторівневу роботу з абстракціями, де моделі, починаючи від концептуальних ескізів і закінчуючи детальними симуляціями, можуть бути узгоджено інтегровані. По-третє, система повинна забезпечувати реагування в реальному часі на вхідні дані користувача та оновлення моделей, особливо в контексті ітеративного проєктування. По-четверте, вона повинна пропонувати надійні механізми для управління невідповідностями та конфліктами, що виникають через розбіжності в моделях. Насамкінець адаптивність є основним чинником: система повинна розвиватися з упровадженням нових парадигм моделювання, знань про предметну область і вимог користувачів, підтримувана механізмами навчання та адаптації на основі поведінки й результатів користувачів.

Практичні реалізації архітектур СППР на основі семантичного узгодження моделей можна знайти в різних галузях інженерії. Наприклад, в аерокосмічній інженерії такі системи використовуються для координації аеродинамічних, структурних моделей та моделей систем управління на етапі попереднього проєктування, забезпечуючи досягнення цільових показників при збереженні технологічності. У мехатроніці семантичне узгодження моделей дозволяє інтегрувати механічні, електричні та програмні моделі, підтримуючи спільне моделювання та паралельне проєктування. У будівельній інженерії системи ISPR були розроблені для координації інформаційних моделей будівель (BIM) з інструментами оцінки вартості, енергетичного аналізу та оцінки стійкості. Майбутні напрямки впровадження передбачають розгортання цих систем у цифрових середовищах-двійниках, де семантична координація віртуальних і фізичних моделей у реальному часі підтримує прийняття оперативних рішень і прогнозоване обслуговування. Крім того, планується, що інтеграція з хмарними платформами й архітектурами мікросервісів підвищить масштабованість і функціональну сумісність рішень ISPR.

У таблиці 2 узагальнено основні архітектурні компоненти СППР на основі семантичної моделі, а також їхні ролі та типові технології реалізації.

Таблиця 2

Архітектурні компоненти СППР на основі семантичної моделі

Компонент	Роль в архітектурі	Технології реалізації
Модуль інтерфейсу користувача	Забезпечує взаємодію з користувачем, візуалізацію моделі та управління сценаріями	Вебдашборди, фреймворки з графічним інтерфейсом (наприклад, Qt)
Менеджер онтологій	Керує доменними онтологіями та їх еволюцією	OWL, RDF, Protégé, SPARQL
Механізм семантичного зіставлення	Вирівнює та відображає семантичні елементи між моделями	Логіка опису, NLP, алгоритми зіставлення графів
Репозиторій моделей	Зберігання та керування інженерними моделями та метаданими	NoSQL/графові бази даних, сховище на основі XML/JSON
Механізм трансформації	Перетворює та синхронізує моделі різних форматів та рівнів	Мови перетворення моделей (ATL, QVT)
Механізм виведення рішень	Виводить рішення, використовуючи міркування та оптимізацію	Розв'язувачі обмежень, ML-моделі, механізми виведення

Джерело: авторська розробка

Перспективи розвитку та інтеграції інтелектуальних систем у платформи цифрового проєктування відображають активний розвиток технологій, зміну інженерних парадигм та зростання попиту на швидкі, адаптивні та високоточні проєктні рішення. Сучасні середовища цифрового інжинірингу зазнають фундаментальної трансформації, зумовленої конвергенцією обчислювальних потужностей, доступності даних і штучного інтелекту, що в сукупності переосмислюють те, як створюються, оцінюються і вдосконалюються складні технічні об'єкти. Ця трансформація більше не обмежується автоматизацією рутинних завдань моделювання, а поширюється на інтелектуальне доповнення творчих та аналітичних процесів, дозволяючи платформам для проєктування діяти не просто як інструменти, а як проактивні учасники процесу прийняття інженерних рішень.

Однією з основних тенденцій розвитку цифрових інженерних середовищ є перехід до інтегрованих хмарних платформ, що забезпечують безперебійну сумісність між різними інструментами проєктування та дисциплінами. Ці середовища призначені для підтримки безперервних спільних робочих процесів, де моделі обмінюються, модифікуються і перевіряються в режимі реального часу розподіленими командами. Поява системної інженерії на основі моделей (MBSE) як основної методології посилила потребу в платформах, що можуть керувати і пов'язувати абстракції системного рівня протягом усього життєвого циклу продукту. У цьому контексті цифрові потоки і цифрові двійники набули популярності як організаційні конструкції, що дозволяють відстежувати кожне проєктне рішення, зміну параметрів і результат моделювання протягом усього процесу розробки й впровадження. Ці конструкції вимагають надійних цифрових платформ, що не лише зберігають і обмінюються даними, але й розуміють семантичні зв'язки між інженерними артефактами.

Штучний інтелект та аналітика об'ємних даних відкривають нові можливості в цьому новому просторі. Застосування штучного інтелекту в платформах цифрового дизайну варіюється від генеративного дизайну, коли алгоритми автоматично створюють і оцінюють тисячі варіантів дизайну, до інтелектуальних рекомендаційних систем, що допомагають інженерам у виборі оптимальних матеріалів, компонентів або конфігурацій на основі минулих успіхів і невдач. Моделі машинного навчання, особливо глибокого навчання і навчання з підкріпленням, все частіше використовуються для вилучення закономірностей з історичних проєктних даних, результатів моделювання і зворотного зв'язку з експлуатації. Цей інтелект, заснований на даних, дозволяє платформам прогнозувати результати роботи, оцінювати надійність і передбачати потенційні недоліки дизайну ще до того, коли будуть створені фізичні прототипи. Технології об'ємних даних підтримують агрегацію, обробку та візуалізацію величезних масивів даних, отриманих у результаті високоточного моделювання, сенсорних мереж та взаємодії з користувачами. Ці технології дозволяють виявляти приховані інсайти, що допомагають приймати кращі проєктні рішення, оптимізувати розподіл ресурсів та скорочувати час і витрати на інженерні ітерації.

У цій цифровій екосистемі, що розвивається, інтелектуальні системи підтримки прийняття рішень, як-от ISPR, відіграють головну роль у трансформації процесів проєктування. Ці системи більше не

є ізольованими аналітичними інструментами, а вбудовуються в ширші платформи проєктування для надання контекстно орієнтованих, адаптивних та інтерактивних рекомендацій. Їх інтеграція дозволяє цифровим платформам підтримувати не лише виконання рішень, але і їх формулювання, допомагаючи користувачам визначати проблемні області, виявляти обмеження та оцінювати компроміси між конкурентними цілями. Системи ISPR діють як інтелектуальні посередники між людьми-дизайнерами й величезним простором проєктування, обмежуючи можливі варіанти та пропонуючи інноваційні альтернативи, що можуть бути не одразу очевидними за допомогою традиційних методів. Їх здатність міркувати в умовах невизначеності, інтегрувати знання з різних галузей і вчитися на основі відгуків користувачів тісно пов'язана з потребами інженерного проєктування наступного покоління.

Інтеграція ISPR у цифрові платформи підтримує перехід до гнучких, паралельних і орієнтованих на користувача методологій проєктування. У таких умовах проєктні рішення мають прийматися швидко, часто на основі неповної інформації й з урахуванням міждисциплінарних особливостей. Інтелектуальні системи підтримки прийняття рішень сприяють цьому, підтримуючи узгодженість між взаємозалежними моделями, виявляючи конфлікти та забезпечуючи простежуваність вибору. Вони також посилюють персоналізацію процесу проєктування, адаптуючи свою поведінку та результати до вподобань, досвіду та цілей окремих користувачів або проєктних команд. Оскільки цифрове середовище дизайну все більше охоплює парадигми співпраці, ISPR сприяє управлінню когнітивним навантаженням розподіленого прийняття рішень та підтримці переговорів і досягнення консенсусу між стейкхолдерами.

У перспективі передбачається, що розвиток інтелектуальних систем на цифрових платформах продовжуватиметься у взаємодії з прогресом у сфері когнітивних обчислень, пояснювального штучного інтелекту та графів знань. Когнітивні обчислення дозволять системам імітувати аспекти людського мислення, що дасть їм змогу вести діалоги з дизайнерами, обґрунтовувати рекомендації та адаптуватися до мінливого контексту в режимі реального часу. Пояснювальний ШІ підвищить довіру та прозорість автоматизованих рішень, полегшивши користувачам розуміння та перевірку обґрунтувань, що лежать в основі результатів роботи системи. Графи знань, структуровані представлення взаємопов'язаних понять і сутностей стануть семантичною основою, що пов'язуватиме інженерні дані, моделі та рішення, сприяючи глибшій і точнішій інтеграції знань. Водночас упровадження периферійних обчислень і технологій 5G дозволить підтримувати прийняття рішень у режимі реального часу навіть у мобільних або обмежених ресурсами середовищах, розширюючи сферу застосування і швидкість реагування інтелектуальних систем у практичних інженерних робочих процесах.

Висновки та перспективи подальших досліджень. У результаті проведеного дослідження було обґрунтовано доцільність використання семантичних технологій у побудові архітектури інтелектуальної системи підтримки прийняття рішень. Запропонована архітектура орієнтована на підвищення контекстної точності прийняття рішень шляхом забезпечення семантичного узгодження моделей, що відображають різні рівні знань, даних і потреб користувачів. Детальний аналіз сучасних підходів дозволив визначити основні функціональні компоненти системи, серед яких важливе місце займають модулі онтологічного узгодження, управління знаннями та інтерактивної взаємодії з користувачем.

Сформульовані вимоги до архітектури ІСППР на основі семантичного моделювання створюють концептуальну основу для реалізації адаптивних і масштабованих рішень, здатних ефективно функціонувати в умовах змінного інформаційного середовища. Запропонована структура дозволяє гнучко інтегрувати різноманітні джерела даних та забезпечити цілісність інформаційного простору для підтримки рішень. Перспективами подальших наукових досліджень є створення прототипу системи та проведення експериментальної перевірки її ефективності в прикладних галузях.

Список використаних джерел:

1. Горда О., Лященко Т., Хроленко В., Тихонова О. Особливості інформаційного моделювання на основі метафор роїв. *Управління розвитком складних систем*. 2023. № 56. С. 92–96. DOI: <https://doi.org/10.32347/2412-9933.2023.56.92-96>
2. Євланов М. В., Мороз Б. І., Лучицький В. Я. Інформаційна технологія виявлення термінів та артефактів проєкту у вимогах до інформаційної системи. *АСУ та прилади автоматики*. 2024. № 182. С. 73–93. DOI: <https://doi.org/10.30837/0135-1710.2024.182.073>
3. Кравчук Я., Ніколаєнко Д., Воробйов І. Інноваційні методи інтеграції інтернету речей у комп'ютерні технології. *Наука і техніка сьогодні*. 2024. № 11(39). DOI: [https://doi.org/10.52058/2786-6025-2024-11\(39\)-898-913](https://doi.org/10.52058/2786-6025-2024-11(39)-898-913)
4. Палагін О. Трансдисциплінарна інтелектуальна інформаційно-аналітична система супроводження процесів реабілітації при пандемії (TISP). *Український журнал фізичної і реабілітаційної медицини*. 2021. № 9 (3-4). С. 31–37. URL: <https://surl.gd/lfyzhr> (дата звернення: 18.04.2025)
5. Петренко М. Г., Бойко М. О., Малахов К. С. Комп'ютерні системи подання та оброблення предметно-орієнтованих знань. *Automation of technological and business processes*. 2024. № 16 (1). С. 42–51. DOI: <https://doi.org/10.15673/atbp.v16i1.2771>

6. Пономаренко М. В. Державна політика України в сфері інтелектуальної власності в контексті обміну даними в міжвідомчій комунікації. *Часопис Київського університету права*. 2023. № 3. С. 148–152. DOI: <https://doi.org/10.36695/2219-5521.3.2023.29>

7. Ткаченко К., Байдак А. Онтологічне моделювання інтелектуальної системи підтримки прийняття рішень під час аналізу ризиків у інноваційно-інвестиційній сфері. *Цифрова платформа: інформаційні технології в соціокультурній сфері*. 2021. № 4 (1). С. 66–78. DOI: <https://doi.org/10.31866/2617-796X.4.1.2021.236948>

8. Тріщ В., Богдан О., Пашковський В., Попович О., Фудорук Е. Засоби підтримки сховищ знань комп'ютерних систем нафтогазової галузі. *Herald of Khmelnytskyi National University. Technical sciences*. 2024. № 335 (1). С. 269–279. DOI: <https://doi.org/10.31891/2307-5732-2024-335-3-36>

9. Чалий С. Ф., Лещинська І. О. Екстерналізація наявних знань в ментальній моделі користувача системи штучного інтелекту. *Bulletin of National Technical University "KhPI". Series: System Analysis, Control and Information Technologies*. 2024. № 1 (11). С. 91–96. DOI: <https://doi.org/10.20998/2079-0023.2024.01.15>

Дата надходження статті: 10.06.2025

Дата прийняття статті: 24.06.2025

Опубліковано: 23.09.2025

UDC 004.056:004.774:621.9.048.5

DOI <https://doi.org/10.32689/maup.it.2025.2.25>

Inna ROZLOMII

Candidate of Technical Sciences, Associate Professor,
Associate Professor at the Department of Information Security and Computer Engineering,
Cherkasy State Technological University, inna-roz@ukr.net
ORCID: 0000-0001-5065-9004

Serhii NAUMENKO

Lecturer at the Department of Information Technologies,
Bohdan Khmelnytsky National University of Cherkasy, naumenko.serhii1122@vu.cdu.edu.ua
ORCID: 0000-0002-6337-1605

Volodymyr SYMONYUK

Candidate of Technical Sciences, Associate Professor,
Associate Professor at the Department of Automation and Computer-Integrated Technologies,
Lutsk National Technical University, v.symonyuk@lntu.edu.ua
ORCID: 0000-0002-7624-4760

Vitalii PTASHENCHUK

Candidate of Technical Sciences, Associate Professor,
Associate Professor at the Department of Automation and Computer-Integrated Technologies,
Lutsk National Technical University, v.ptashenchuk@lntu.edu.ua
ORCID: 0000-0003-1570-7570

Vitalii ZAZHOMA

Candidate of Technical Sciences, Associate Professor,
Associate Professor at the Department of Information Systems and Organization of Civil Protection Measures,
National University of Civil Protection of Ukraine, zazhoma_vitalii@nuczu.edu.ua
ORCID: 0000-0003-2083-2472

SIMPLIFIED CRYPTOGRAPHY FOR SECURE VIBRATION PARAMETERS IN POST-PROCESSING OF 3D-PRINTED PARTS

Abstract. In modern manufacturing, 3D printing technologies are gaining increasing popularity due to their ability to rapidly create complex and customized parts. However, to achieve high quality and meet industrial standards, 3D-printed parts require post-processing. Vibration containers are widely used for grinding, polishing, and other processing methods, improving the appearance and mechanical properties of products. This highlights the need to protect data generated during the vibration process, which is transmitted to central servers for analysis and management.

The aim of this study is to develop a framework for a secure post-processing system for 3D-printed parts, which includes an encryption module based on lightweight ciphers to protect vibration parameters and other operational data.

Methodology. This article proposes the use of lightweight cryptography to ensure the security of vibration parameters in the post-processing of parts printed using 3D printing methods. Lightweight cryptography allows for efficient data encryption with minimal resource consumption, which is critical for systems with limited computational capabilities. The paper discusses various post-processing methods – such as mechanical grinding and polishing – and their impact on the surface quality and compliance with standards. The cryptographic protocols used to protect data during its transmission from vibration containers to central servers are described in detail.

The results can contribute to the advancement of 3D printing technologies and the creation of new post-processing methods that meet contemporary requirements for quality and information security.

Scientific novelty. The use of the SPECK cipher has proven to be the optimal solution for resource-constrained systems, as it provides high encryption speed with low energy consumption, which is critical for microcontrollers and embedded systems used in post-processing operations.

Conclusion. The proposed solutions provide reliable protection against unauthorized access and manipulation, improving overall manufacturing security and efficiency. Special attention is given to the integration of cryptographic methods into monitoring and control systems for vibration processes, ensuring high accuracy and reliability in data processing.

Key words: 3D printing, post-processing, vibration containers, lightweight cryptography, data encryption, mechanical grinding, polishing.

Інна РОЗЛОМІЙ, Сергій НАУМЕНКО, Володимир СИМОНЮК, Віталій ПТАШЕНЧУК, Віталій ЗАЖОМА. ПОЛЕГШЕНА КРИПТОГРАФІЯ ДЛЯ БЕЗПЕКИ ПАРАМЕТРІВ ВІБРАЦІЇ В ПОСТОБРОБЦІ 3D-ДРУКОВАНИХ ДЕТАЛЕЙ

Анотація. У сучасному виробництві технології 3D-друку набувають все більшої популярності завдяки можливості швидкого створення складних та індивідуальних деталей. Однак, для досягнення високої якості та відповідності промисловим стандартам, 3D-друковані деталі потребують постобробки. Вібраційні контейнери широко використовуються для шліфування, полірування та інших методів обробки, покращуючи зовнішній вигляд і механічні властивості виробів. При цьому виникає потреба у захисті даних, які генеруються під час процесу вібрації та передаються на центральні сервери для аналізу та управління.

Метою цього дослідження є розробка структури захищеної системи постобробки деталей, надрукованих методом 3D-друку, яка включає модуль шифрування на основі полегшених шифрів для забезпечення захисту параметрів вібрації та інших операційних даних.

Методологія. У цій статті пропонується використання полегшеної криптографії для забезпечення безпеки параметрів вібрації в процесі постобробки деталей, надрукованих методом 3D-друку. Полегшена криптографія дозволяє ефективно шифрувати дані з мінімальними витратами ресурсів, що є критичним для систем з обмеженими обчислювальними можливостями. Розглядаються різні методи постобробки, включаючи механічне шліфування та полірування, а також їх вплив на якість поверхні та відповідність стандартам. Детально описуються криптографічні протоколи, які використовуються для захисту даних під час їх передачі від вібраційних контейнерів до центральних серверів.

Результати можуть бути використані для подальшого вдосконалення технологій 3D-друку та розробки нових методів постобробки, що відповідають сучасним вимогам до якості та інформаційної безпеки.

Наукова новизна. Використання шифру SPECK виявилось оптимальним рішенням для систем з обмеженими ресурсами, оскільки він забезпечує високу швидкість шифрування при низькому енергоспоживанні, що є критично важливим для мікроконтролерів та вбудованих систем, які використовуються в процесах постобробки.

Висновки. Запропоновані рішення забезпечують надійний захист від несанкціонованого доступу та маніпуляцій, підвищуючи загальну безпеку та ефективність виробничих процесів. Особлива увага приділяється інтеграції криптографічних методів у системи моніторингу та контролю вібраційних процесів, що дозволяє забезпечити високу точність та надійність обробки даних.

Ключові слова: 3D-друк, постобробка, вібраційні контейнери, полегшена криптографія, шифрування даних, механічне шліфування, полірування.

Introduction. A prominent trend in modern manufacturing is the implementation of technologies related to 3D printing, which enables the creation of parts with various levels of complexity at high speed and precision [5]. Additionally, 3D printing technology often requires an additional stage to produce of high-quality printed parts, known as post-processing [16].

An important aspect of the 3D printing process is post-processing, as it enhances the quality, appearance, and mechanical properties of printed parts [3, 5]. Different 3D printing technologies, such as stereolithography (SLA), fused deposition modeling (FDM), and selective laser sintering (SLS), require specialized post-processing methods to address various issues, such as layer lines, support marks, and rough surfaces. Resolving these issues enables printing parts that meet industry standards and achieve the desired functionality and aesthetics.

According to research [1], the 3D printing market continues to grow across various sectors. For example, in the medical field, 3D printing is actively used to create prosthetics, implants, and surgical instruments, contributing to its rapid growth. In the aerospace industry, 3D printing is used to produce complex and lightweight components, driving significant growth in this segment. The automotive industry utilizes 3D printing for prototyping and small-scale production of parts, allowing for a substantial reduction in production time and costs. Figure 1 shows a graph depicting the growth in the number of 3D-printed parts across various sectors, including aerospace, automotive, medical industries, and consumer goods, based on real data [1].

As shown, all sectors demonstrate significant growth in the number of manufactured parts, particularly in the automotive and medical industries. 3D printing is becoming increasingly popular due to its ability to create complex and customized parts [4, 18]. The quality and appearance of a 3D-printed object often depend on the technology and post-processing stages [7]. Post-processing technology not only improves the aesthetics of the printed part but can also enhance its mechanical properties according to technical requirements [9].

Post-processing plays an important role in ensuring that 3D-printed parts meet the necessary quality and aesthetic standards. For instance, the SLS 3D printing process may result in visible layer lines on the surface of a part, which can diminish its overall aesthetic quality. A rough surface texture is also a common issue with parts printed on a 3D printer. However, post-processing methods such as grinding, polishing, and painting can effectively eliminate or minimize these defects [7].

In industries such as aerospace, automotive, and medical, there are often strict regulations and standards regarding the quality, performance, and appearance of parts [9]. For example, manufacturers typically use ASME Y14.36 or ISO 21920-1 to specify surface texture. Post-processing technologies, such as grinding, polishing,

smoothing, and coating, help ensure that 3D-printed parts comply with these standards, making them suitable and attractive for a variety of applications [6].

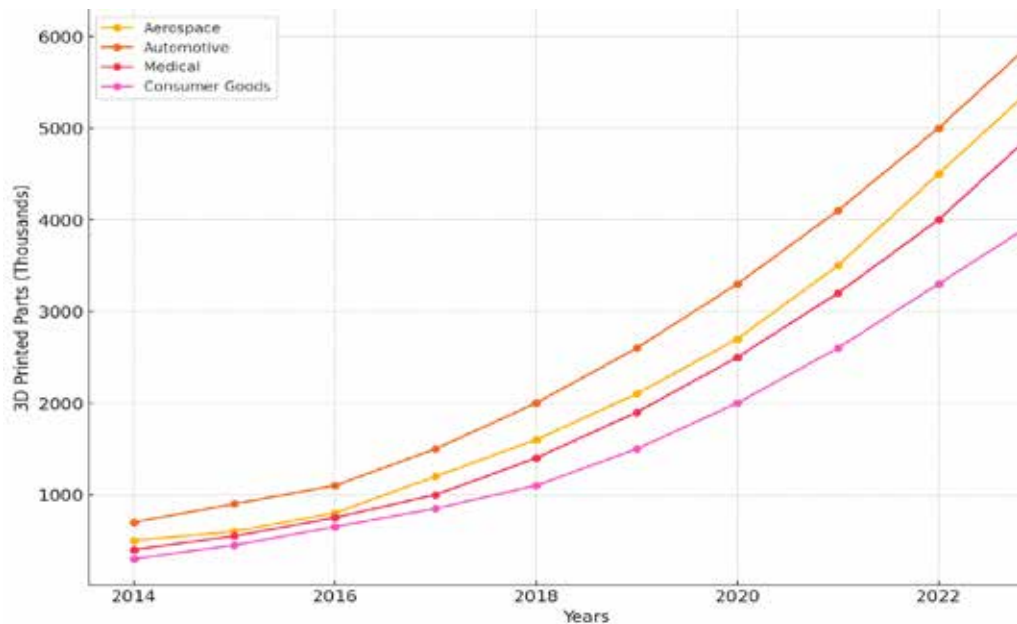


Fig. 1. Growth in the Number of 3D-Printed Parts Across Various Sectors

In addition to improving the quality and appearance of parts, protecting the information generated during the post-processing process is also crucial. Vibration parameters and other information should be encrypted before being transmitted to the central server for analysis and management. This ensures data security and prevents unauthorized access.

The aim of this study is to develop a framework for a secure post-processing system for 3D-printed parts, which includes an encryption module based on lightweight ciphers to protect vibration parameters and other operational data.

Related Works. The scientific community's interest in 3D printing technologies and post-processing methods for printed parts has significantly increased in recent years. In particular, several authors in their studies [6, 7] have analyzed various post-processing methods to improve the surface quality of 3D-printed parts, including grinding and polishing. They found that the application of mechanical grinding can significantly reduce surface roughness and enhance the overall aesthetics of the parts.

In the study [12], an analysis was conducted on the impact of different vibration modes on the quality of parts processed using abrasive vibration methods. The authors discovered that optimizing vibration parameters can greatly improve the efficiency of grinding and polishing.

The importance of cryptographic data protection in industrial applications is also a subject of research for many scientists. In [15], modern methods of lightweight cryptography that can be applied to protect information in IoT devices are discussed. The authors propose the use of efficient cryptographic protocols to protect transmitted data in resource-constrained systems. In [8], the application of cryptographic methods for data protection in industrial systems, including vibration technologies, is explored.

However, the integration of cryptographic methods into the post-processing of 3D-printed parts, particularly the protection of vibration parameters during grinding and polishing, remains underexplored. An important aspect is the development of effective methods for encrypting data generated by vibration containers and transmitting it to central servers for further analysis and management. This would ensure a high level of data security and improve the overall efficiency of manufacturing processes.

Therefore, there is a need for further research aimed at integrating cryptographic protocols into the post-processing of 3D-printed parts, particularly the development of methods to protect vibration parameters and other critical data.

Research Methodology. The research methodology is based on the use of various materials and methods for post-processing 3D-printed parts, as well as the development of data protection methods using lightweight cryptography. The main stages of the research include sample preparation, post-processing, data collection and analysis, and ensuring data security.

The study utilizes samples printed using SLA, FDM and SLS technologies. Each sample undergoes different post-processing methods, such as grinding, polishing, and painting, to determine the optimal processing parameters. The samples are made from various materials, including photopolymers, thermoplastics, and metal powders.

The research methodology includes the following stages:

1. Sample preparation. In the first stage, samples are printed using three different 3D printing technologies: SLA, FDM, and SLS [13]. The initial surface parameters for each sample are determined, including surface roughness, visible layer lines, and other defects.
2. Sample post-processing. In the second stage, the samples undergo various post-processing methods, including mechanical grinding, polishing, and painting. Different abrasive materials and polishing pastes are used for each method. The vibration modes for the containers in which the processing takes place are optimized to achieve the best results.
3. Data collection and analysis. After the post-processing is completed, the surface quality of the samples is analyzed. Parameters such as roughness, visibility of layer lines, and overall aesthetic appearance are evaluated. Data is collected using specialized equipment, including profilometers and optical microscopes.
4. Data protection. An important aspect of the research is ensuring the security of the data generated during the post-processing process. Lightweight cryptography methods are used to effectively encrypt the data with minimal resource expenditure. The data is encrypted before being transmitted to central servers for further analysis and management.
5. Integration of cryptographic protocols. In the final stage, cryptographic protocols are developed and implemented to protect the data during transmission from the vibration containers to the central servers. The choice of encryption methods is based on their efficiency and ability to integrate into monitoring and control systems for vibration processes.

This research methodology ensures high quality of 3D-printed parts and reliable protection of data generated during post-processing. The results of the study can be used to improve 3D printing technologies and develop new post-processing methods that meet modern quality and security standards.

Vibration Processing Techniques. To ensure high quality of 3D-printed parts, several post-processing methods are used, including grinding, polishing, and painting. Grinding and polishing help to remove surface roughness and visible layer lines that form during printing [6]. Various abrasive materials and polishing pastes are employed to achieve the desired level of smoothness and aesthetic appearance.

An important aspect of post-processing is selecting the correct vibration modes for the containers in which the parts are processed [12]. Optimal vibration modes help to improve the quality of grinding and polishing by ensuring even surface treatment of the parts. Mechanical grinding and polishing are most effective when processing multiple parts simultaneously, as this is economically advantageous. However, considering certain design features of the parts being processed, their size, production scale, and other technical, organizational, and economic factors of the technological process, it may also be feasible to process one or more parts at a time.

Processing occurs in a specialized vibratory bowl filled with processing media [15]. For simultaneous processing of a large number of parts, bowls with ring-shaped, parabolic, or toroidal geometry (Fig. 2: a, b, c), and other rounded forms are preferred.

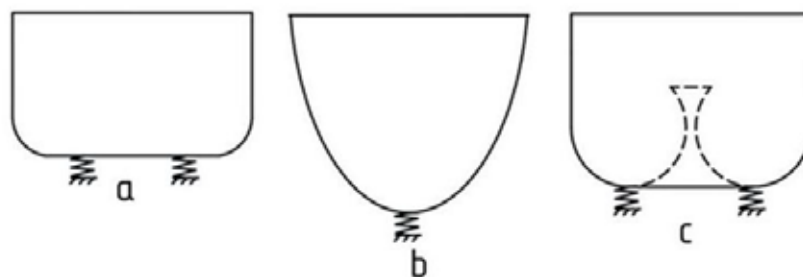


Fig. 2. Diagrams of the Most Common Shapes of Vibratory Bowls

To process a part of the surface of the components, it is necessary to protect the unprocessed part with certain means or special fixtures (Fig. 3).

The processing media in the vibratory bowl can include various types of abrasive materials or their mixtures combined with softening agents and liquid solutions. The selection and application of these media depend on numerous technical, economic, scientific, research, and other factors.

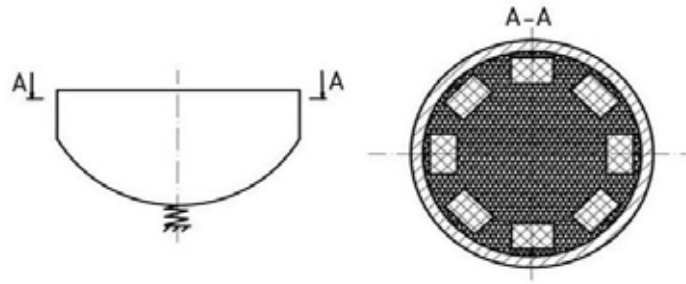


Fig. 3. Diagram of Part Placement with Fixtures in a Vibratory Bowl with Processing Media

Structure of a Secured Post-Processing System for 3D-Printed Parts. During the post-processing of 3D-printed parts, there is a need to protect the data generated by vibratory containers and transmitted to central servers for further analysis and management. Lightweight cryptography methods are used to effectively encrypt the data with minimal resource expenditure.

Vibration parameters and other information must be encrypted before transmission to the central server for analysis and management. This ensures the confidentiality and integrity of the transmitted data and prevents unauthorized access and modification.

Data transmission and processing in the 3D-printed parts post-processing system occur in several stages, each of which is crucial for ensuring the security and quality of the finished products [8]. As shown in the structural diagram, data collected by sensors is processed by a local controller for preliminary processing, then encrypted and transmitted via a secure communication channel to the central server. On the server, the data is decrypted and analyzed to detect potential defects or deviations, allowing for timely corrective actions. This process ensures data confidentiality and enhances production efficiency.

Figure 4 presents the structural diagram of the 3D-printed parts post-processing system, which includes an integrated encryption module to ensure data transmission security.

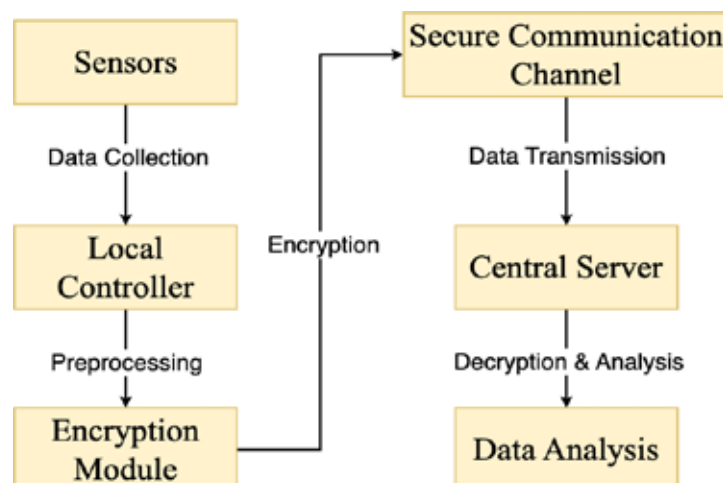


Fig. 4. Structural Scheme of the Secured 3D-Printed Parts Post-Processing System

The post-processing of 3D-printed parts involves the following stages:

1. Sensors collect vibration parameters and other necessary data during the post-processing process.
2. A local controller receives data from the sensors and performs preliminary processing, filtering, and aggregation of the data.
3. After preliminary processing, the data is transmitted to the encryption module, where it is encrypted to ensure security.
4. The encrypted data is transmitted via a secure communication channel to the central server.
5. On the central server, the data is decrypted and further analyzed.
6. The data analysis process allows for the detection of potential defects or deviations, after which necessary corrective actions can be taken.

Upon completing all stages of post-processing, the system not only ensures high quality of the 3D-printed parts but also guarantees the security and confidentiality of technological data. With the integration of the encryption module, the collected data is protected from unauthorized access, preserving the integrity of the information during transmission and subsequent processing on the central server. This enhances the reliability of the entire post-processing process and minimizes risks associated with information leakage or external threats.

Integration of Cryptographic Protocols. Cryptographic protocols provide robust protection against unauthorized access and tampering, which is critical for systems with limited computational capabilities. The use of lightweight encryption algorithms ensures data security without significantly impacting system performance. A key aspect is the integration of these protocols into vibration process monitoring and control systems, enabling high accuracy and reliability in data processing.

Encrypting data is essential for protecting sensitive information from cyberattacks [13]. Vibration parameters and other data collected during the post-processing of 3D-printed parts may contain confidential technological information that needs to be protected from competitors and malicious actors. Encryption ensures that even if the data is intercepted, it remains inaccessible to unauthorized parties.

In post-processing 3D-printed parts, devices with limited resources, such as microcontrollers or embedded systems, are often used. Lightweight ciphers, unlike traditional ones, are specifically designed to operate under conditions of limited computational power and energy consumption. They provide the necessary level of security while not overburdening the system or affecting performance.

Encryption algorithms in such systems are implemented both in hardware and software [17]. Hardware modules, such as specialized chips or embedded cryptographic processors, can provide high encryption and decryption speeds with minimal energy consumption. This is especially important for devices that run on batteries or have limited energy resources. On the other hand, software implementation of encryption algorithms allows for greater flexibility in adapting to the specific requirements of the system [2].

In both cases, the main goal is to ensure data security without significantly impacting system performance. This enables secure data processing in low-resource systems.

In environments with limited computational resources, such as post-processing of 3D-printed parts, it is crucial to choose cryptographic algorithms that are both lightweight and effective. Several low-resource ciphers have been developed specifically for such applications, providing a balance between security and performance.

In such systems, it is important to encrypt not only vibration parameters but also other critical data generated during the post-processing of 3D-printed parts. This includes temperature readings, the speed and direction of material movement, equipment technological parameters, configuration files and programs that define system operating modes, and the results of surface quality measurements, including roughness and layer line visibility. Additionally, data on energy consumption during processing also needs to be encrypted to prevent unauthorized access to confidential information that may contain important technological details.

The application of cryptographic algorithms in post-processing systems for 3D-printed parts is a crucial step in ensuring information security. Given the limited computational resources and energy consumption of such systems, special attention must be paid to selecting ciphers that can provide reliable data protection without overburdening the system. Some of these ciphers include SPECK, SIMON, and PRESENT, which have been specifically designed for efficient operation in resource-constrained environments [14]. They provide a high level of security with minimal energy consumption, making them an ideal choice for use in embedded systems and microcontrollers used in post-processing operations.

SPECK is a lightweight block cipher optimized for software implementations in resource-constrained environments, such as post-processing systems for 3D-printed parts [11]. Due to its simple structure and flexibility, SPECK provides high encryption speed with minimal computational resource consumption, making it an ideal choice for data protection in devices that operate on batteries or have limited energy and computational capabilities. It employs a key schedule (1). In this study, the SPECK64/96 configuration was selected, which operates on a 64-bit data block and uses a 96-bit key. This configuration offers an excellent trade-off between security level and efficiency in constrained environments.

$$K_i = \left(K_{(i-1)} \oplus F_{(i-1)} \left(K_{(i-2)}, K_{(i-3)} \right) \right) \quad (1)$$

where $F_{i-1}(K_{i-2}, K_{i-3})$ is a round function that combines keys from previous rounds.

The encryption round function is represented by the expression (2).

$$F(x, y) = ((x \gg 8) + y) \oplus K_i \quad (2)$$

This function shifts the value of x 8 bits to the right, adds it to y , and the result is XOR-ed with the key K_i .

SPECK is characterized by low energy consumption and high data processing speed, making it ideal for use in battery-powered devices. Cryptanalysis studies have demonstrated that SPECK64/96 offers sufficient resistance against differential and linear attacks in non-critical embedded systems.

SIMON is a lightweight block cipher designed for hardware implementations with minimal requirements for logic elements and power consumption. It provides high efficiency and flexibility in various embedded systems, particularly where computational resources are limited. SIMON supports multiple block and key sizes, allowing it to be adapted to specific security and performance requirements. SIMON is particularly well-suited for hardware implementations, offering low logic element count and energy consumption, making it ideal for battery-powered devices in post-processing systems. The cipher uses a key schedule (3).

$$K_{(i+1)} = K_i \oplus \left(K_{(i-1)} \oplus \text{RotateLeft}(K_{(i-2)}, 8) \right) \quad (3)$$

Here, the next round key is obtained by XOR-ing the current key with the result of rotating the key from two rounds ago by 8 bits and then XOR-ing it with the previous key.

The encryption round function is represented by the expression (4).

$$F(x) = ((x \gg 1) \wedge (x \gg 8)) \oplus (x \gg 2) \quad (4)$$

This function processes the value of x by shifting it 1 and 8 bits to the right and applying the AND operation. Then the result is XOR-ed with x shifted by 2 bits to the right. The compact design of SIMON ensures its efficient integration into embedded systems with minimal impact on overall performance.

PRESENT is a lightweight block cipher designed for use in resource-constrained environments, such as embedded systems and microcontrollers. It operates with a fixed block length of 64 bits and a key of either 80 or 128 bits. PRESENT is particularly effective in environments where memory and computational power are severely limited [10]. PRESENT uses a key schedule (5).

$$K_i = \text{SubBytes}(\text{RotateLeft}(K_{(i-1)}, 61)) \oplus i \quad (5)$$

The key is processed by the SubBytes, function, which is applied to the key, then rotated 61 bits to the right, and the result is XOR-ed with the round number i .

The round function for encryption is represented by expression (6).

$$F(x) = \text{SBox}(x) \oplus \text{Perm}(x) \quad (6)$$

In this function, the value x is first passed through the SBox (substitution table), and then the result is XOR-ed with the permuted value of x .

Given the need for high processing speed and flexibility in software implementations, which are characteristic of post-processing systems for 3D-printed parts, the SPECK cipher is the most suitable for protecting vibration parameters and other operational data. SPECK provides the necessary level of security with minimal impact on computational resources, which is critical for such systems.

Results and Discussion. As a result of the research, a post-processing system for 3D-printed parts with an integrated encryption module has been proposed. The post-processing system includes vibration containers equipped with sensors that measure vibration parameters and other technological indicators during the processing of 3D-printed parts. The collected data is sent to a local controller, where it undergoes preprocessing, including filtering and aggregation. The data is then encrypted before being transmitted to a central server for further analysis.

The integrated encryption module is implemented as a separate component at the preprocessing stage. This stage handles the data generated during vibration processing. The main function of the module is to protect confidential technological information from unauthorized access during its transmission to the central server for further analysis. The encryption module is integrated into the local controller, which receives data from the sensors in the vibration containers. In this research, the encryption module was implemented on a resource-constrained microcontroller platform, specifically the **STM32F103** based on ARM Cortex-M3 architecture. This controller is widely used in industrial applications due to its low power consumption and real-time processing capabilities. The SPECK cipher was embedded into the firmware using a lightweight cryptographic library optimized for this microcontroller. Several lightweight cryptographic algorithms were considered for encrypting the data. Among them, SPECK was selected due to its optimal balance of software implementation efficiency, encryption speed, and minimal RAM/ROM footprint, making it highly suitable for microcontrollers commonly used in post-processing control systems. Its performance was benchmarked against SIMON and PRESENT in terms of energy efficiency and throughput, where SPECK showed a favorable trade-off for systems without hardware acceleration. It provides high encryption speed, low energy consumption, and minimal

computational resource requirements, which is critical for systems with limited capabilities, such as microcontrollers used in post-processing of 3D-printed parts. SPECK effectively protects confidential technological data without impacting the overall system performance, making it an ideal solution for such conditions.

Conclusion. As a result of the conducted research, a structure for a secure post-processing system for 3D-printed parts has been developed, which includes an integrated encryption module based on lightweight cryptographic algorithms. The proposed system ensures high surface processing quality by optimizing vibration process parameters and effectively encrypting data generated during post-processing. The use of the SPECK cipher has proven to be the optimal solution for resource-constrained systems, as it provides high encryption speed with low energy consumption, which is critical for microcontrollers and embedded systems used in post-processing operations. The results can contribute to the advancement of 3D printing technologies and the creation of new post-processing methods that meet contemporary requirements for quality and information security.

Bibliography:

1. Fortune Business Insights. (n.d.). 3D printing market size, share & COVID-19 impact analysis. URL: <https://www.fortunebusinessinsights.com/industry-reports/3d-printing-market-101902> (February 10, 2025)
2. Hozdić E. Characterization and Comparative Analysis of Mechanical Parameters of FDM-and SLA-Printed ABS Materials. *Applied Sciences*. 2024. № 14(2). P. 649. <https://doi.org/10.3390/app14020649>
3. Huang J., Qin Q., Wang J. A review of stereolithography: Processes and systems. *Processes*. 2020. № 8(9). P. 1138. <https://doi.org/10.3390/pr8091138>
4. Kafle A., Luis E., Silwal R., Pan H. M., Shrestha P. L., Bastola A. K. 3D/4D printing of polymers: fused deposition modelling (FDM), selective laser sintering (SLS), and stereolithography (SLA). *Polymer*. 2021. № 13(18). P. 3101. <https://doi.org/10.3390/polym13183101>
5. Karakurt I., Lin L. 3D printing technologies: techniques, materials, and post-processing. *Current Opinion in Chemical Engineering*. 2020. № 28. P. 134–143. <https://doi.org/10.1016/j.coche.2020.04.001>
6. Kopar M., Yildiz A. R. Experimental investigation of mechanical properties of PLA, ABS, and PETG 3-d printing materials using fused deposition modeling technique. *Materials Testing*. 2023. № 65(12). P. 1795–1804. <https://doi.org/10.1515/mt-2023-0202>
7. Kristiawan R. B., Imaduddin F., Ariawan D., Ubaidillah Arifin, Z. A review on the fused deposition modeling (FDM) 3D printing: Filament processing, materials, and printing parameters. *Open Engineering*. 2021. № 11(1). P. 639–649. <https://doi.org/10.1515/eng-2021-0063>
8. Lim J. X. Y., Pham Q. C. Automated post-processing of 3D-printed parts: artificial powdering for deep classification and localisation. *Virtual and Physical Prototyping*. 2021. № 16(1). P. 1–14. <https://doi.org/10.1080/17452759.2021.1927762>
9. Pavlenko P., Teslia I., Khlevna I., Yehorchenkov O., Yehorchenkova N., Kataieva Y., Khlevnyi A., Veretelnik V., Latysheva T., Kubiavka L. Development of a concept of combined project-production activities planning using digital twins. *Eastern-European Journal of Enterprise Technologies*. 2024. № 5 (131(3)). P. 6–17. <https://doi.org/10.15587/1729-4061.2024.311751>
10. Rashidi B. Flexible structures of lightweight block ciphers PRESENT, SIMON and LED. *IET Circuits, Devices & Systems*. 2020. № 14(3). P. 369–380. <https://doi.org/10.1049/iet-cds.2019.0363>
11. Rashidi B. High-throughput and flexible ASIC implementations of SIMON and SPECK lightweight block ciphers. *International Journal of Circuit Theory and Applications*. 2019. № 47(8). P. 1254–1268. <https://doi.org/10.1002/cta.2645>
12. Raza A., Markovic N., Wolf T., Romahn P., Zinn A. H., Kolossa D. Glass Container Fill Level Measurement via Vibration on a Low-Power Embedded System. In *2023 IEEE International Conference on Omni-layer Intelligent Systems (COINS)*, July 2023, pp. 1–6.
13. Rozlomii I., Yarmilko A., Naumenko S. Analysis of Information Security Issues in Balancing Multiple Independent Containers on a Single Server. *CEUR Workshop Proceedings*. 2023. Vol. 3628. P. 450–461. URL: <https://ceur-ws.org/Vol-3628/paper26.pdf>
14. Rozlomii I., Yarmilko A., Naumenko S. Data security of IoT devices with limited resources: challenges and potential solutions. *CEUR Workshop Proceedings*. 2024. Vol. 3666. P. 85–96. URL: <https://ceur-ws.org/Vol-3666/paper13.pdf>
15. Symoniuk V., Denysiuk V., Lapchenko Y., Kaidyk O., Ptachenchuk V. About Trimming Processes of Parts in the Shock-Impulse Load of Vibrobunker. In *Grabchenko's International Conference on Advanced Manufacturing Processes*, September 2019. Cham: Springer International Publishing, 2019, pp. 321–330.
16. Symonyuk V. P., Ptashenchuk V. V., Tymoshchuk A. A. To analysis of the technical state of the equipment for manufacturing parts using 3D printing. *Promising technologies and devices*. 2022. № 21. P. 113–118.
17. Voievodin Y., Rozlomii I. Application Security Optimization in Container Orchestration Systems Through Strategic Scheduler Decisions. *CEUR Workshop Proceedings*. 2024. Vol. 3654. P. 471–478. URL: <https://ceur-ws.org/Vol-3654/short16.pdf>
18. Xu X., Goyanes A., Trenfield S. J., Diaz-Gomez L., Alvarez-Lorenzo C., Gaisford S., Basit A. W. Stereolithography (SLA) 3D printing of a bladder device for intravesical drug delivery. *Mater Sci Eng C Mater Biol Appl*. 2021 Jan; 120:111773. <https://doi.org/10.1016/j.msec.2020.111773>. Epub 2020 Dec 4. PMID: 33545904.

Дата надходження статті: 01.07.2025

Дата прийняття статті: 11.07.2025

Опубліковано: 23.09.2025

UDC 004.438.52:004.415.3:004.421.2
DOI <https://doi.org/10.32689/maup.it.2025.2.26>

Oleh SYPIAHIN

Master's Degree, Information Security,
Full Stack Engineer, floLive (Israel),
fed4wet@gmail.com
ORCID: 0009-0008-7565-3221

Dmytro POLTAVSKYI

Bachelor's Degree,
Engineering Team Lead, Uplandme Inc. (USA),
poltavskyi.dmytro@gmail.com
ORCID: 0009-0009-0387-6677

Roman MARTYNIENKO

Master's Degree,
Senior Software Engineer, Henry AI, Inc. (USA),
worldofwbdesign@gmail.com
ORCID: 0009-0005-4663-8530

EXPLORING THE ADVANTAGES AND LIMITATIONS OF THE FEATURE-SLICED DESIGN ARCHITECTURAL METHODOLOGY IN COMMERCIAL FRONTEND PROJECTS

Abstract. The purpose of this study is to identify the advantages and limitations of the Feature-Sliced Design (FSD) architectural methodology in the development of commercial frontend applications, particularly using Angular-based projects that require a high degree of modularity, adaptability to change, and preservation of structural integrity throughout the software product life cycle. Special attention is given to assessing the feasibility of implementing this approach within small and medium-sized development teams operating in fast-paced environments with short release cycles and high demands for maintainable code.

The research methodology is based on a structural analysis of the theoretical foundations of FSD, practical modeling of software architecture using an open-source ToDo application implemented according to the full hierarchical structure—app, processes, pages, widgets, features, entities, and shared—and a comparative analysis with leading alternative architectural approaches such as Atomic Design, MVVM, and MVC. The comparison was conducted based on criteria such as scalability, cohesion, module decoupling, reusability of components, structural transparency, and support for continuous integration.

The scientific novelty lies in a comprehensive description of the core advantages of FSD, including business logic encapsulation, minimized inter-module dependencies, predictable structural hierarchy, and transparent component interaction logic. For the first time, the study systematizes the practical constraints associated with FSD implementation: the complexity of initial configuration, fragmented documentation, high architectural discipline requirements, and the onboarding difficulties for new developers.

Conclusions. The study confirms the effectiveness of FSD as an architectural paradigm for building stable, scalable, and long-term maintainable frontend systems. However, successful adoption requires a well-prepared development team, comprehensive internal documentation, standardized structuring practices, and strict adherence to established architectural principles. The practical significance of the research lies in offering recommendations for integrating FSD into commercial web applications with elevated requirements for architectural manageability, maintainability, and scalability.

Key words: Feature-Sliced Design, Angular, frontend architecture, scalability, layered structure, modularity, ToDo application, architectural methodology, technical debt, UX composition, project structure.

Олег СИПЯГІН, Дмитро ПОЛТАВСЬКИЙ, Роман МАРТИНЕНКО. ДОСЛІДЖЕННЯ ПЕРЕВАГ ТА ОБМЕЖЕНЬ АРХІТЕКТУРНОЇ МЕТОДОЛОГІЇ FEATURE-SLICED DESIGN У КОМЕРЦІЙНИХ FRONTEND-ПРОЄКТАХ

Анотація. Метою дослідження є виявлення переваг та обмежень архітектурної методології Feature-Sliced Design (FSD) у розробці комерційних frontend-застосунків, зокрема на прикладі Angular-проектів, які вимагають високого ступеня модульності, адаптивності до змін і підтримки структурної цілісності на всіх етапах життєвого циклу програмного продукту. Особливу увагу приділено визначенню доцільності впровадження даного підходу у командах малого та середнього масштабу, що функціонують у середовищі інтенсивної розробки, з короткими релізними циклами та високими вимогами до підтримованості коду.

Методологія дослідження базується на структурному аналізі теоретичних засад FSD, практичному моделюванні архітектури програмного забезпечення на прикладі ToDo-додатку з відкритим кодом, реалізованого

© O. Sypiahin, D. Poltavskyi, R. Martynenko, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

відповідно до всієї ієрархії рівнів – *app, processes, pages, widgets, features, entities* та *shared*. Також здійснено порівняльний аналіз із провідними альтернативними архітектурними підходами – *Atomic Design, MVVM* і *MVC* – на основі критеріїв масштабованості, когезії, рівня зв'язаності модулів, можливості повторного використання компонентів, структурної прозорості та підтримки безперервної інтеграції.

Наукова новизна полягає в комплексному описі ключових переваг FSD, серед яких – інкапсуляція бізнес-функціоналу, мінімізація міжмодульних залежностей, передбачувана ієрархія структури та прозора логіка компонентної взаємодії. Вперше узагальнено обмеження, характерні для практичного застосування методології: складність початкового впровадження, фрагментарність документації, високі вимоги до архітектурної дисципліни та труднощі адаптації нових розробників у команду.

Висновки. Результати дослідження підтверджують ефективність FSD як архітектурної парадигми для побудови стабільних, масштабованих і довготривало підтримуваних *frontend-систем*. Водночас успішне застосування підходу вимагає високої підготовленості команди, внутрішньої документації, уніфікованих підходів до структурування та чіткого дотримання прийнятих принципів. Практична значущість дослідження полягає у формулюванні рекомендацій щодо інтеграції FSD у комерційні вебзастосунки з підвищеними вимогами до архітектурної керованості, підтримованості й масштабованості.

Ключові слова: *Feature-Sliced Design, Angular, frontend-архітектура, масштабованість, шарова структура, модульність, ToDo-додаток, архітектурна методологія, технічний борг, UX-композиція, структура проекту.*

Problem Statement. In the context of increasing complexity in modern frontend systems and growing demands for scalability, modularity, and maintainability of source code, there is an urgent need for architectural approaches that support the effective organization of web application structure. One such approach is Feature-Sliced Design (FSD), an architectural methodology that involves dividing a software product into hierarchical layers (*app, processes, pages, widgets, features, entities, shared*) with clearly defined responsibilities. Despite the rising popularity of this approach within open-source communities, its practical implementation in commercial environments raises several concerns related to initial implementation costs, scalability challenges, onboarding of new team members, and alignment with business requirements. The problem addressed in this study lies in the critical analysis of the strengths and weaknesses of the FSD methodology when applied in real-world frontend development projects. The aim of the study is to identify the architectural, technical, and organizational conditions under which the use of FSD ensures maximum development efficiency and to determine the limitations that may arise in a commercial setting.

Analysis of Recent Studies and Publications. In the current environment of rapid frontend technology development, there is increasing interest in optimizing architectural solutions, particularly those that are modular and scalable, enabling effective management of complexity in large-scale projects. A significant position in this context is occupied by the Feature-Sliced Design (FSD) architectural methodology, which proposes structural separation of code across domain, functional, and interface levels. The relevance of this topic is confirmed by the emergence of studies analyzing both its potential and limitations in practical application. One of the conceptually closest approaches is HOFA, presented in the monograph by H. Ben Khalfallah. This approach emphasizes the creation of clean architecture in JavaScript and React applications, where maintaining isolation, separation of responsibilities, and modular organization is critically important for code maintainability [2]. The HOFA methodology closely corresponds to the principles of FSD, particularly in terms of component autonomy and the structuring of business functionality. Similarly, the study by M. Kolomoyets and Y. Kynash focuses directly on the architectural design of frontend applications. It emphasizes the importance of maintaining a clear hierarchy in building web interfaces, which largely aligns with the core logic of FSD. The authors highlight the need for implementing structural decomposition that enables functionality to scale without compromising the integrity of the system [8]. An interesting interdisciplinary example of structured architecture is demonstrated by the team of X. Kong et al., who developed the RFSD-YOLO model for detecting prohibited items in X-ray images. Although this system belongs to the field of computer vision, its name (Refined Feature Sliced Design) and structural principles reflect the adaptation of modularity and independent component processing concepts to machine learning tasks [9].

The study by N. Rašović in the field of additive manufacturing also demonstrates the application of multi-attribute analysis methods and architectural structuring for technical decision-making. Although not directly related to web development, the approach to parameter formalization and modular task structuring indicates a broader trend of transferring architectural thinking into adjacent technical domains [10].

S. Gao et al., in their article, describe the architecture of an inspection neural network with a dynamic feature extraction module and a task alignment mechanism. The structure of this model reflects an intention to isolate subtasks and build efficient, predominantly independent modules, which fully aligns with the compositional principles of the FSD approach [5].

Therefore, recent publications indicate a growing demand for architectural methodologies characterized by clear decomposition and control over inter-component dependencies. The Feature-Sliced Design methodology

is increasingly viewed as a promising standard for structuring frontend applications; however, its practical implementation in commercial environments requires further investigation.

Formulation of the Research Objective. The primary objective of this study is to comprehensively examine the architectural methodology known as Feature-Sliced Design (FSD) within the context of commercial frontend development, with a focus on its practical feasibility, advantages, and structural limitations. The central emphasis is placed on evaluating the potential of FSD to enhance scalability, modularity, and maintainability of web applications, taking into account the specifics of business logic, team management, and the project lifecycle.

The technological goals of the study are as follows: to analyze the architectural principles underlying FSD and compare them with traditional models (MVC, MVVM, Atomic Design); to model the structural framework of a frontend application using FSD stratification (app / processes / pages / widgets / features / entities / shared); to implement a demonstration project (ToDo Application) in Angular, incorporating the core principles of FSD to assess the effectiveness of the approach; to verify technical performance, structural flexibility, and scalability potential using business logic scenarios.

The practical goals of the study are as follows: to identify the conditions under which the application of Feature-Sliced Design is optimal in commercial environments; to assess the impact of the FSD methodology on reducing technical debt, accelerating the onboarding of new developers, and improving team collaboration quality; to develop recommendations for implementing FSD in small and medium-sized IT teams.

The expected technical outcomes of the study are as follows: structured documentation and a visual model of the frontend application architecture built using the FSD methodology; a functional prototype of an Angular application implemented in accordance with FSD stratification; an analytical conclusion regarding the advantages and limitations of FSD compared to alternative approaches.

The study holds significant importance for the advancement of architectural design practices in the field of frontend development. It enables the evaluation not only of the effectiveness of FSD as a methodology but also of its relevance to real-world business environments, making the findings valuable for teams seeking structural clarity, flexibility, and scalability of software interfaces. The innovative aspect of this work lies in the practical application of FSD principles within the Angular environment, followed by an analysis of its performance in a business context.

Presentation of the main research material. The Feature-Sliced Design (FSD) architectural methodology, which is formed at the intersection of structural engineering and a domain-oriented approach to frontend development, is gradually gaining popularity among teams working with Angular, React, and other modern frameworks. The main idea is to divide an application not by technical characteristics (e.g., components, services, or templates), but by functional scenarios and business logic, which allows you to maintain the scalability and manageability of the project [4].

The methodology is based on the principle of a multi-level hierarchy: the application is structured into layers (app, processes, pages, widgets, features, entities, shared), each of which performs its own autonomous role. Lower layers are unaware of the existence of higher layers – for example, shared components have no idea how they are used at the page or feature level. This approach ensures one-way encapsulation of dependencies, i.e., direct links are formed only down the structure, which reduces the number of non-obvious links and facilitates maintenance [3]. Technically, this means that a lower layer (e.g., shared) does not have access to layers above or next to it. Higher layers (e.g., entities, features, pages) can use everything below them, but not vice versa – similar to rivers that flow into the sea but do not flow in the opposite direction. This logic of relationships between FSD layers is reflected in (Tab. 1).

Table 1

Rules for interaction between layers in Feature-Sliced Design architecture

Layer	Can be used	Can be used
app	shared, entities, features, widgets, pages, processes	–
processes	shared, entities, features, widgets, pages	app
pages	shared, entities, features, widgets	processes, app
widgets	shared, entities, features	pages, processes, app
features	shared, entities	widgets, pages, processes, app
entities	shared	features, widgets, pages, processes, app
shared	–	entities, features, widgets, pages, processes, app

The methodology pays special attention to the so-called “removability” of modules: each functional block must be implemented in such a way that it can be completely removed without compromising the integrity of

the system. This is achieved through high component cohesion and minimal interdependence between them. By analogy with a water system, each fragment (slice) is an independent “river” that flows into the common “ocean” of the application logic, but does not intersect with other direct flows [4].

As part of the research, a prototype ToDo application was implemented on Angular, built according to FSD principles. The corresponding structure looked as follows:

```
bash
src/
├── app/           # глобальні налаштування маршрутизації та сховища
├── pages/         # окремі сторінки (tasks-list, task-details, not-found)
├── features/      # сценарії користувача (tasks-filter, toggle-task)
├── entities/     # бізнес-сутності (task-card, task-row)
└── shared/       # UI-бібліотеки, API, спільні сервіси
```

Each page contains only the composition of components, while the interaction logic is moved to separate “features.” For example, interaction with the task list is implemented as a separate function tasks-filter, and the change in execution status is implemented in toggle-task. The essence of a task is formed as a data structure with its own visual components, implemented as task-card and task-row [3]. Figure 1 shows how the application structure is built according to the FSD methodology: from pages to shared components. Each layer is clearly separated, and the arrows illustrate the permissible direction of interaction between layers.



Fig. 1. Layer hierarchy in Feature-Sliced Design based on the ToDo application example

It is important to note that the implementation of the ToDo application provided an opportunity to empirically validate the key advantages of FSD. First, the project's structural organization facilitated the distribution of roles within the development team. Second, during the testing phase, a reduction in git merge conflicts was recorded, attributed to the separation of responsibilities across distinct project segments. Third, it became possible to easily remove or replace any functional block without disrupting the overall application logic, which is a critical factor for projects with frequent releases and evolving requirements [6, 7]. Thus, the application of Feature-Sliced Design in commercial frontend projects enhances the manageability of software architecture, although it requires the team to have a clear understanding of the paradigm and to maintain strict discipline in adhering to structural rules. While the methodology is not a universal solution, its benefits become particularly evident in medium- to large-scale development environments, where the cost of architectural debt is especially high.

Conclusions. The conducted study provided a theoretical and applied analysis of the Feature-Sliced Design (FSD) architectural methodology in the context of its application to commercial frontend projects. Based on the client's materials, the architecture of a ToDo application developed in Angular was reconstructed in accordance with FSD principles, which enabled the demonstration of real implementation mechanisms.

The main focus of the study was on the structural organization of the project through the stratification of logic: app, processes, pages, widgets, features, entities, and shared. Each layer was assigned specific functional responsibilities, permitted interdependencies (according to the rules of unidirectional encapsulation), and a defined role in ensuring modularity. Visual diagrams and integrated tables illustrate the direction of dependency flows and the logic underlying the construction of the interface layer of the product.

The research concluded that FSD ensures high structural flexibility, simplifies the distribution of responsibilities among developers, and facilitates the refactoring of individual code segments without compromising architectural integrity. Additional advantages of the methodology include reduced technical debt and logical isolation of functional modules. At the same time, the study identified several limitations, including the need for preliminary architectural planning, standardization of approaches within the team, and appropriate developer competencies.

The innovative contribution of this work lies in the analytical adaptation of the modern FSD methodology to a concrete application example, making the results applicable as a conceptual model for implementation in small and medium-sized projects. The resulting structure demonstrates potential for scalability and flexible development without compromising manageability.

Future research directions include the formalization of criteria for evaluating architectural efficiency in comparison with alternative approaches (such as MVVM or Atomic Design), as well as the extension of FSD methodology for use in multifunctional applications with complex business logic.

Bibliography:

1. Ben Khalfallah H. CRISP: Clean, Reliable, Integrated Software Process. *Crafting Clean Code with JavaScript and React*. Berkeley, CA : Apress, 2024. URL: https://doi.org/10.1007/979-8-8688-1004-6_6. (date of access: 08.06.2025).
2. Ben Khalfallah H. HOFA: The Path Toward Clean Architecture. *Crafting Clean Code with JavaScript and React: A Practical Guide to Sustainable Front-End Development*. Berkeley, CA : Apress, 2024. P. 233–287. URL: https://doi.org/10.1007/979-8-8688-1004-6_4. (date of access: 08.06.2025).
3. Feature-Sliced Design – modern Front End Architectural Methodology on Angular. *Medium*. URL: <https://medium.com/@fed4wet/feature-sliced-design-modern-architectural-methodology-on-angular-d0ef705ef598>. (date of access: 08.06.2025).
4. Feature-Sliced Design (FSD). URL: <https://medium.com/@jstify.community/feature-slice-design-%D1%89%D0%BE-%D1%86%D0%B5-%D1%82%D0%B0%D0%BA%D0%B5-8cb86d059c16>. (date of access: 08.06.2025).
5. Gao S., Xia T., Hong G., Zhu Y., Chen Z., Pan E., Xi L. An inspection network with dynamic feature extractor and task alignment head for steel surface defect. *Measurement* 2024. Vol. 224. P. 113957. URL: <https://doi.org/10.1016/j.measurement.2023.113957>. (date of access: 08.06.2025).
6. Goh H. A., Ho C. K., Abas F. S. Front-end deep learning web apps development and deployment: a review. *Applied Intelligence* 2023. Vol. 53, No. 12. P. 15923–15945. URL: <https://doi.org/10.1007/s10489-022-04278-6>. (date of access: 08.06.2025).
7. Hidayat D. C., Atmaja I. K. J., Sarasvananda I. B. G. Analysis and Comparison of Micro Frontend and Monolithic Architecture for Web Applications. *Jurnal Galaksi* 2024. Vol. 1, No. 2. P. 92–100. URL: <https://doi.org/10.70103/galaksi.v1i2.19>. (date of access: 08.06.2025).
8. Kolomoiets M., Kynash Y. Front-End web development project architecture design. *2023 IEEE 18th International Conference on Computer Science and Information Technologies (CSIT)*. 2023. P. 1–5. URL: <https://doi.org/10.1109/CSIT61576.2023.10324238>. (date of access: 08.06.2025).
9. Kong X., Li A., Li W., Li Z., Zhang Y. RFSD-YOLO: An Enhanced X-Ray Object Detection Model for Prohibited Items. *2024 IEEE International Conference on Systems, Man, and Cybernetics (SMC)*. October 2024. P. 4412–4418. URL: <https://doi.org/10.1109/SMC54092.2024.10831942>. (date of access: 08.06.2025).
10. Rašović N. Recommended layer thickness to the powder-based additive manufacturing using multi-attribute decision support. *International Journal of Computer Integrated Manufacturing* 2021. Vol. 34, No. 5. P. 455–469. URL: <https://doi.org/10.1080/0951192X.2021.1891574>. (date of access: 08.06.2025).

Дата надходження статті: 25.06.2025

Дата прийняття статті: 30.06.2025

Опубліковано: 23.09.2025

UDC 004.652.4:004.738.5

DOI <https://doi.org/10.32689/maup.it.2025.2.27>

Khrystyna TERLETSKA

Bachelor's Degree in "Documentation and Information Services",

Senior Software Engineer at Datadog,

Graduate of Lviv Polytechnic National University, khrystynaterletska1@gmail.com

ORCID: 0009-0002-7302-2625

ENHANCING REAL-TIME DATA REPLICATION EFFICIENCY THROUGH OPTIMIZATION OF CHANGE DATA CAPTURE METHODS

Abstract. Relevance of the research is driven by the growing need for efficient real-time data replication in distributed and hybrid environments, where system scalability, consistency, and low latency are critical. Traditional change propagation mechanisms are increasingly inadequate under high-frequency workloads, highlighting the need for a re-evaluation and modernization of Change Data Capture (CDC) methods.

The aim of the article is to provide a scientific rationale and develop approaches to improving the efficiency of real-time data replication through the optimization of Change Data Capture methods, taking into account the architectural and operational specifics of modern distributed computing environments.

The research methodology is based on systems analysis of CDC implementations in cloud and hybrid infrastructures, architectural modeling of replication flows, comparative evaluation techniques, typological classification of functional characteristics, and a criterion-based approach to assessing efficiency in various data update scenarios.

The research results include the identification of key functional features of CDC mechanisms, the classification of architectural replication models, and the substantiation of performance criteria such as latency, throughput, consistency, scalability, connectivity flexibility, and resource efficiency. Key technical constraints were identified in high-change-rate environments, particularly performance instability, access conflicts, and compatibility issues with traditional database systems and streaming platforms. It was demonstrated that event-driven architectures with asynchronous processing and adaptive buffering yield better performance when properly tuned.

In the conclusions, the relevance of hybrid CDC models is substantiated, the dependency between replication efficiency and architectural processing models is confirmed, and common implementation barriers in conventional database and stream-processing ecosystems are outlined.

The research perspectives include the development of intelligent CDC strategies, dynamic change flow orchestration, and the improvement of cross-platform interoperability in multi-cloud infrastructures.

Key words: asynchronous processing, streaming architectures, transaction consistency, change buffering, in-memory analytics.

Христина ТЕРЛЕЦЬКА. ПІДВИЩЕННЯ ЕФЕКТИВНОСТІ РЕПЛІКАЦІЇ ДАНИХ У РЕЖИМІ РЕАЛЬНОГО ЧАСУ ЧЕРЕЗ ОПТИМІЗАЦІЮ МЕТОДІВ CHANGE DATA CAPTURE

Анотація. Актуальність дослідження зумовлено необхідністю підвищення ефективності реплікації даних у режимі реального часу в умовах стрімкого зростання обсягів інформації, розширення розподілених інфраструктур і динамічного навантаження на системи обробки. Традиційні механізми передачі змін дедалі частіше виявляються недостатньо гнучкими або ресурсно затратними, особливо у високонавантажених середовищах, що актуалізує завдання переосмислення методів Change Data Capture (CDC).

Метою статті є наукове обґрунтування та розробка підходів до підвищення ефективності реплікації даних у режимі реального часу шляхом оптимізації методів Change Data Capture з урахуванням специфіки сучасних розподілених обчислювальних середовищ.

Методологія дослідження базується на системному аналізі CDC-процедур у хмарних та гібридних інфраструктурах, моделюванні архітектурних схем реплікації, застосуванні методів порівняльного оцінювання, типологізації функціональних характеристик, а також критеріального підходу до визначення ефективності процедур у різних сценаріях оновлення даних.

Результати дослідження відображаються у визначенні ключових функціональних ознак CDC-реалізацій, класифікації архітектурних моделей реплікації, обґрунтуванні релевантних критеріїв оцінки (затримка, узгодженість, масштабованість, пропускна здатність, гнучкість підключення), а також у виявленні основних технологічних обмежень у високонавантажених середовищах. Встановлено, що подієво-орієнтовані архітектури з асинхронною обробкою змін демонструють кращу продуктивність за умов правильної буферизації і контролю черг подій.

У висновках обґрунтовано доцільність гібридного підходу до CDC у складних архітектурах, підтверджено залежність ефективності від поєднання архітектурної моделі та режиму обробки, а також окреслено типові технічні обмеження, що виникають у процесі впровадження CDC у традиційні СКБД і потокові сервіси.

Перспективами подальших досліджень визначено розробку інтелектуальних CDC-стратегій, адаптивного керування потоками змін і розширення міжплатформної сумісності компонентів у мультихмарних середовищах.

Ключові слова: асинхронна обробка, потокові архітектури, узгодженість транзакцій, буферизація змін, аналітика в оперативній пам'яті.

© K. Terletska, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

Problem statement. In today's environment of rapid growth in data volumes and demands for fast information processing, ensuring effective real-time data replication is becoming particularly important. Multi-organizational infrastructures, distributed computing systems, and cloud services require constant synchronous exchange of updates between sources and analytical platforms, which necessitates high-performance methods for detecting changes in data. Change Data Capture (CDC) technology is a key tool for detecting, capturing, and transferring changed data without the need to scan the entire source. However, in practice, its effectiveness is limited by a number of technical and logistical factors, including unstable performance when processing large transaction volumes, lack of universal compatibility with different types of data sources, difficulties in reconciling changes in distributed environments, and transmission channel overload. These challenges are particularly acute in systems where the delay between the occurrence of a change and its reflection in the target environment is critical, such as in financial, medical, or logistics platforms. In view of this, the scientific task is to find models that can adaptively respond to load changes, minimize duplicate operations, optimize resource consumption, and ensure transaction consistency. The practical significance of the research lies in the creation of CDC mechanisms that can be integrated into existing infrastructures without loss of compatibility, while complying with requirements for scalability, continuity, and transactional integrity, which has a direct impact on the stability and efficiency of real-time information systems.

Analysis of recent studies and publications. Analysis of current research on improving the efficiency of real-time data replication through optimization of Change Data Capture (CDC) methods allows us to identify four main areas of focus.

The first area focuses on the development of theoretical and applied concepts of CDC as a key replication tool. In their work, L. Hao, T. Jiang, Y. Lin, and Y. Lu [8] propose a classification of approaches to solving the CDC problem from the perspective of change source types—transaction logs, triggers, and timestamps. D. Seenivasan and M. Vaithianathan [14] analyze in detail the adaptation of CDC to the architecture of modern real-time computer systems, emphasizing the balance between the frequency of change capture and the load on the system. F. M. Imani, Y. D. L. Widyasari, and S. P. Arifin [12] consider CDC as a means of optimizing ETL processes, in particular to reduce data duplication and speed up analytical processing. T. Zheng, G. Chen, X. Wang, C. Chen, X. Wang, and S. Luo [20] explore CDC in the context of intelligent real-time big data processing, demonstrating the effectiveness of a coordinated environment for interaction between data sources, processing platforms, and applications. H. Chandra [4] conducted an experimental comparison of the effectiveness of CDC implementation in different types of data structures, allowing CDC to be adapted to the specifics of storage systems. Further research into adaptive algorithms for constructing change dependency graphs to optimize the CDC process in large distributed systems is promising.

The second area focuses on the security aspects of data replication and ensuring protection when implementing CDC on the front end. In his publication, Y. Horbenko [11] develops the concept of a secure automated framework with built-in encryption on the client side and implementation of the Zero Trust approach to API, which minimizes the risks of confidentiality loss during replication. In a subsequent study, Y. Horbenko [10] shows how the use of confidential computing technologies, in particular secure computing enclaves and homomorphic encryption, increases the security of front-end components during the exchange process. A promising direction in this area is the development of CDC processes with end-to-end encryption and anonymization at the source level, which complicates unauthorized analysis of change streams.

The third area covers architectural and infrastructure solutions for implementing CDC in cloud and distributed environments. R. B. Yurinec and I. B. Pirko [1] proposed innovative methods for integrating CDC into data warehouse filling procedures, emphasizing adaptability to changes in sources. D. Leschert, K. Sommerstad, J. Janzen, and C. Wester [3] analyze the difficulties of configuring CDC in industrial distributed environments, where configuration accuracy critically affects replication quality. B. Vagadia [17] considers CDC in the broader context of digital disintermediation and business process transformation, where replication performs the function of real-time updating of digital copies. In W. Q. Yan's monograph [18], CDC is described as a basic element of streaming data collection in intelligent video surveillance systems, which requires high performance and fault tolerance. A. Zakiah, R. Yusuf, and A. S. Prihatmanto [19] analyze the role of CDC in ensuring high data availability in the digital age at the conference. The development of CDC as a microservice with dynamic routing of changes in a multi-cloud environment with limited resources is promising.

The fourth area is related to the support of CDC in document management, analytics, and data governance within complex digital environments. The study by P. Dhakal, M. Munikar, and B. Dahal [6] presents an approach to automated data extraction from documents using a template-matching method, which can be integrated into CDC processes. The research conducted by P. A. Harris, R. Taylor, B. L. Minor, V. Elliott, M. Fernandez, L. O'Neal, L. McLeod, G. Delacqua, F. Delacqua, J. Kirby, and S. N. Duda [9] describes the work of the REDCap consortium, where CDC ensures the synchronization of changes in medical records across platforms. A promising direction

is the integration of CDC with semantic data tagging mechanisms for enhanced contextual analysis in health-care and public administration systems.

Despite significant progress in the field of incremental data loading, a number of aspects of the overall problem remain unresolved. In particular, the specifics of working with streaming sources in a cloud environment, where the dynamics of changes, uneven data flow, and parallel access conflicts create significant difficulties for stable processing, have not been sufficiently researched. Existing approaches do not fully address the need for transactionality, consistency, and version control in distributed analytics systems. There is also a lack of uniform criteria for evaluating the effectiveness of incremental update procedures in terms of performance, scalability, and data integrity. Much of the research relies on limited experimental scenarios or demonstrates applied implementations without proper generalized analytics.

The proposed study aims to partially fill these gaps by systematically analyzing theoretical and applied approaches to incremental loading, identifying key limitations of their application in cloud environments, and substantiating criteria for the effectiveness of adaptive change processing. Particular attention is paid to the potential of using platforms such as Delta Lake from the perspective of transactional updates, version control, and integration with streaming mechanisms. This allows us to expand the existing understanding of the functional capabilities of CDC in complex infrastructures and form a basis for further applied research and project solutions.

The **purpose of this article** is to provide scientific justification and develop approaches to improving the efficiency of real-time data replication by optimizing Change Data Capture methods, taking into account the specifics of modern distributed computing environments.

To achieve this goal, the following tasks are to be performed:

1. Analyze technical approaches and architectural models for implementing CDC in cloud and hybrid environments in order to assess their impact on latency and computing costs.
2. Justify the criteria for the effectiveness of CDC procedures, taking into account source types, processing modes, consistency, and scalability.
3. Identify problems with the application of CDC in high-load systems and formulate recommendations for their elimination based on adaptive buffering, asynchronous logging, and in-memory processing.

Presentation of the main material. CDC is the basis of modern synchronization processes in real-time systems. Its key purpose is to ensure fast and efficient transfer of only changed records from data sources to analytical platforms or target repositories without the need for a full table scan. This approach reduces latency, optimizes the use of network and computing resources, and ensures data relevance in heterogeneous environments. The development of cloud computing, streaming analytics, and microservice architecture has led to the emergence of various technical implementations of CDC, which differ in data sources, integration mechanisms, and configuration flexibility (Tab. 1).

Table 1

Technical Approaches to CDC Implementation and Their Functional Characteristics

Technical Approach	Description of Operation Mechanism	Functional Features
Triggers in Database Management Systems (DBMSs)	Utilization of SQL triggers that respond to insert, update, or delete operations within tables	Rapid implementation in traditional DBMSs; ensures transactional consistency
Transaction Log Analysis	Captures changes directly from the DBMS transaction log (e.g., WAL in PostgreSQL)	Minimal impact on the data source; high accuracy of changes; preserves temporal sequence
Snapshot-Based CDC	Compares the current table state with a previous version using control values	DBMS-independent; suitable for legacy systems
Event Stream Processing	Detects changes through streaming platforms (e.g., Apache Kafka, Apache Pulsar, Debezium)	Supports scalable real-time change processing; integrates with analytical pipelines
Platform-Oriented CDC Services	Employs managed services (e.g., Google BigQuery Data Transfer, Snowflake Streams, Azure CDC)	Enables automation; integrates with cloud-based analytics services; reduces operational workload

Source: compiled by the author based on [3; 4; 5; 7; 13]

In real-world conditions, preference is usually given to combinations of approaches, in particular the combination of transaction log analysis with event processing using Apache Kafka or Debezium, which allows large volumes of data to be processed in real time without overloading the sources. This solution provides low latency, scalability, and asynchronous data delivery to multiple consumers [5]. In cloud environments, where the “data as a service” model prevails, platform solutions such as Google Cloud BigQuery or Snowflake are

increasingly being used, in which CDC functions are implemented at the infrastructure level as managed services, which greatly simplifies the deployment of complex ETL (Extract, Transform, Load) processes and ensures resilience to peak loads [13]. In distributed architectures and microservice approaches, event-driven CDC is the most promising for building adaptive synchronization systems capable of responding to data changes with minimal processing costs [16].

In modern information systems, data replication using CDC technology has become an integral part of ensuring consistency, continuity of processing, and scalability of computing processes. This is especially true in cloud and hybrid environments, where data is located in different regions, different types of storage and services are used, and the load on systems changes dynamically. CDC minimizes the amount of data transferred by focusing only on changed records, which is especially important when processing large streams of information in real time. However, the effectiveness of such replication largely depends on the architectural model of its implementation, which must take into account not only the computing capabilities of the system, but also the requirements for latency, consistency, and infrastructure costs (Tab. 2).

Table 2

Architectural Models for Data Replication Using CDC in Cloud and Hybrid Environments

Architectural Model	Description of Structure and Component Interaction	Potential for Latency and Cost Optimization
Push-Based Replication Model	The data source independently initiates the transmission of changes to the replica or handler	Suitable for low-latency scenarios but requires a stable connection; limited flow control
Pull Model with Periodic Polling	A CDC agent or replica periodically queries the source for updates	Lower processing overhead but higher latency, especially with large data volumes
Event-Driven Model with Event Broker	Changes are written to a log (e.g., Kafka), from which they are delivered to consumers	Highly scalable; ensures low latency; enables real-time event processing
CDC via Embedded Cloud Services	Replication is managed through cloud-based platforms (e.g., Google Dataflow, Azure Synapse, Snowpipe)	Minimal maintenance effort; optimized for cloud provider infrastructure

Source: compiled by the author based on [2; 6; 7; 17; 19; 20]

In practice, event-driven models are increasingly favored, wherein changes are captured in real time using event brokers such as Apache Kafka or Google Pub/Sub and then distributed to consumers or processors through streaming systems like Apache Flink or Google Dataflow [7]. This approach significantly reduces latency under high volumes of change while ensuring scalability and component isolation. Cloud-managed services, such as Snowpipe in Snowflake or CDC pipelines in Microsoft Azure Synapse, automate event processing by dynamically allocating computing resources based on actual workloads. This minimizes administrative overhead and helps prevent system downtime [15]. In complex hybrid architectures, hybrid models are being adopted more frequently. These allow for differentiated change capture frequencies based on data type or priority. For instance, critical transactions are recorded using event-driven methods, while secondary analytical data is updated in batches [2]. Consequently, the architecture of CDC-based replication must remain flexible and adaptive to meet contemporary demands for availability, processing speed, and infrastructure cost efficiency.

The effectiveness of implementing Change Data Capture (CDC) procedures in complex information environments depends not only on the choice of technology or tool, but primarily on the alignment of CDC mechanisms with data source types, processing modes, and specific update scenarios. There is no universal approach to CDC implementation, as different systems (transactional databases, analytical warehouses, streaming sources, or non-relational NoSQL systems) have distinct requirements regarding latency, consistency, and availability. In real-time contexts, it is particularly important to ensure consistency guarantees, adapt to variable workloads, and maintain scalability without compromising data integrity. Therefore, establishing clear criteria for evaluating the efficiency of CDC procedures is a necessary prerequisite for their optimal use in modern heterogeneous architectures.

In practical settings, the implementation of CDC procedures often involves finding a compromise between latency and consistency, depending on the characteristics of the data source. For instance, financial transactional systems prioritize strict transactional consistency, where minor delays are acceptable, but any loss or inconsistency of changes is intolerable. In streaming architectures with high-frequency updates, the priority is minimizing latency, even at the cost of temporary incompleteness, which is later resolved by event processors [11]. In cloud environments, scalability and resource efficiency are of primary importance. Platforms such as Snowflake employ adaptive CDC mechanisms with integrated services like Snowpipe Streaming, which provide

low latency and automatic load balancing under distributed conditions. This approach enables the integration of heterogeneous data sources and optimizes replication channels without requiring modifications to the architecture of the source system.

Table 3

**Efficiency Criteria for CDC Implementation in the Context of Data Source Types,
Processing Modes, and Update Scenarios**

Efficiency Criterion	Description	Dependence on Source or Scenario
Transmission latency	The time between a change in the source and its appearance in the target system	Critical for streaming data and OLTP systems (Online Transaction Processing)
Throughput	The number of changes the system can process per unit of time	Important for replicating large data batches, archives, and batch updates
Change consistency	Guarantee of the correct order and completeness of changes in the target replica	Essential for financial and healthcare systems where incomplete replication is unacceptable
Scalability	The ability of the CDC architecture to support increasing data volume and number of consumers	Crucial for cloud environments and multi-component microservice architectures
Connectivity flexibility	The ability to integrate CDC with diverse sources including relational, NoSQL, and API-based systems	Important for hybrid architectures, integration platforms, and ETL pipelines
Resource utilization efficiency	The ratio between processing costs and replication speed	A key parameter in systems with limited computing resources

Source: compiled by the author based on [1; 4; 6; 8; 10; 12; 20]

In systems with high-frequency data changes, the implementation of Change Data Capture (CDC) procedures encounters a range of systemic challenges that significantly affect the stability, reliability, and efficiency of real-time data processing. One of the primary issues is performance instability caused by the overloading of change transmission channels, especially during peak loads or when changes arrive unevenly. In high-transaction environments, CDC systems based on triggers or periodic scanning can place considerable strain on the source database, reducing overall throughput and increasing response times [16]. Even when using event-driven architectures with message brokers, there remains a risk of exceeding queue limits, which may result in delays or data loss under constrained computational resources [7].

Another critical issue involves conflicts arising from concurrent data access, which occur when multiple subsystems simultaneously read, process, and update records. In such cases, the integrity of the transactional context may be compromised, particularly in systems lacking centralized control over change streams. This becomes especially problematic in distributed databases or environments with microservice architectures, where CDC data consumers operate independently. Such independence can lead to state desynchronization or duplicated processing [8]. Even when transaction logs are used as the primary source of change data, it is not always possible to guarantee the correct order of events, especially in the presence of network delays or node failures [20]. An additional challenge lies in the compatibility limitations of CDC solutions with traditional database management systems (DBMS) and streaming platforms. Some legacy DBMSs, particularly those that do not support external access to transaction logs or rely on outdated locking mechanisms, significantly complicate the implementation of CDC without substantial infrastructure modification [5]. At the same time, many streaming platforms, such as Apache Flink or Amazon Kinesis, impose specific requirements on event formats and delivery order that do not always align with the capabilities of the data source. This creates the need for supplementary adapters, converters, or buffering layers, which increase system complexity, introduce potential points of failure, and reduce overall system stability [18].

Enhancing Change Data Capture (CDC) procedures in environments characterized by high-frequency changes and variable workloads requires the implementation of adaptive technological solutions capable of maintaining system resilience during peak loads, improving data consistency, and optimizing the use of computational resources. One of the key areas of improvement involves the adoption of adaptive change buffering mechanisms, which enable the dynamic adjustment of the volume of events accumulated prior to processing based on the current system state and available resources. Such buffering can be implemented either at the data source (for instance, in the form of commit logs with timestamps) or within the intermediate event broker layer, where intelligent queue management with prioritization is applied.

Another important component of CDC optimization is the use of asynchronous change logging. Offloading change capture to a separate process helps prevent blocking the main transaction processing flow, enables

independent scaling of processing components, and reduces latency. Asynchronous change handling, when combined with timestamps and transaction identification mechanisms, ensures the orderly delivery of changes to replicated target systems or analytical services. This approach is particularly effective in systems with a high volume of parallel transactions, where synchronous CDC logging can negatively impact overall performance.

The third direction involves integrating CDC mechanisms with in-memory analytics, allowing modified data to be processed directly in main memory prior to final storage. This approach enables near-real-time analytics and forecasting without the need to wait for the completion of the full ETL cycle. In cloud environments, such integration is typically implemented through the use of in-memory computational clusters, such as Apache Ignite or Google BigQuery BI Engine, where modified data is delivered from CDC channels immediately following event processing. This makes it possible not only to reduce system response latency to data changes, but also to enable adaptive scaling of computational resources based on the context of the analytical query.

Collectively, these approaches provide a multi-level improvement in the efficiency of CDC procedures by reducing the impact of workload on performance, enhancing transactional consistency, minimizing analytical processing latency, and improving adaptability to the current demands of real-time distributed computing.

Conclusions. This study substantiates the role of Change Data Capture (CDC) technology as a key component of effective real-time data replication within modern cloud and hybrid architectures. It has been established that the efficiency of CDC implementation depends not only on the technical solution itself but also on its alignment with the type of data source, the nature of data changes, and the requirements for consistency, scalability, and latency. The analysis identified the most relevant architectural models for CDC deployment, described the criteria for their effectiveness, and examined technical constraints that arise in systems with high-frequency data changes. The main challenges include performance instability, conflicts during concurrent data access, and compatibility issues with certain database management systems and streaming platforms. The study provides a scientific rationale for the adoption of adaptive buffering mechanisms, asynchronous change logging, and integration with in-memory analytics as strategic directions for enhancing the reliability and performance of CDC infrastructures. Future research should focus on intelligent strategies for the dynamic optimization of CDC processes based on workload context and business process characteristics.

Bibliography:

1. Юринєць Р. Б., Пірко І. Б. Інноваційні методики інтегрування даних для оптимізації процесу наповнення сховища даних. *Науковий вісник НЛТУ України*. 2024. Вип. 34, № 6. С. 101–105. DOI: <https://doi.org/10.36930/40340614>.
2. Beyond Teradata Migration: Implement a Modern Data Warehouse in Azure Synapse Analytics. *Microsoft Azure: website*. 2023. URL: <https://learn.microsoft.com/en-us/azure/synapse-analytics/migration-guides/teradata/7-beyond-data-warehouse-migration> (date of access: 13.06.2025).
3. Capture Configuration Challenges. 2023 IEEE IAS Petroleum and Chemical Industry Technical Conference (PCIC), New Orleans, USA. IEEE, 2023. P. 195–203. DOI: 10.1109/PCIC43643.2023.10414316.
4. Chandra H. Experimental results on change data capture methods implementation in different data structures to support real-time data warehouse. *International Journal of Business Information Systems*. 2020. Vol. 34, No. 3. P. 373. DOI: <https://doi.org/10.1504/ijbis.2020.108651> (date of access: 13.06.2025).
5. Debezium Documentation. *Debezium: website*. 2024. URL: <https://debezium.io/documentation> (дата звернення: 06.06.2025).
6. Dhakal P., Munikar M., Dahal B. One-Shot Template Matching for Automatic Document Data Capture. 2019 Artificial Intelligence for Transforming Business and Society (AITB), Kathmandu, Nepal. IEEE, 2019. P. 1–6. DOI: 10.1109/AITB48515.2019.8947440.
7. Event-Driven Architecture. Confluent Developer: website. 2024. URL: <https://developer.confluent.io/courses/microservices/event-driven-architecture/> (date of access: 13.06.2025).
8. Hao L., Jiang T., Lin Y., Lu Y. Methods for Solving the Change Data Capture Problem. In: Xiong N., Li M., Li K., Xiao Z., Liao L., Wang L. (eds). *Advances in Natural Computation, Fuzzy Systems and Knowledge Discovery*. ICNC-FSKD 2022. *Lecture Notes on Data Engineering and Communications Technologies*. Cham: Springer, 2023. Vol. 153. P. 1025–1034. DOI: https://doi.org/10.1007/978-3-031-20738-9_87.
9. Harris P. A., Taylor R., Minor B. L., Elliott V., Fernandez M., O'Neal L., McLeod L., Delacqua G., Delacqua F., Kirby J., Duda S. N. The REDCap Consortium: Building an International Community of Software Platform Partners. *Journal of Biomedical Informatics*. 2019. Vol. 95. Article 103208. DOI: <https://doi.org/10.1016/j.jbi.2019.103208>.
10. Horbenko Y. Confidential Computing in Front-End: Enhancing Data Security with Secure Enclaves and Homomorphic Encryption. *International Journal of Advanced Multidisciplinary Research and Studies*. 2025. Vol. 5, No. 3. P. 308–321. URL: <https://www.multiresearchjournal.com/admin/uploads/archives/archive-1747130538.pdf> (date of access: 13.06.2025).
11. Horbenko Y. Secure Front-End Automation Framework: A Novel Approach to Client-Side Data Encryption and Zero Trust API Interaction. *Asian Journal of Research in Computer Science*. 2025. Vol. 18, No. 6. P. 177–193. DOI: <https://doi.org/10.9734/ajrcos/2025/v18i6690>.
12. Imani F. M., Widyasari Y. D. L., Arifin S. P. Optimizing Extract, Transform, and Load Process Using Change Data Capture. *6th International Seminar on Research of Information Technology and Intelligent Systems (ISRITI)*, Batam, Indonesia. IEEE, 2023. P. 266–269. DOI: 10.1109/ISRITI60336.2023.10468009.

13. Import Data into a Secured BigQuery Data Warehouse. *Google: website*. 2023. URL: <https://cloud.google.com/architecture/blueprints/confidential-data-warehouse-blueprint> (date of access: 13.06.2025).
14. Seenivasan D., Vaithianathan M. Real-Time Adaptation: Change Data Capture in Modern Computer Architecture. *ESP International Journal of Advancements in Computational Technology*. 2023. Vol. 1, No. 2. P. 49–61. URL: <https://www.researchgate.net/publication/381225264> (date of access: 13.06.2025).
15. Snowpipe Streaming – Snowflake Documentation. *Snowflake: website*. 2024. URL: <https://docs.snowflake.com/en/user-guide/data-load-snowpipe-streaming-overview>
16. Track Data Changes (CDC) in SQL Server. *Microsoft: website*. 2024. URL: <https://learn.microsoft.com/en-us/sql/relational-databases/track-changes/track-data-changes-sql-server> (дата звернення: 06.06.2025).
17. Vagadia B. Data Capture and Distribution. In: *Digital Disruption. Future of Business and Finance*. Cham: Springer, 2020. P. 45–66. DOI: https://doi.org/10.1007/978-3-030-54494-2_4.
18. Yan W. Q. Introduction to Intelligent Surveillance: Surveillance Data Capture, Transmission, and Analytics. Cham: Springer, 2019. 425 p. URL: <https://books.google.com.ua/books?id=pheJDwAAQBAJ> (date of access: 13.06.2025).
19. Zakiah A., Yusuf R., Prihatmanto A. S. The Benefits of Change Data Capture in Enhancing Data Availability in the Digital Transformation Era. *Widyatama International Conference on Engineering 2024 (WICOENG 2024)*. Atlantis Press, 2024. P. 302–308. DOI: https://doi.org/10.2991/978-94-6463-618-5_32.
20. Zheng T., Chen G., Wang X., Chen C., Wang X., Luo S. Real-time intelligent big data processing: technology, platform, and applications. *Science China Information Sciences*. 2019. Vol. 62. Article 82101. DOI: <https://doi.org/10.1007/s11432-018-9834-8>.

Дата надходження статті: 17.06.2025

Дата прийняття статті: 23.06.2025

Опубліковано: 23.09.2025

УДК 004.415.3:004.41

DOI <https://doi.org/10.32689/maup.it.2025.2.28>

Данило ТРОЩИЙ

аспірант кафедри інженерії програмного забезпечення,
Національний аерокосмічний університет «Харківський авіаційний інститут»,
d.troshchyi@khai.edu

ORCID: 0009-0008-3160-6060

Юлія КУЗНЕЦОВА

кандидат технічних наук, доцент кафедри інженерії програмного забезпечення,
Національний аерокосмічний університет «Харківський авіаційний інститут»,
y.kuznetsova@khai.edu

ORCID: 0000-0001-9416-6981

ОГЛЯД ТА АНАЛІЗ АСПЕКТІВ ПОВТОРЮВАНOSTІ РУТИННИХ ЗАВДАНЬ У ПРОЦЕСІ РОЗРОБЛЕННЯ МОБІЛЬНИХ ДОДАТКІВ

Анотація. Метою дослідження є системний аналіз рутинних повторюваних завдань, що виникають у процесі розроблення мобільних застосунків, а також визначення їхнього впливу на ефективність роботи розробницьких команд. Особлива увага приділяється виявленню підходів до зниження обсягу рутинних дій за рахунок автоматизації, стандартизації та повторного використання рішень.

Методологія. У роботі застосовано комплексний підхід, що поєднує аналіз наукових публікацій, огляд галузевих звітів, а також систематизацію типових рутинних завдань на різних етапах життєвого циклу мобільного застосунку. Проведено порівняльний аналіз існуючих технічних рішень: інструментів автоматизованого тестування, генерації користувацького інтерфейсу, застосування практик DevOps і інтеграції віртуальних помічників.

Наукова новизна. Запропоновано узагальнену класифікацію рутинних завдань мобільної розробки з урахуванням специфіки Android та iOS-платформ. Вперше комплексно проаналізовано прогалини у застосуванні сучасних інструментів автоматизації саме у мобільній сфері. Визначено, що більшість існуючих підходів не враховують особливості мобільних платформ або мають вузький технічний фокус.

Висновки. Результати дослідження підтверджують актуальність необхідності уніфікації та автоматизації повторюваних завдань для підвищення ефективності мобільної розробки. Запропоновано перспективні напрями подальших досліджень, зокрема розроблення спеціалізованих інструментів, шаблонів та практик, адаптованих до викликів мобільної інженерії, що дозволить скоротити витрати часу та ресурсів, зменшити кількість помилок і прискорити вихід продукту на ринок.

Ключові слова: мобільний застосунок, розробка додатків, рутина, повторюваність, програмне забезпечення, автоматизація, життєвий цикл програмного забезпечення.

Danylo TROSHCHYI, Yuliya KUZNETSOVA. REVIEW AND ANALYSIS OF REPETITIVE ROUTINE TASKS IN THE MOBILE APPLICATION DEVELOPMENT PROCESS

Abstract. The purpose of the study is to provide a systematic analysis of repetitive routine tasks in the development of mobile applications and to determine their impact on the efficiency of development teams. Special attention is paid to identifying approaches to reducing routine workload through automation, standardization, and reuse of solutions.

Methodology. The study applies a comprehensive approach that combines analysis of scientific publications, review of industry reports, and systematization of typical routine tasks at different stages of the mobile application lifecycle. A comparative analysis of existing technical solutions is conducted, including automated testing tools, user interface generation, DevOps practices, and the integration of virtual assistants.

Scientific novelty. A generalized classification of routine tasks in mobile development is proposed, taking into account the specifics of Android and iOS platforms. For the first time, existing gaps in the application of modern automation tools in the mobile domain are comprehensively analyzed. It is identified that most existing approaches either overlook platform-specific requirements or have a narrowly focused technical application.

Conclusions. The research results confirm the relevance of unifying and automating repetitive tasks to improve the efficiency of mobile development. Promising directions for further research are proposed, including the development of specialized tools, templates, and practices tailored to the challenges of mobile engineering, which will reduce time and resource consumption, minimize errors, and accelerate product delivery to the market.

Key words: mobile application, app development, routine tasks, repetition, software development, automation, software lifecycle.

Вступ. У сучасному світі мобільних технологій розробка додатків стає все більш популярною та необхідною. З кожним днем зростає попит на якісні мобільні рішення, що спонукає розробників працювати над новими проектами та вдосконалювати існуючі. Особливої популярності набувають mobile-first підходи, коли саме мобільна версія продукту стає основною точкою взаємодії з користувачем. Приклади таких продуктів – український застосунок Diia, який кардинально змінив сприйняття державних послуг у цифровому просторі, або BetterMe, що зумів вийти на глобальний ринок з рішеннями для фізичного та ментального здоров'я. Також варто згадати стартапи на кшталт Reface та Liki24, які демонструють, як мобільний застосунок може стати ядром продукту та драйвером бізнесу.

Успішність таких проектів залежить не лише від інноваційної ідеї, а й від ефективності реалізації – і саме тут виявляється важливість системного підходу до розробки. У практиці мобільного розробника значна частина часу витрачається на вирішення типових задач, які з проекту в проект мають подібну структуру й логіку. Йдеться про такі етапи, як створення архітектури проекту, підключення бібліотек, організація взаємодії з API, тестування, підтримка тощо.

Постановка проблеми. Незважаючи на те, що рутинні завдання часто сприймаються як другорядні, їх повторюваність і трудомісткість можуть істотно впливати на загальну продуктивність команди розробників. Кожен проект починається з набору типових дій, які вимагають витрат часу і уваги, але при цьому не додають унікальної цінності продукту. Така ситуація стимулює пошук способів автоматизації, стандартизації та оптимізації цих процесів.

В останні роки у сфері мобільної розробки з'явилися численні інструменти та підходи, що дозволяють зменшити навантаження на розробника, скоротити час виходу продукту на ринок і підвищити якість кінцевого продукту. Проте їх застосування часто фрагментарне, а відсутність системного підходу до повторюваних задач призводить до дублювання зусиль та ускладнень при масштабуванні командної роботи.

У цьому контексті постає потреба у більш глибокому аналізі повторюваних завдань, властивих мобільній розробці, з метою подальшого формування рекомендацій щодо їх ефективної організації, оптимізації або автоматизації. Зокрема, важливо враховувати специфіку життєвого циклу мобільного застосунку – від формування вимог до підтримки продукту в пострелізному періоді.

Аналіз досліджень і публікацій. Розробка мобільних застосунків охоплює цілий спектр технічних і організаційних процесів, що повторюються з проекту в проект. Такі задачі, попри свою технічну простоту, становлять основу щоденної роботи розробника і мають суттєвий вплив на загальну продуктивність команди. Незважаючи на свою буденність, ці задачі впливають на ефективність команди, якість коду, швидкість релізів і загальну стабільність продукту.

Рутинні завдання в мобільній розробці включають (але не обмежуються): налаштування середовища, підключення бібліотек, створення типових компонентів інтерфейсу, конфігурацію API, написання юніт- та UI-тестів, запуск тестів, аналіз логів, оновлення залежностей, складання збірки, запуск симуляторів або емуляторів, конфігурацію CI/CD, перевірку пул реквестів тощо. Часто саме ці дії вимагають повторення, не додаючи новизни, але залишаючись критично важливими для успішного запуску і підтримки застосунку.

Важливо відзначити, що характер рутинних задач значною мірою залежить від масштабу проекту, використовуваних технологій та особливостей команди розробників. Наприклад, у невеликих проектах частина цих завдань може бути мінімізована завдяки простоті структури, тоді як у масштабних корпоративних рішеннях складність і кількість рутинних процесів зростає в рази.

За деякими оцінками, на виконання рутинних задач мобільні розробники витрачають до 60% свого часу. Такий обсяг роботи часто недооцінюється, хоча саме від ефективності їх виконання залежить швидкість виведення продукту на ринок, якість коду і подальша підтримка [6, 7].

Аналіз публікацій у сфері програмної інженерії свідчить про наявність значної уваги до рутинних задач, особливо в контексті автоматизації, повторного використання коду, застосування шаблонів проектування тощо.

Однак переважна частина таких досліджень стосується загальної розробки програмного забезпечення і недостатньо фокусуються саме на особливостях мобільної розробки, де специфіка платформ, залежність від магазинів застосунків і взаємодія з мобільними пристроями часто залишається поза межами уваги дослідників.

На практиці команди створюють власні бібліотеки, використовують локальні шаблони, формалізують окремі етапи CI/CD, але без єдиної уніфікованої методології такі дії залишаються точковими й погано масштабуються на інші проекти або команди.

Таким чином, наявність повторюваних рутинних задач у мобільній розробці є системною проблемою, що напряду впливає на ефективність роботи, підтримку якості та швидкість оновлення продукту.

Її ігнорування – один із бар'єрів для росту та зрілості команд, тому подальше дослідження цієї теми є вкрай актуальним.

Упродовж останніх років дослідницький інтерес до оптимізації процесів розроблення програмного забезпечення постійно зростає, що зумовлено як технологічним прогресом, так і підвищеними вимогами до продуктивності команд розробників. Особливу увагу в цьому контексті привертають саме повторювані, рутинні дії, які супроводжують проєктування, реалізацію, тестування й обслуговування мобільних застосунків.

Відповідно до звіту GitLab DevSecOps Report 2024 [1], значна частина часу фахівців витрачається не на створення нового функціоналу, а на підтримку інфраструктури, усунення технічного боргу, рев'ю коду, налаштування середовища, оновлення бібліотек тощо. У тому ж звіті зазначено, що 27% розробників вважають збільшення рівня автоматизації ключовим чинником, який міг би суттєво підвищити їхню задоволеність роботою, що ще раз підкреслює важливість зменшення обсягу повторюваних ручних дій у щоденних процесах мобільної розробки.

Gitlab та JetBrains у свої дослідницьких роботах [1, 3] акцентують увагу на обмеженнях, що виникають внаслідок великого обсягу рутинної праці. Зокрема, дослідники підкреслюють, що повторюваність дій у мобільній розробці негативно впливає на стабільність, швидкість постачання оновлень і ментальне здоров'я інженерів. Проблема поглиблюється в умовах обмежених ресурсів, типової для невеликих стартапів або інді-команд.

В межах статті [10] автори досліджують вплив віртуальних персональних асистентів на ефективність виконання рутинних завдань. Зокрема, йдеться про використання Siri, Google Assistant та інших аналогів для автоматизації базових дій. Хоча підхід є цікавим, більшість кейсів орієнтовані на побутові або загальні офісні сценарії, а не на процеси, специфічні для мобільної розробки. Крім того, практичні аспекти інтеграції таких рішень у робоче середовище розробника залишаються недостатньо розкритими.

Робота [5] присвячена темі автоматизації DevOps-задач у процесі розробки програмного забезпечення. Автор окреслює базові принципи CI/CD, тестування та управління релізами, що безперечно стосується рутинних задач. Водночас, ця праця має досить загальний характер – більшість висновків стосуються лише традиційної бекенд-розробки та не враховуються особливості мобільних платформ. Крім того, приклади мають обмежену практичну цінність для тих, хто працює у сфері Android або iOS.

У дослідженні [9] розглядаються особливості впровадження CI/CD у мобільній розробці в умовах суворих регуляторних вимог. Автори аналізують виклики, пов'язані з дотриманням стандартів безпеки та конфіденційності, та пропонують стратегії для успішного впровадження CI/CD у таких середовищах. Однак, дослідження обмежується специфікою високорегульованих галузей, що може зменшити його застосовність у ширшому контексті мобільної розробки.

Публікація [11] представляє підхід MOBICAT, спрямований на автоматичну генерацію GUI-коду для Android-додатків за допомогою методів інженерії, керованої моделями. Цей підхід дозволяє зменшити ручну працю при створенні інтерфейсів користувача. Проте, дослідження зосереджене на специфічному аспекті розробки – генерації GUI-коду, що обмежує його застосовність у вирішенні ширшого спектру рутинних завдань у мобільній розробці.

У статті [8] аналізується застосування інструментів автоматизованого тестування, таких як Selenium, Appium і TestComplete, у проєктуванні та підтримці якісних мобільних застосунків. Автори розглядають ключові характеристики інструментів, зокрема їхню гнучкість та підтримку різних платформ. Незважаючи на корисний огляд, дослідження побудовано лише в межах тестування та не містить інформації щодо покращення інших типових задач під час розробки мобільного застосунку.

Інша публікація [4] також присвячена автоматизованому тестуванню. У ній автори описують специфіку вибору інструментів залежно від типу застосунку та бюджету команди. Стаття корисна для загального ознайомлення з підходами до автоматизації, але не акцентує увагу на мобільних розробниках як цільовій аудиторії. Крім того, більша частина прикладів пов'язана з веб-тестуванням, що зменшує релевантність у контексті мобільної інженерії.

Ще одне дослідження [2] зосереджене на автоматизації рутинних дій фронтенд-розробників. У статті запропоновано концепцію «єдиної платформи» – середовища, де об'єднано типові інструменти, що найчастіше використовуються під час повсякденної розробки. Автор демонструє можливість створення кастомізованого робочого середовища з підтримкою повторного використання компонентів і збору фідбеку для їх адаптивного оновлення. Платформа реалізована з урахуванням Lean Startup-моделі, що дозволяє швидко адаптувати рішення під потреби користувача. Водночас, попри цінність запропонованої архітектури, робота не охоплює проблеми мобільної розробки, зосереджуючись переважно на веб-контексті. Крім того, аналіз обмежується лише розробкою інтерфейсів, без розгляду супутніх рутинних процесів, таких як CI/CD, тестування чи збір аналітики.

Таким чином, більшість доступних публікацій або не приділяють достатньо уваги мобільній розробці, або ж досліджують продуктивність у загальному контексті, зосереджуючись на високорівневих підходах – архітектурі, тестуванні або застосуванні DevOps-практик. Інша частина робіт має вкрай вузьку спрямованість (наприклад, тестування) і стосується окремих технічних рішень, що не завжди масштабуються або є релевантними в довгостроковій перспективі. Це створює потребу в більш глибокому аналізі повторюваних задач саме в контексті мобільної розробки – з урахуванням її особливостей на різних етапах життєвого циклу програмного забезпечення.

Аналіз рутинних завдань під час розробки застосунку. Огляд рутинних задач доцільно здійснювати у зв'язку з етапами життєвого циклу програмного забезпечення. Це дозволить точніше визначити вузькі місця, що найменше автоматизовані.

На етапі збирання та уточнення вимог розробник часто стикається з повторюваними комунікаціями з менеджером або замовником, необхідністю формалізувати неструктуровану інформацію, адаптувати її до технічних обмежень платформи. Незважаючи на наявність систем управління вимогами, велика частка роботи тут залишається ручною.

Конфігурація проєкту на початку розробки також передбачає багато рутинної роботи: налаштування структури модулів, інтеграція SDK, налаштування CI/CD, тестових середовищ, схем збірки, signing configs тощо. Ці задачі часто повторюються від проєкту до проєкту з мінімальними відмінностями, однак потребують значного часу та уваги.

Під час проєктування інтерфейсу рутинними завданнями стають перенесення дизайн-макетів до коду, адаптація їх до різних екранів, відсутність узгоджених UI-компонентів або стилів. Ці завдання потребують багаторазових дрібних змін, перевірок і уточнень, що уповільнює роботу.

На етапі архітектурного проєктування розробник постійно повторює схожі дії: створення шаблонних елементів, підключення бібліотек, реалізація типових взаємодій із бекендом, написання однакових структур даних. Більшість команд не мають власних генераторів коду або DSL-інструментів, тому багато з цих процесів виконуються вручну.

Протягом реалізації функціоналу, типовими стають завдання інтеграції з API, обробка помилок, створення UI тощо. Ці задачі рідко є творчими, але вимагають значного часу та зусиль.

На етапі тестування розробник стикається з низкою рутинних дій, пов'язаних із забезпеченням якості продукту. Насамперед ідеться про написання юніт-тестів для перевірки логіки окремих компонентів та UI-тестів, які автоматично відтворюють сценарії взаємодії користувача з інтерфейсом. Окрім самого написання, важливо регулярно запускати ці тести, аналізувати результати, усувати причини можливих збоїв, а також підтримувати тести в актуальному стані в міру зміни функціоналу. Також слід враховувати, що навіть за наявності автоматизованих тестів, частина перевірок виконується вручну – однак ця ручна перевірка, як правило, відбувається ще на етапі розробки, коли інженер перевіряє працездатність нового функціоналу локально.

Що ж до фази релізу, то основна частина налаштувань вже має бути виконана раніше, на етапі конфігурації збірки та налаштування CI/CD-процесів. Втім, тут залишається чимало повторюваних завдань, пов'язаних із доналаштуванням параметрів релізної збірки, перевіркою цільових середовищ (наприклад, конфігурації build flavors), підготовкою скріншотів, метаданих, списків змін, а також запуском альфа- та бета-версій для обмеженого кола користувачів.

На етапі супроводу й обслуговування розробник витрачає час на перевірку аналітики, crash-репортів, відстеження проблем, які проявилися тільки в користувачів. Окрім того, виникає потреба в періодичному оновленні бібліотек, зміні політик платформ (наприклад, Google Play), які потребують швидкого реагування і спричиняють додаткові рутинні навантаження.

Окремо варто виділити роботу з фідбеком, як із боку кінцевих користувачів, так і замовників. Збір, інтерпретація, пріоритезація побажань і проблем це постійна діяльність, яка нерідко не має чіткої структури і здійснюється вручну через кілька джерел одночасно (відгуки в маркеті, email, месенджери, внутрішні трекери тощо).

Таким чином, майже кожен етап життєвого циклу мобільного застосунку супроводжується значним обсягом повторюваних завдань. Частина з них є технічно тривіальною, але вимагає витрат часу й уваги, відволікає від креативної або аналітичної діяльності, а в довгостроковій перспективі – знижує загальну ефективність команди.

Висновки. У ході аналізу було виявлено, що значна частина завдань, які щодня виконують мобільні розробники, має рутинний характер і часто залишається поза увагою дослідників. На основі систематизації процесів, характерних для всіх етапів життєвого циклу мобільного застосунку – від ініціалізації до підтримки та супроводу – можна зробити висновок, що існує потреба у переосмисленні підходів до оптимізації саме повторюваних дій. Попри широке впровадження засобів автоматизації, окремі сфери,

як-от локалізація чи фідбек-менеджмент, залишаються недостатньо дослідженими з точки зору їх впливу на загальну ефективність.

Поточне дослідження має аналітичний характер і не претендує на повне охоплення теми. Натомість його метою було окреслити карту повторюваних задач та виявити ділянки, що потребують подальшого вивчення. З огляду на це, перспективним напрямком наступного дослідження є поглиблений розгляд складнощів, які виникають під час вирішення типових рутинних завдань, а також систематизація існуючих підходів, їх переваг, обмежень та доцільності використання у контексті мобільної розробки.

Список використаних джерел:

1. GitLab Global DevSecOps Survey. URL: <https://about.gitlab.com/developer-survey>
2. Gupta S. Design of Controlled and Customizable Platform for Frontend Developer's Repetitive Tasks. *International Research Journal of Modernization in Engineering Technology and Science*. 2021. Vol. 3, Issue 11. С. 245–249. DOI: 10.13140/RG.2.2.17892.09604
3. JetBrains Developer Ecosystem Survey 2023. URL: <https://www.jetbrains.com/lp/devecosystem-2023/>
4. Muhammed Suhail TS. Automation: Deep Dive Into Web, Mobile & API – Part 1; *International Journal of Scientific and Research Publications (IJSRP)* 12(12) (ISSN: 2250-3153). 2022 DOI: 10.29322/IJSRP.12.12.2022.p13233
5. Nguyen Ba Long. Enterprise-grade CI/CD pipeline for mobile development: Bachelor's Thesis. – Metropolia University of Applied Sciences. 2022. URL: https://www.theseus.fi/bitstream/handle/10024/753156/Nguyen_Ba_Long.pdf
6. Repetitive Tasks at Work: Research and Statistics 2024. URL: <https://www.processmaker.com/blog/repetitive-tasks-at-work-research-and-statistics-2024/>
7. Time Spent on Recurring Tasks. Clockify Blog. URL: <https://clockify.me/time-spent-on-recurring-tasks>
8. Venkata Amarnath Rayudu Amisetty. AI-driven mobile and web automation: The CI/CD integration. revolution. *World Journal of Advanced Research and Reviews*, 2025, 26(02), 3554–3562. DOI: 10.30574/wjarr.2025.26.2.2039.
9. Vijay Bhasker Reddy Bhimanapati, Om Goel. Implementing CI/CD for Mobile Application Development in Highly Regulated Industries. *International Journal of Novel Research and Development*. 2023. DOI: 10.1234/jis.2020.0301
10. Wali, A., Mahamad, S., Sulaiman, S. Task Automation Intelligent Agents: A Review. *Future Internet* 2023, 15, 196. 2023. DOI: 10.3390/fi15060196
11. Zafar H, Ur Rehman Khan S, Mashkoor A and Nisa HU MOBICAT: a model-driven engineering approach for automatic GUI code generation for Android applications. *Front. Comput. Sci.* 2024 6:1397805. DOI: 10.3389/fcomp.2024.1397805

Дата надходження статті: 26.06.2025

Дата прийняття статті: 02.07.2025

Опубліковано: 23.09.2025

UDC 004.9

DOI <https://doi.org/10.32689/maup.it.2025.2.29>

Tamara FRANCHUK

Candidate of Economic Sciences,

Associate Professor at the Department of Accounting, Auditing and Taxation,

National Academy of Statistics, Accounting and Auditing, Tamara_Franchuk@ukr.net

ORCID: 0000-0001-7615-1276

Dmytro TYSHCHENKO

Candidate of Economic Sciences, Associate Professor,

Associate Professor at the Department of Software Engineering and Cybersecurity,

State University of Trade and Economics, tyshchenko_d@knute.edu.ua

ORCID: 0000-0002-2193-9012

Alona DESIATKO

Candidate of Technical Sciences, Associate Professor,

Acting Head of the Department of Software Engineering and Cybersecurity,

State University of Trade and Economics, desyatko@knute.edu.ua

ORCID: 0000-0002-2284-3418

Dmytro HNATCHENKO

Senior Lecturer at the Department of Software Engineering and Cybersecurity,

State University of Trade and Economics, hnatchenko@knute.edu.ua

ORCID: 0000-0002-6584-4525

ASSESSMENT OF THE USE OF CLOUD TECHNOLOGIES IN ACCOUNTING AS AN INNOVATIVE TOOL

Abstract. *The purpose* of the work is to determine the main levers for the application of cloud technologies in accounting as an innovative tool.

The methodology used in the work consists of identifying effective means of automatic data collection and tax reporting, as well as uploading bank statements in the required format. It has been established that developers of so-called digital asset technologies and software products require constant training and some testing of all service organization processes.

The scientific novelty of the work lies in identifying information literacy and effective financial management as key factors in increasing the level of financial stability and independence in the context of the development of the cloud environment.

Conclusions. It has been proven that financial awareness and effective financial management are key factors in achieving financial stability and independence, and therefore, with the help of cloud technologies available to us today, creating innovative tools that promote effective financial management is a necessary tool for the effective functioning of an enterprise. The main advantages of using cloud technologies in management accounting are investigated, taking into account the relevance of data protection issues, company efficiency, efficiency and mobility of management decision-making. The requirements for training and advanced training of industry specialists are considered. The peculiarities of the domestic, European and global markets, in particular legislative requirements and accounting standards, are analyzed. The criteria on which the organization of the accounting system in the cloud environment depends are substantiated. The current aspects of the creation and implementation of cryptocurrency as one of the digital assets that occupy a certain place in the formation of modern accounting and auditing, which is considered as one of the possible types of digital currency, which is designed to perform the functions of ensuring optimality and efficiency, are investigated. The main components of the functionality of software products to simplify accounting were analyzed. The functionality of software products, the features of working with different types of organizations, companies, enterprises, their interaction, integration with financial institutions, for effective optimization of data exchange between companies and specialists were studied.

Key words: cloud technologies, digitalization, software product, accounting, financial analysis.

Тамара ФРАНЧУК, Дмитро ТИЩЕНКО, Альона ДЕСЯТКО, Дмитро ГНАТЧЕНКО. ОЦІНКА ЗАСТОСУВАННЯ ХМАРНИХ ТЕХНОЛОГІЙ В БУХГАЛТЕРСЬКОМУ ОБЛІКУ ЯК ІНОВАЦІЙНОГО ІНСТРУМЕНТАРІЮ

Анотація. *Метою* роботи є визначення основних важелів застосування хмарних технологій в бухгалтерському обліку як інноваційного інструментарію.

Методологія, використана в роботі, полягає у визначенні ефективних засобів автоматичного збору даних та формування податкової звітності, а також завантаження банківських виписок у потрібному форматі. Встановлено, що розробники так званих технологій цифрових активів та програмних продуктів потребують постійного навчання та певного тестування всіх процесів організації сервісу.

© T. Franchuk, D. Tyshchenko, A. Desiatko, D. Hnatchenko, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

Наукова новизна роботи полягає у визначенні інформаційної обізнаності та ефективного управління фінансами як ключових факторів підвищення рівня фінансової стабільності та незалежності в умовах розвитку хмарного середовища. Обґрунтовано критерії, від яких залежить організація системи бухгалтерського обліку в хмарному середовищі.

Висновки. Доведено, що фінансова свідомість та ефективне управління фінансами є ключовими факторами для досягнення фінансової стабільності та незалежності, а отже за допомогою хмарних технологій, доступних нам сьогодні, створення інноваційних інструментів, що сприяють ефективному керуванню фінансами, є необхідним інструментом ефективного функціонування підприємства. Досліджено наслідки застосування хмарних технологій, з урахуванням актуальності питання захисту даних, ефективності роботи компанії, її оперативності та мобільності. Обґрунтовано вимоги до підготовки та підвищення кваліфікації фахівців галузі. Проаналізовано особливості вітчизняного, європейського та світового ринку, зокрема законодавчі вимоги та стандарти бухгалтерського обліку. Встановлено актуальні аспекти створення та впровадження криптовалюти як одного з цифрових активів, які займають певне місце у становленні сучасного обліку та аудиту, яку розглядають як один із можливих видів цифрової валюти, яка покликана виконувати функції забезпечення оптимальності та ефективності. Запропоновано вдосконалення компонентів функціональності програмних продуктів для спрощення бухгалтерського обліку. Обґрунтовано запропонований функціонал програмних продуктів, особливості роботи з різними типами організацій, компаній, підприємств, їх взаємодія, інтеграція з фінансовими установами, для ефективно оптимізації обміну даними між компаніями, фахівцями.

Ключові слова: хмарні технології, цифровізація, програмний продукт, бухгалтерський облік, фінансовий аналіз.

Introduction. In today's world, where financial awareness and effective financial management are becoming key factors in achieving financial stability and independence, the development of a virtual financial assistant is becoming particularly important. With the help of the technologies available to us today, the creation of innovative tools that contribute to effective financial management has become possible. Based on the research chosen by the authors, its purpose is to analyze and improve the use of cloud technologies for accounting and optimize financial reporting in the process of qualitative modification of cloud services.

Accounting application solutions require hardware for storing and processing information, while server capacities – computers – are used to host automated management and accounting systems. The server requires special placement conditions – a specially ventilated room, a staff of specialists to service it, etc. Conventionally, servers for business management automation systems can be divided into two groups: local equipment and cloud capabilities. Regarding local equipment, we note that it is computing hardware on which specialized programs are installed to perform computing operations.

The software operates locally and does not require an Internet connection. This is its advantage. Such machines can be: office servers, which can be used for software with low resource requirements; high-power personal computers, which are actually used as servers. Larger businesses with extensive IT infrastructure use large server machines and a staff of specialists to maintain them as computing equipment. Cloud computing is the same servers, super-powerful computer machines. All hardware is maintained by technical specialists and programmers. Such clouds have many advantages for accounting programs, compared to local servers. Like any other software, management, accounting, or other accounting systems can be stored and administered locally or on remote servers – clouds.

Analysis of recent research and publications. When studying modern domestic, European and global market features, in particular legislative requirements and accounting standards, it should be noted that the issue of using cloud technologies in accounting is a relevant issue and requires a comprehensive and thorough analysis to study and offer recommendations for optimizing accounting activities, which is reflected in recent publications by researchers. Shish, A. explores cloud technologies in accounting analysis in Ukraine and analyzes the main differences and strategies for adaptation to the local context [4]. Pramuka B., Pinasti M. consider due to the relevance of small business requests and comprehensively explore the main levers that rely on the use of cloud and network accounting systems, which is due to the relevance of small business requests a ways to address them [10]. Berdichevskaya V. considers the main aspects of using cloud technologies in accounting of organizations, examining their areas of application, main advantages, disadvantages, as well as potential problems that arise during their operation [6]. Fahmi M., Muda I., Kesuma S. analyze digitization technologies and the place of the enterprise in the process of introducing its developments into practical actions in the process of using accounting and auditing software products [7]. Ghatrifi M. O. M., Amairi J. S. S., Thottoli M. Investigating international experience in improving accounting education and training during the development of modern cloud technologies [9]. Mykhaylovyna S. O., Matros O. M., Polishchuk O. M. investigate food technologies as an important aspect of the development of the accounting and taxation system [3]. The problems of implementing modern cloud technologies and their widespread application in the field of accounting and analysis of financial activities often fall into the focus of scientific research by modern and foreign scientists, the results of theoretical research and practical achievements are presented by scientists in their works [1, 2, 5, 8, 11, 12].

Main part. The active development of innovative technologies and other changes that have swept the economy indicate that they are primarily due to the introduction of cloud technologies and the emergence

of digital assets such as cryptocurrencies. Today, researchers argue that cloud technologies are changing the theory and practice of financial reporting and accounting, which cannot be considered outside of the modern digitalization of information and other flows. Information technology specialists develop special software products, taking into account that these are new solutions that can help increase the efficiency of enterprise management, accounting and financial activities using digitalization tools. We can argue that there are a number of advantages of using cloud technologies, the most important of which are presented on (Fig. 1).

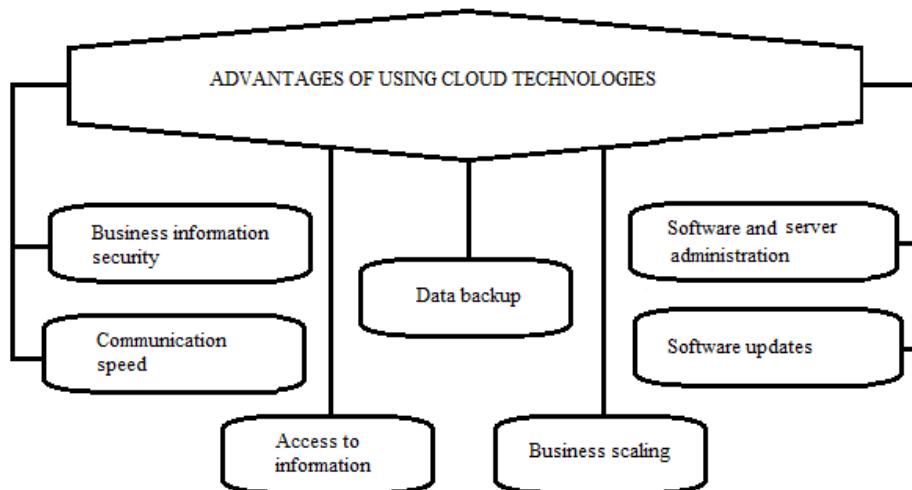


Fig. 1. Advantages of using cloud technologies

Accounting is a component without which the activities of any modern enterprise, institution or organization cannot do. The main task of modern accounting, in conditions of rapid changes and innovations, is to provide interested parties of the company with timely, fast and in the shortest possible time, namely: managers, investors, participants, shareholders, and management, the necessary and truthful information to make accurate, balanced, and relevant management decisions. An accountant must possess certain digital skills that relate to information literacy, communication and interaction, security and problem solving. Within the framework of communication and interaction, an accountant must be knowledgeable in modern technologies, be able to use the latest platforms and services, while interacting with information users in compliance with the rules of conduct within the framework of communication of this format. The criteria on which the organization of the accounting system depends are presented on (Fig. 2).

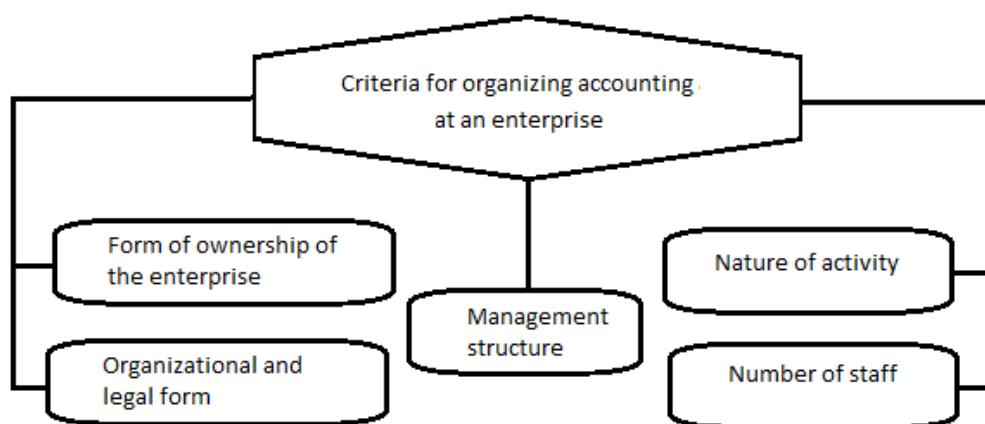


Fig. 2. The criteria on which the organization of the accounting system

When it comes to information literacy, an accountant is required to have a significant number of skills that go beyond the skills of the profession, such as: the ability to qualitatively select and filter data among the arrays of information in the digital space, recognize business processes and increase the use of modern information technologies in working with accounting information, etc. Let's analyze a main components of this programs functionality to simplify accounting, which represented on Figure 3 [5].

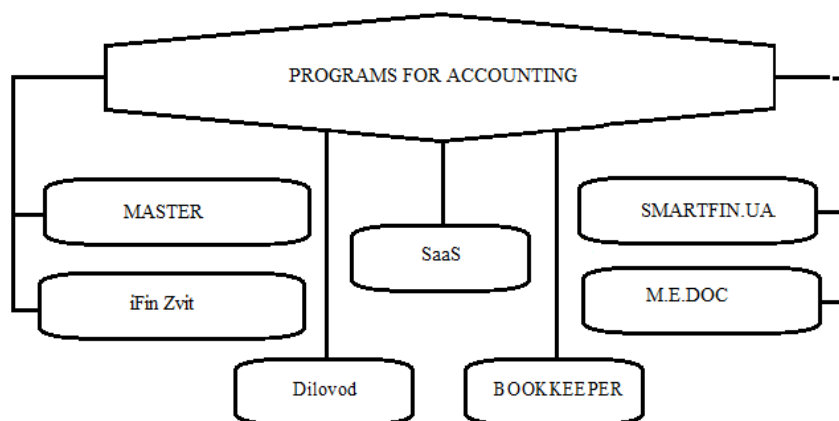


Fig. 3. The criteria on which the organization of the accounting system

Functionality of My Electronic Document (M.E.DOC): universality: suitable for use by entrepreneurs of any type of activity and various tax systems; modular structure: The program has a modular structure that allows you to customize individual functionality depending on the needs of the enterprise. Features include personnel records, payroll calculation, connection of the System of Mandatory and Centralized Accounting and Reporting, registration of excise invoices and transport consignment notes (TCN), as well as a full cycle of work with account books and tax invoices; electronic exchange: simplifies the process of exchanging electronic documents with counterparties and supports online reporting; electronic reporting: allows you to submit reports online to all regulatory authorities. Dilovod functionality: combining accounting and management records, control of turnover, registration of single contribution payers (SCP) and the ability to submit electronic reports – all these features in one program; the ability to keep records of several sole proprietors and legal entities on one database; automatic archiving of copies of documents and display of all monetary circulation in the system; an important bonus – automatic checking for errors in reporting; the ability to submit reports online, if there is Internet access. It can be an attractive alternative for entrepreneurs who are looking for a convenient and functional tool for accounting and reporting without significant efforts to learn accounting skills. Master functionality: accounting: integration with the client bank simplifies financial transactions on accounts. The program also allows you to keep records of sales, purchases, monitor warehouse balances and accurately determine actual production costs. Financial reports are generated automatically, providing the user with up-to-date information; personalized interface: master provides users with the ability to customize the interface according to their needs and preferences, which makes interaction with the program as convenient and effective as possible; allows reporting to regulatory authorities directly through the software interface, providing quick and convenient access to this option. With this approach, the subscription fee becomes more transparent and decreases proportionally to the number of additional functions installed. Using Bookkeeper the company's employees can easily enter primary documents, thereby leaving the accountant with the tasks of checking the correctness of the data and compiling reports. Functionality: work with different types of organizations, for interaction with non-profit organizations, sole proprietors and legal entities; integration with Privat-24, which allows you to speed up the circulation of documents between you, the accountant and your counterparties; operational and tax accounting of sole proprietors, helps to maintain operational accounting and monitor tax obligations of sole proprietors; automatic collection of tax reports, generates tax reports automatically, simplifying your work; downloading bank statements in DBF format.

The functionality of SaaS depends on the selected version, suitable for different types of enterprises: standard version is suitable for small and medium-sized enterprises; corporate version is an assistant for large corporations; industry version contains specialized functionality for niche enterprises. The Debit Plus system functionality: work with wages, tariffs, automatic surcharges; generation of certificates and export to the bank; work with material assets, stocks, balances; turnover balance sheets and accounting cards; work with balance sheets and generation of forms, statements, general ledger; work with fixed assets, inventory cards, acts, depreciation accrual; work with cash and banking transactions, registration journal, advance reports, cash orders; does not support electronic reporting. Full functionality of iFin Zvit is available even when using gadgets, smartphones and tablets that can be used remotely. Not need to manually update the program, since iFin Zvit automatically adapts to changes in legislation. Functionality: ability to synchronize with the latest version of 1C, if it is installed on the computer, and transfer data and reports; exchange of tax invoices with clients, their registration in the Unified Register of Tax Invoices and submission of requests; conclusion

of contracts with regulatory authorities, warehouse accounting and payroll calculations; provision of expert consultations in cases of unclear situations; the ability to send reports to the Treasury, Pension Fund and Tax Inspectorate. Smartfin.ua allows not only professionals, but also entrepreneurs without special specialists to keep accounting records independently. Smartfin.ua functionality includes: work time accounting; reporting; control over settlements and cash transactions; calculation of salaries, bonuses and allowances; accounting for all trade transactions; maintenance of employment contracts and employee cards; the ability to send electronic reports through your personal account [5].

The studied cloud products are usually complex products, not limited to providing the right to use the software, but also including other additional services depending on the selected cloud services, including ensuring reliable isolation of the resource pool provided to the client from the resource pools of other clients and system technology stack using server and network virtualization tools, confidentiality and protection of customer data, ensuring the continuous functioning of relevant services. Analyzing the ways and prospects for the development of accounting and certification, it should be noted that the convergence of cloud technologies and digital assets means changing the ways of optimizing accounting and conducting financial activities. WThis procedure includes identifying and analyzing all controls, determining which controls should be implemented, and identifying gaps that need to be remedied. During testing, specialists verify and confirm changes in control areas and identify them through walkthroughs, observations, and evidence checks. At the end of testing and other activities, an analysis is performed and necessary recommendations are introduced to improve the efficiency of the system. In fact, clouds are suitable for all participants in business processes, in particular for financial services – optimization and savings on the purchase and maintenance of equipment, salaries for a full-time specialist or an outsourced programmer. The advantages of using cloud services for accounting automation are obvious.

Conclusions. Ukraine's active integration into the European and global economic space contributes to continuous changes and improvements in the system of accounting standards and financial reporting. This is due to both changes in the economy and the fact that the domestic accounting system is systematically combined with international financial reporting standards (IFRS). The IFRS system itself is also constantly being improved, which necessitates the constant adaptation of the national accounting standards system to updated international standards. Under such conditions, the use of software products for accounting control and financial reporting generation aims to ensure high-quality and proper interaction between programming specialists and accounting specialists when developing software products. Server duplication and backup technology allows the user not to worry about the reliability of information, so any data can be restored quite quickly. Access to the accounting program is via remote desktop or via a Web browser. There is no need to download the configuration to your computer or use office facilities. The only condition for a quick response from cloud software is the speed of the Internet connection. In the local version, access to the system is possible only if the host computer is working. The convenience of using cloud technologies is that the capacity can be scaled by increasing the space on the server as needed, or vice versa, reducing it. Permission to the information base can be restricted in a matter of minutes, while this is not easy to implement with local capacities. Quite often, it is information data security that is the condition that restrains small business owners from migrating to the cloud, which requires additional attention from researchers.

Bibliography:

1. Капітон А., Сухоребрий О., Ненич Д. Використання мультимодального штучного інтелекту в економіці, освіті, науці та транспорті. *Інформаційні технології та цифрова економіка*. Київ: ДУІТ, 2024. С. 83–85.
2. Краус К. Цифрові та інституційні зміни в економічних відносинах: компаративний аналіз і оцінювання можливих структурних зрушень. *Галицький економічний вісник*. 2025. № 92. № 1. С. 7–17.
3. Михайловина С. О., Матрос О. М., Поліщук О. М. «Хмарні» технології як важливий аспект розвитку системи бухгалтерського обліку і оподаткування. *Ефективна економіка*, 2021. № 8. URL: <http://www.economy.nayka.com.ua/?op=1&z=9153> (дата звернення: 20 лютого 2025).
4. Шиш А. Хмарні технології у бухгалтерському обліку та фінансовому аналізі в Україні: аналіз відмінностей та стратегії адаптації до місцевого контексту. *Здобутки економіки: перспективи та інновації*, 2024. № 2. URL: <https://doi.org/10.57125/econp.2024.01.29.02> (дата звернення: 20 лютого 2025).
5. AI identification tools URL: <https://thetransmitted.com/ai/instrumenty-identyfikacziyi-shi-zhovten-2024/> (дата звернення: 17 лютого 2025)
6. Berdichevskaya V. Cloud technologies in accounting of organizations: scopes, advantages and problems of use. *Vestnik NSUEM*, 2023. no. 1. URL: <https://doi.org/10.34020/2073-6495-2023-1-099-107> (дата звернення: 20 лютого 2025).
7. Fahmi M., Muda I., Kesuma S. A. Digitization Technologies and Contributions to Companies towards Accounting and Auditing Practices. *International Journal of Social Service and Research*, 2023. Vol. 3, no. 3. URL: <https://doi.org/10.46799/ijssr.v3i3.298> (дата звернення: 20 лютого 2025).

8. Franchuk T., Tyshchenko D., Desiatko A., Karpunin I. Features of accounting digitalization processes. *Galician economic journal*, 2025, vol. 95, no 1, pp. 61–66.
9. Ghatrifi M. O. M., Amairi J. S. S., Thottoli M. M. Surfing the technology wave. *Computers and Education: Artificial Intelligence*. 2023. no. 4. URL: <https://doi.org/10.1016/j.caeai.2023.100144> (дата звернення: 20 лютого 2025).
10. Pramuka B., Pinasti M. Does cloud-based accounting information system harmonize the small business needs? *Journal of Information and Organizational Sciences*, 2020. Vol. 44, no. 1. URL: <https://doi.org/10.31341/jios.44.1.6> (дата звернення: 20 лютого 2025).
11. Tyshchenko D., Franchuk T., Zakharov R., Moskalenko V. Підтримка динамічних потреб безпеки засобами VPN Системи управління, навігації та зв'язку. Полтава: ПНТУ, 2024. Т. 3 (77). С. 149–152.
12. Tyshchenko D., Franchuk T., Stepashkina K., Karpunin, I. Проектування та розробка системи корпоративного електронного документообігу. *Європейський науковий журнал Економічних та Фінансових інновацій*. 2024. № 1(13). С. 200–207.

Дата надходження статті: 11.06.2025

Дата прийняття статті: 30.06.2025

Опубліковано: 23.09.2025

УДК 004.942:614.8

DOI <https://doi.org/10.32689/maup.it.2025.2.30>

Олексій ХОНІН

аспірант,

Приватний вищий навчальний заклад «Європейський університет»,

oleksii.o.khonin@e-u.edu.ua

ORCID: 0009-0007-9328-7870

Олена СКЛЯРЕНКО

кандидат фізико-математичних наук, доцент,

завідувач кафедри математичних дисциплін та інноваційного проектування,

Приватний вищий навчальний заклад «Європейський університет»

ORCID: 0000-0001-6555-1223

Scopus-Author ID: 57192677216

АЛГОРИТМИ ОПТИМАЛЬНОЇ МАРШРУТИЗАЦІЇ ДЛЯ СЛУЖБ ЕКСТРЕНИХ РЕАГУВАНЬ

Анотація. Метою дослідження є питання підвищення ефективності реагування служб надзвичайних ситуацій шляхом оптимізації алгоритмів маршрутизації та розробка математичної моделі маршрутизації та алгоритмів, які забезпечують найкоротші шляхи з урахуванням змін дорожньої ситуації в реальному часі.

Методологія. У дослідженні використано методи теорії графів, евристичні та адаптивні алгоритми, а також симуляційне моделювання для порівняння ефективності різних підходів.

Наукова новизна. Уперше запропоновано комплексний підхід до побудови маршруту з урахуванням динамічної зміни дорожньої обстановки в реальному часі, що базується на адаптивному алгоритмі з механізмом перепланування. На відміну від класичних методів, розроблений алгоритм дозволяє своєчасно коригувати маршрут у разі виявлення заторів, аварій чи перекриттів, забезпечуючи стійкість рішень та мінімізацію часу прибуття. Крім того, дослідження передбачає масштабованість алгоритмів для різних типів населених пунктів – від мегаполісів до сільської місцевості – з урахуванням особливостей дорожньої інфраструктури. Важливою складовою новизни є інтеграція моделі в реальні системи підтримки прийняття рішень, що дає змогу не лише моделювати, але й практично застосовувати результати.

Висновки. Запропонований адаптивний підхід до маршрутизації дозволяє суттєво зменшити середній час реагування служб надзвичайних ситуацій, що підтверджено результатами симуляційного моделювання. Розроблені алгоритми забезпечують гнучке врахування динамічної дорожньої обстановки в режимі реального часу та можуть бути інтегровані у сучасні диспетчерські системи. Масштабованість і технологічна сумісність рішення робить його перспективним для впровадження у різних адміністративно-територіальних умовах України.

Ключові слова: алгоритми маршрутизації, служби надзвичайних ситуацій, оптимізація, час реагування, теорія графів, адаптивні алгоритми, транспортна модель.

Oleksii KHONIN, Olena SKLIARENKO. OPTIMAL ROUTING ALGORITHMS FOR EMERGENCY RESPONSE SERVICES

Abstract. The purpose of this study is to enhance the efficiency of emergency response services by optimizing routing algorithms and developing a mathematical model and algorithms that ensure the shortest paths while accounting for real-time changes in road conditions.

Methodology. The study employs methods of graph theory, heuristic and adaptive algorithms, as well as simulation modeling to compare the effectiveness of various approaches.

Scientific novelty. A comprehensive approach to route planning is proposed for the first time, taking into account dynamic changes in traffic conditions in real time. It is based on an adaptive algorithm with a replanning mechanism. Unlike classical methods, the developed algorithm allows for timely route adjustments in case of traffic jams, accidents, or road closures, ensuring solution robustness and minimizing arrival time. Moreover, the study considers the scalability of the algorithms for various types of settlements – from large metropolitan areas to rural regions – while accounting for specific features of the transportation infrastructure. An important aspect of the novelty lies in the integration of the model into real-world decision support systems, enabling not only simulation but also practical implementation of the results.

Conclusions. The proposed adaptive routing approach significantly reduces the average response time of emergency services, as confirmed by simulation results. The developed algorithms provide flexible handling of real-time traffic dynamics and can be integrated into modern dispatch systems. The solution's scalability and technological compatibility make it promising for implementation in various administrative and territorial contexts of Ukraine.

Key words: routing algorithms, emergency response services, optimization, response time, graph theory, adaptive algorithms, transportation model.

© О. Хонін, О. Скляренко, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

Вступ. Швидкість прибуття служб екстреного реагування є визначальним чинником для збереження життів та мінімізації збитків у надзвичайних ситуаціях. В умовах зростання урбанізації та транспортного навантаження вибір оптимального маршруту ускладнюється динамікою дорожньої обстановки. Актуальність дослідження зумовлена зростанням навантаження на служби реагування в умовах урбанізації, інтенсивного автомобільного руху та частих надзвичайних подій.

Традиційні алгоритми маршрутизації, як-от Дейкстра чи A*, не враховують зміни в реальному часі, що знижує їхню ефективність. Це зумовлює необхідність розробки адаптивних рішень, здатних швидко реагувати на затори, аварії та інші перешкоди.

Метою дослідження є створення алгоритмів маршрутизації для служб 101/112, які забезпечують мінімальний час реагування завдяки врахуванню змін дорожніх умов у реальному часі.

Постановка проблеми. Служби екстреного реагування часто стикаються з проблемою вибору найшвидшого маршруту до місця події в умовах динамічної дорожньої ситуації. Традиційні алгоритми не враховують затори, аварії або перекриття доріг у реальному часі, що призводить до затримок і зниження ефективності реагування. Відсутність адаптивних алгоритмів, здатних оперативно перебудовувати маршрут з урахуванням актуальних даних, обмежує можливості сучасних диспетчерських систем. Це створює потребу у нових підходах до оптимізації маршрутів із використанням інтелектуальних методів і симуляційного моделювання [1; 4].

Аналіз досліджень і публікацій. Проблема побудови ефективних маршрутів у транспортних мережах розглядається в науковій літературі вже понад півстоліття. Серед класичних підходів до вирішення задач маршрутизації найпоширенішими є алгоритм Дейкстри, алгоритм A* та алгоритм Флойда-Уоршелла [8]. Ці алгоритми забезпечують знаходження найкоротшого шляху в зважених графах, однак мають певні обмеження щодо масштабованості та адаптації до динамічних змін трафіку.

У сучасних дослідженнях усе більше уваги приділяється адаптивним та інтелектуальним методам оптимізації. Зокрема, розглядаються підходи на основі динамічних графів, у яких ваги ребер змінюються в режимі реального часу, залежно від навантаження та заторів [3]. Використання методів машинного навчання для прогнозування затримок, вибору маршруту або оцінки ризиків дозволяє підвищити точність та ефективність планування [2].

Особливої актуальності набувають алгоритми маршрутизації в контексті систем екстреного реагування. У таких системах важливим є не лише скорочення відстані, а й мінімізація часу з урахуванням пріоритетів, типу надзвичайної ситуації та ресурсів підрозділів. У літературі представлено різноманітні підходи: багатокритеріальна оптимізація, моделі із часовими вікнами, використання геоінформаційних систем (ГІС) у поєднанні з GPS [5; 6; 8].

Попри значні напрацювання, досі спостерігається низка обмежень: недостатня адаптація алгоритмів до умов міських мереж з високим рівнем нестабільності; обмежена інтеграція з оперативними диспетчерськими системами; відсутність повноцінної перевірки рішень у реальних сценаріях, зокрема в українських умовах [1; 4].

Це вказує на необхідність розроблення нових або вдосконалених моделей, що враховують особливості роботи служб реагування в умовах обмеженого часу, високого навантаження на дороги та змінного середовища.

Мета та завдання дослідження. Метою статті є підвищення ефективності реагування служб надзвичайних ситуацій шляхом розробки та впровадження алгоритмів маршрутизації, здатних адаптуватися до змін дорожніх умов у реальному часі [9; 10].

Для досягнення поставленої мети сформульовано наступні завдання:

- побудова математичної моделі маршрутизації, яка враховує часові затримки, завантаженість доріг та тип події;
- аналіз ефективності класичних і сучасних алгоритмів на основі експериментального моделювання;
- розробка адаптивного алгоритму маршрутизації для екстрених служб з динамічною переоцінкою ваг ребер у графі;
- реалізація програмного прототипу системи маршрутизації;
- оцінка масштабованості рішень для умов різних міських і сільських локацій.

Виклад основного матеріалу. У рамках дослідження було розроблено адаптивний алгоритм маршрутизації, який поєднує класичні методи графового пошуку з механізмами динамічного оновлення даних про дорожню ситуацію в режимі реального часу. В основі рішення – модифікований алгоритм A*, розширений блоком оцінювання актуального стану транспортної мережі, що базується на телематичних та GPS-даних.

Для перевірки ефективності було побудовано транспортну модель міської мережі, що враховує типи доріг, обмеження швидкості, ймовірність виникнення заторів та непередбачених подій (аварій,

ремонтів). Проведено серію симуляцій з порівнянням роботи класичних та адаптивних алгоритмів на різних сценаріях. Результати демонструють скорочення середнього часу прибуття на 12–27%, що підтверджує ефективність запропонованого підходу.

Запропоноване рішення може бути інтегроване в системи диспетчеризації служб 101/112 та масштабовано для застосування в містах з різною інфраструктурною складністю.

У роботі використано орієнтований граф $G = (V, E)$ $G = (V, E)$ $G = (V, E)$, де:

- V – множина вершин, що представляють ключові точки дорожньої мережі (перехрестя, вузли доступу);

- E – множина дуг з динамічними вагами $w_e(t)$, які змінюються в залежності від поточної ситуації (затор, обмеження швидкості, аварії).

Ваги дуг визначаються як:

$$w_e(t) = d_e + \delta_e(t) \quad w_e(t) = \frac{d_e}{v_e(t)} + \delta_e(t) \quad w_e(t) = v_e(t) d_e + \delta_e(t), \text{ де:}$$

- d_e – довжина ділянки;
- $v_e(t)$ – фактична швидкість на ділянці у момент t ;
- $\delta_e(t)$ – затримка, зумовлена особливими обставинами (наприклад, аварією або перекриттям).

Псевдочод запропонованого адаптивного алгоритму:

Input: граф $G(V, E)$, початкова вершина s , цільова вершина t

Ініціалізувати чергу пріоритетів Q :

Для кожної вершини $v \in V$: $\text{dist}[v] \leftarrow \infty$

$\text{dist}[s] \leftarrow 0$

$Q.\text{enqueue}(s, \text{priority}=0)$

Поки Q не порожня:

$u \leftarrow Q.\text{dequeue}()$

Якщо $u == t$: завершити

Для кожного сусіда v з u :

$w \leftarrow \text{обчислити_вагу}(u, v, t)$

Якщо $\text{dist}[v] > \text{dist}[u] + w$:

$\text{dist}[v] \leftarrow \text{dist}[u] + w$

$Q.\text{enqueue}(v, \text{priority}=\text{dist}[v])$

Якщо під час руху виявлено перекриття ребра:

оновити вагу $w_e \rightarrow \infty$

повторити розрахунок маршруту з поточної вершини

Адаптація маршруту в реальному часі:

- кожні NNN секунд система отримує оновлені GPS- або трафік-дані;
- переоцінює ваги найближчих ребер;
- у разі перевищення допустимого порогу затримки – виконує replanning (повторне планування маршруту).

Симуляційне моделювання виконувалося з використанням SUMO (Simulation of Urban Mobility) на основі реальних карт міст України, імпортованих з OpenStreetMap. Було змодельовано понад 500 сценаріїв руху для служб реагування у різних умовах.

Програмний прототип алгоритму реалізовано у вигляді модуля, що інтегрується із диспетчерськими системами та працює у зв'язці із симулятором трафіку. Проведено серію експериментів, де порівнювалися класичні алгоритми маршрутизації з адаптивним.

Скорочення часу реагування склало від 1,0 до 2,4 хвилин у середньому. У 83% випадків маршрут, обраний адаптивним алгоритмом, не потребував корекції навіть при зміні умов.

Також була протестована інтеграція з прототипом диспетчерської системи, яка дозволяла у реальному часі візуалізувати маршрут і час прибуття.

Отримані результати підтверджують ефективність адаптивного алгоритму у порівнянні з класичними підходами. Зниження часу реагування на 12–27% є суттєвим показником, особливо в умовах, коли кожна хвилина може мати вирішальне значення для порятунку життя чи мінімізації шкоди. Запропонована методика виявилась ефективною як у міських умовах із високим трафіком, так і в умовах сільської місцевості з обмеженим вибором маршрутів.

Таблиця 1

Середній час прибуття в різних підходах (в хвилинах)

Алгоритм	Час (середній)	Відхилення (%)
Дейкстра	9,8	+ 21
A*	9,1	+ 12
Floyd-Warshall	10,5	+ 32
Запропонований	8,1	0

Джерело: сформовано авторами на основі [1; 4; 7]

Завдяки динамічній зміні ваг графа в режимі реального часу, система здатна оперативно адаптуватися до аварій, заторів та перекриттів. Переваги розробленого підходу:

- висока швидкість розрахунку маршрутів навіть при великій кількості вершин (до 5000);
- стійкість до змін у дорожній ситуації;
- сумісність з сучасними геоінформаційними системами;
- можливість використання у вже існуючих диспетчерських центрах.

Обмеження:

- необхідність наявності надійних джерел даних про стан дорожньої мережі;
- неповне врахування людського фактора (наприклад, поведінка інших учасників руху);
- потреба у регулярному оновленні бази даних маршрутів.

Для перевірки масштабованості та стабільності роботи алгоритму проведено тестування на картах міст різного розміру та щільності забудови: умовно мале (100 тис. жителів), середнє (300 тис.) і велике (понад 500 тис.). При збільшенні кількості вершин графа з 1000 до 5000 час обчислення маршруту залишався в межах 1,2–2,3 с, що є допустимим у реальному часі.

Було також проведено серію симуляцій у сільських районах з мінімальною кількістю альтернативних маршрутів. У таких умовах алгоритм демонстрував менший приріст ефективності (приблизно 8–12%), але зберігав стабільну роботу та адаптацію до змін. Запропоноване рішення також протестовано для: швидкої медичної допомоги – з врахуванням типу виклику; ДСНС – з пріоритетами для пожеж, обвалів, затоплень; поліції – у випадках термінового затримання або охорони об'єктів; газових та електроаварійних служб – у разі обривів мереж або загрози вибухів.

Висновки. У межах даного дослідження було розроблено і протестовано адаптивний алгоритм маршрутизації для служб екстреного реагування. Результати показали зменшення середнього часу реагування на 12–27% порівняно з традиційними методами (алгоритми Дейкстри та A*), що підтверджує ефективність запропонованого підходу. Практична цінність роботи полягає у можливості інтеграції алгоритмів у системи диспетчеризації служб 101/112, а також у масштабованості рішення для різних міст і регіонів України.

Запропонований підхід дозволяє враховувати динамічні зміни дорожньої ситуації, забезпечує зменшення часу реагування на надзвичайні події і демонструє хорошу масштабованість. Наукова новизна полягає в поєднанні класичних евристичних методів (Dijkstra, A*) з динамічним переобрахунком ваг графа та процедурою replanning. Практична цінність – у можливості впровадження в системи диспетчеризації служб 101/112 або міських геоінформаційних платформ. Напрями подальших досліджень:

- розширення алгоритму з урахуванням погодних умов, часу доби, сезонності;
- інтеграція з розумними транспортними системами SmartCity;
- створення повноцінного веб-інтерфейсу для оперативного доступу диспетчерів;
- тестування у реальних умовах спільно з обласними та міськими підрозділами служб.

Список використаних джерел:

1. Hojat B., Ilbeigi M. Adaptive Emergency Evacuation Routing: A Graph-Based Approach. Tongji University Press / Elsevier, 2024, 292 p.
2. Lakhno V. A., Kasatkin D. Y., Skliarenko O. V., Kolodinska Y. O. Modeling and Optimization of Discrete Evolutionary Systems of Information Security Management in a Random Environment. *Machine Learning and Autonomous Systems. Smart Innovation, Systems and Technologies*, 2022, Vol. 269, pp. 9–22. Springer, Singapore. DOI: https://doi.org/10.1007/978-981-16-7996-4_2
3. Li X., Shahidehpour M. Urban emergency routing and traffic resilience with real-time data. *IEEE Transactions on Smart Transportation Systems*. 2021. Vol. 2, No. 3. P. 214–225.
4. Mahmoud K., Taha M. Dynamic route optimization for emergency vehicle dispatch using real-time traffic data. *Egyptian Informatics Journal*. 2022. Vol. 23, No. 1. P. 27–35.
5. Shiri D., Akbari V., Salman F. S. Online algorithms for ambulance routing in disaster response with time-varying victim conditions. *OR Spectrum*. 2024. Vol. 46, pp. 785–819. doi:10.1007/s00291-024-00744-4

6. Talaat F. M., Gamel S. A. Smart navigation system for emergency vehicles (SNSEV): utilizing fog and cloud computing technology for real-time traffic management. *Neural Computing and Applications*. 2025. Vol. 37, pp. 15547–15571.
7. Taylor & Francis Group, Evolutionary Optimization Algorithms. Chichester, UK: Taylor & Francis Group, 2021.
8. Toth P., Vigo D. Vehicle Routing: Problems, Methods, and Applications. SIAM, 2024 (updated edition), 527 p.
9. Trisolvena M. N., Wattimena F. Y., Untajana P. P. Logistics Efficiency in Product Distribution with Genetic Algorithms for Optimal Routes. *International Journal of Software Engineering and Computer Science (IJSECS)*, 2024, Vol. 4, No. 1, pp. 247–262. Accessed on: Mar. 22, 2024. [Online]. Available: <https://journal.lembagakita.org/index.php/ijsecs/article/view/2045>
10. Zverovich V. Modern Applications of Graph Theory. Oxford University Press, 2021, 368 p.

Дата надходження статті: 18.06.2025

Дата прийняття статті: 25.06.2025

Опубліковано: 23.09.2025

UDC 004.056

DOI <https://doi.org/10.32689/maup.it.2025.2.31>**Maksym CHEREMNOV**

Postgraduate Student at the Department of Computer Engineering and Innovative Technologies,
International Humanitarian University,
cheremnovmaksym@gmail.com
ORCID: 0009-0006-4936-4747

TRANSFORMATION OF INTERNATIONAL CYBERSECURITY STANDARDS AS A RESPONSE TO TECHNOLOGICAL AND GEOPOLITICAL CHALLENGES

Abstract. The article is aimed at analyzing the transformation of international cybersecurity standards (ISO/IEC 27001, NIST Cybersecurity Framework, GDPR) as a response to technological (AI, IoT, 5G) and geopolitical challenges (state-sponsored cyberattacks, regulatory fragmentation).

The study aims to identify key areas of evolution of standards, assess their adaptation to current threats, and suggest ways to improve global cyber resilience through cross-sectoral cooperation.

The study is based on the analysis of verified data from reputable sources, including IBM Security Cost of a Data Breach 2024, ENISA 2023, Verizon DBIR 2024, Microsoft Threat Intelligence Report 2023, as well as official documents of international organizations (EU, NIST, ISO). A systematic approach was used to assess technological, geopolitical and regulatory factors affecting cybersecurity standards. A comparative analysis of standards (ISO/IEC 27001:2022, NIST CSF 2.0, GDPR, NIS2) allowed us to identify their updates and limitations. Additionally, cross-sectoral cooperation initiatives such as the Cybersecurity Tech Accord and the ENISA MeliCERTes platform were analyzed to assess their role in the implementation of standards.

The article offers a comprehensive analysis of the transformation of cybersecurity standards with a focus on their adaptation to new technological threats (AI, IoT) and geopolitical realities (state-sponsored attacks). The novelty lies in the consideration of cross-sectoral cooperation as a key mechanism for implementing standards, which has not been sufficiently covered in the literature. The study also systematizes standard updates (e.g., the "Govern" feature in NIST CSF 2.0, the integration of DevSecOps in ISO/IEC 27001:2022) and assesses their impact on global harmonization, highlighting the barriers associated with regulatory fragmentation.

The transformation of international cybersecurity standards is necessary to counter modern threats. Updates to ISO/IEC 27001:2022 and NIST CSF 2.0 address risks from AI, IoT, and supply chains, while GDPR and NIS2 strengthen data protection and incident response. However, political differences, such as between Western and Chinese standards (GB/T), complicate global harmonization. The human factor, which accounts for 68% of data breaches, calls for increased cyber literacy. Cross-sectoral cooperation, such as the ENISA and CISA initiatives, is critical to the practical implementation of standards. In the future, cyber resilience will depend on flexible standards, coordination between states, the private sector and international organizations, and investment in education and technology.

Key words: cybersecurity, international standards, AI, IoT, geopolitical challenges, ISO/IEC 27001, NIST.

Максим ЧЕРЕМНОВ. ТРАНСФОРМАЦІЯ МІЖНАРОДНИХ СТАНДАРТІВ КІБЕРБЕЗПЕКИ ЯК ВІДПОВІДЬ НА ТЕХНОЛОГІЧНІ ТА ГЕОПОЛІТИЧНІ ВИКЛИКИ

Анотація. Стаття спрямована на аналіз можливості та доцільності змін в міжнародних стандартах кібербезпеки (ISO/IEC 27001, NIST Cybersecurity Framework, GDPR) з огляду на технологічні (ШІ, IoT, 5G) та геополітичні чинники.

Мета роботи. полягає у визначенні ключових напрямів еволюції стандартів, оцінці рівня їхньої адаптації до сучасних загроз і формування пропозицій щодо підвищення глобальної кіберстійкості через міжсекторальну співпрацю.

Методологія. В дослідженні використано системний підхід до оцінки технологічних, геополітичних і регуляторних факторів, що впливають на стандарти кібербезпеки. Проведено порівняльний аналіз стандартів ISO/IEC 27001:2022, NIST CSF 2.0, GDPR, NIS2, який дозволив виявити їхні обмеження в контексті змін зовнішнього середовища.

Наукова новизна полягає, насамперед, у фокусуванні дослідження на міжсекторальній співпраці як ключового механізму реалізації стандартів, що раніше висвітлювалося недостатньо. В цьому контексті проаналізовано ініціативи міжсекторальної співпраці, такі як Cybersecurity Tech Accord і ENISA MeliCERTes, для оцінки їхньої ролі в можливій трансформації наявних стандартів. Визначено, що останні оновлення ISO/IEC 27001:2022 і NIST CSF 2.0 враховують ризики від ШІ, IoT і ланцюгів постачання, тоді як GDPR і NIS2 посилюють захист даних і реагування на інциденти. Проте політичні розбіжності, зокрема між західними та китайськими стандартами (GB/T), ускладнюють глобальну гармонізацію. Паралельно із цим визначено, що людський фактор, що спричиняє 68% витоків даних, вимагає посилення кіберграмотності.

Висновки. За підсумками дослідження зроблено висновок, що у майбутньому глобальна кіберстійкість залежатиме від гнучкості стандартів та ефективної координації між державами, приватним сектором і міжнародними організаціями за умови належних інвестицій у освіту й технології.

Ключові слова: кібербезпека, міжнародні стандарти, ШІ, IoT, геополітичні виклики, ISO/IEC 27001, NIST.

© M. Cheremnov, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

Problem statement. In the modern world, cyberspace has become an environment for economic, social and political activities. At the same time, the rapid development of technologies, including artificial intelligence, the Internet of Things and 5G, as well as the growth of cyberattacks (primarily strategic ones, as an element of military and geopolitical confrontations), create new challenges for cybersecurity. Therefore, existing international standards, such as ISO/IEC 27001, NIST Cybersecurity Framework, and GDPR, need to be transformed to effectively respond to these challenges. The key problems are insufficient adaptation of standards to technological innovations, regulatory fragmentation, and geopolitical barriers that complicate global harmonization and reduce cyber resilience.

Analysis of recent research and publications. A number of scientific works are devoted to research on standardization in the context of ensuring effective cyber defense. Thus, S. Arsila, N. Pritam, S. Kepple [1] in their report on the investigation of data leaks drew attention to the need to revise approaches to risk assessment. The model for assessing risk management processes proposed by B. Barafor, A. L. Mesquida, A. Mas [2] is based on the ISO 31000 standard and integrates it into the context of multiple standards, which promotes an integrated approach to cybersecurity. At the same time, M. Edwards [7] emphasized the importance of updating the ISO 27001:2022 standard, in particular Annex A 5.23, which regulates the security of cloud services. Tanvir [20] notes that reliable communication frameworks for IoT play a critical role in ensuring data security in networks with a large number of connected devices. In addition, the GB/T 22239-2019 standard [9] establishes basic requirements for classified cybersecurity protection, which is relevant, in particular, for Asian countries.

Regarding the coverage of geopolitical factors, A. Venkat [21] emphasized that geopolitical conflicts are increasingly becoming a catalyst for cyberattacks, which prompts the creation of international initiatives such as the Cooperative Cyber Defense Cooperation [12]. K. Siglik and J. Gehring [5] propose a multilateral approach to ensuring peace in cyberspace, emphasizing the need for cooperation between states, the private sector, and international organizations. In this context, the NIS 2 Directive [15] and Regulation (EU) 2019/881 [17] set higher standards for the protection of critical infrastructure and cybersecurity certification in the EU. Increased cyber threats, including ransomware, also affect standards. S. Morgan [14] predicts that global losses from ransomware will reach \$57 billion in 2025, which emphasizes the need to strengthen incident response standards such as ISO/IEC 27035-1:2023 [11]. A. Ribeiro [19] in the ENISA 2024 report points to the growth of attacks on availability and data, which requires the adaptation of standards to new challenges. IBM's 2024 report [6] confirms that data breaches remain one of the most costly problems for organizations, prompting the implementation of standards such as ISO 27001:2022 [10].

At the legislative level, cybersecurity standards are also evolving. L. Y. Chang [4] analyzes the Budapest Convention as a basis for combating cybercrime, which is especially relevant for Asian countries. In turn, the California Personal Data Protection Act of 2018 [3] establishes new requirements for data processing, affecting global standards. I. Relekar [18] notes that the updated NIST Cybersecurity Framework 2.0 helps organizations adapt to new threats through a flexible approach to risk management.

D. Parsons [16] emphasizes the growth of attacks on industrial control systems, which requires strengthening the standards for protecting critical infrastructure. Microsoft's 2023 report [13] emphasizes the need to integrate artificial intelligence into cybersecurity standards to counter sophisticated attacks.

A. Folorunso, W. Mohammed, I. Wada, B. Samuel [8] prove that the implementation of ISO standards significantly increases the level of cybersecurity of organizations, providing a structured approach to information security.

The purpose of the article is to analyze the transformation of international cybersecurity standards (ISO/IEC 27001, NIST Cybersecurity Framework, GDPR) in response to technological (AI, IoT, 5G) and geopolitical (state-sponsored cyberattacks, regulatory fragmentation) challenges, as well as to identify key areas for their adaptation, assess the role of cross-sectoral cooperation, and develop recommendations for improving global cyber resilience.

Summary of the main material. Cyberspace is a key element of the modern world, where technological innovations and geopolitical tensions pose complex security challenges. The growth of cyberattacks, the proliferation of artificial intelligence (AI), the Internet of Things (IoT), 5G networks, and state-sponsored cyber operations require the adaptation of international cybersecurity standards. This transformation is aimed at responding to new threats, harmonizing regulatory requirements, and taking into account political realities.

Technological advances are radically changing the nature of cyber threats. The IBM Security Cost of a Data Breach 2024 report reports that the average global cost of a data breach reached USD 4.88 million, which is 10% more than in 2023. Phishing (16% of cases) and the use of stolen credentials (15%) remain the main causes of leaks. Artificial intelligence is opening up new opportunities for attacks, such as automated phishing campaigns and deepfakes, which make them more difficult to detect [6]. Statista predicts that by the end of 2025, the number of IoT devices will reach 75.44 billion, which creates additional vulnerabilities due to weak security and protocol heterogeneity [20]. In response, ISO/IEC 27001:2022 has been updated to include requirements for risk management in cloud, AI, and IoT. The standard emphasizes the integration of cybersecurity into software

development processes through DevSecOps and continuous risk monitoring to adapt to rapidly changing technologies [10; 7].

The fact that modern cyber threats are caused mainly by geopolitical challenges is confirmed by the fact that most cyber attacks on critical infrastructure in the EU showed signs of state support, in particular from Russia, China and North Korea [21]. The Microsoft Threat Intelligence Report 2023 indicates that Russia is responsible for 58% of state-sponsored attacks in the world. This prompts the development of standards that take into account political aspects [13]. The EU Cybersecurity Act of 2019 established a framework for the certification of products and services, promoting the unification of requirements in the EU [17]. However, global harmonization is complicated by differences between the approaches of the United States (NIST), China (GB/T standards) and other countries. For example, Chinese standards GB/T 22239 provide for data localization, which contradicts Western principles [9].

Regulatory changes are an important driver of transformation (Tab. 1). For example, the GDPR, introduced in 2018, set a global standard for data protection, influencing legislation, including the California Consumer Protection Act (CCPA) [10]. In 2023, the European Commission proposed the NIS2 Directive, which covers the energy, transport, healthcare, and digital services sectors, strengthening the requirements for incident response. ENISA estimates that NIS2 will affect 160 thousand organizations in the EU, which is twice as many as the previous NIS Directive [15]. The NIST Cybersecurity Framework 2.0, published in February 2024, introduces the “Govern” function for strategic cybersecurity management and recommendations for protecting supply chains and countering AI attacks [18].

Table 1

Key cybersecurity standards and frameworks

Standard	Organization	Year of update	Main changes
ISO/IEC 27001	ISO/IEC	2022	Integration of DevSecOps, AI risk management, cloud technologies, and IoT
NIST CSF 2.0	NIST	2024	Implementation of the “Govern” function, supply chain protection, countering AI attacks
GDPR	EU	2018	Establishment of global data protection standards, fines up to 4% of annual turnover
EU Cybersecurity Act	EU	2019	Certification of cybersecurity products, services, and processes
NIS2	EU	2023 (offer)	Expansion into new sectors, strengthening incident response requirements

Джерело: складено за даними [10; 17; 18]

It is the flexibility of standards that can be considered a key component for an effective response to challenges. NIST CSF 2.0 uses a modular approach, allowing organizations to adapt the framework to their needs [18]. ISO/IEC 27001:2022 supports integration with other standards, such as ISO 31000 (risk management), which facilitates implementation in different jurisdictions [2].

The main priority is to prepare for cyber incidents. For example, Cybersecurity Ventures noted that in 2021, ransomware attacks occurred every 11 seconds, causing \$20 billion in losses, which in turn will lead to attacks every second in 2031 [14]. ISO/IEC 27035:2023 provides guidance on coordination between organizations and government agencies during incidents [11]. The MITRE ATT&CK framework, which includes more than 200 attack techniques, is used to model threats. According to the SANS Institute, in 2023, 45% of organizations in the United States used MITRE ATT&CK to test systems [16].

It is also important to note that the human factor poses the greatest risks. The Verizon DBIR 2024 report indicates that 68% of data breaches are due to human error, including unintentional disclosure (43%) and weak passwords (25%) [1]. ISO/IEC 27002:2022 includes recommendations for two-factor authentication and staff training.

In general, effective cyber risk management requires a comprehensive approach that combines technological, organizational and human aspects, and takes into account geopolitical factors. To systematize the relationships between these elements and relevant standards, we have developed the diagram shown in (Fig. 1), which illustrates the integration of technological, geopolitical challenges and cybersecurity standards in a single conceptual model.

International cybersecurity standards continue to evolve to meet new realities. Updates to ISO/IEC 27001:2022 and NIST CSF 2.0 demonstrate progress in countering AI attacks, protecting IoT and supply chains [8]. GDPR and NIS2 set high standards for data protection and incident response, influencing global practices. However, regulatory fragmentation and political differences, especially between the West and China, make

harmonization difficult. Investments in cyber literacy and coordination between states and organizations are critical to improving global cyber resilience.

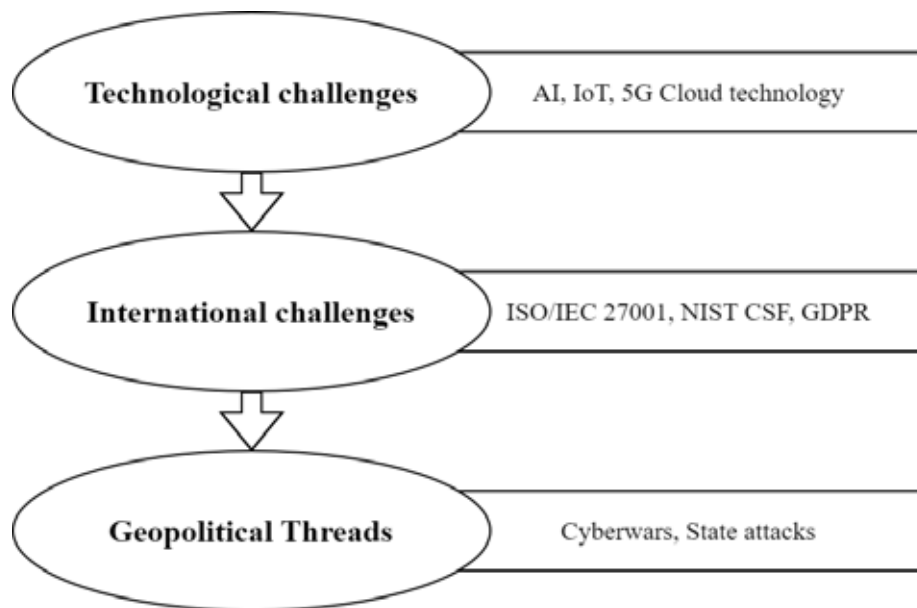


Fig. 1. Transformation of cybersecurity standards

Effective transformation of international cybersecurity standards is impossible without close cooperation between the public and private sectors, as well as international organizations. Most organizations around the world consider cross-sectoral partnerships to be critical to countering cyber threats. For example, the Cybersecurity Tech Accord initiative, signed by more than 150 technology companies, including Microsoft and Cisco, aims to jointly develop secure products and share information about threats [5]. Within the EU, ENISA coordinates with national authorities to facilitate the exchange of incident data through the MeliCERTes platform. In 2023, this platform processed more than 10,000 cyber incident reports, which is 25% more than in 2022, according to ENISA [19].

In addition, international organizations such as the ITU (International Telecommunication Union) and the OECD play a key role in harmonizing standards. For example, in 2023, the ITU published the X.1205 guidelines, which clarify approaches to cybersecurity in the IoT, facilitating their integration with ISO/IEC 27001. The OECD, in turn, updated its digital security guidelines in 2024, emphasizing the need for a global framework for managing risks in supply chains. However, cooperation is complicated by geopolitical barriers. For example, China's refusal to participate in international initiatives such as the Budapest Convention on Cybercrime limits global information sharing on cyber threats. According to Interpol, in 2023, only 40% of member countries regularly exchanged data on cyber incidents, which emphasizes the need for greater coordination [7].

Practical implications of cross-sectoral cooperation include the creation of joint cyber threat response centers. For example, in the United States, CISA (the Cybersecurity and Infrastructure Security Agency) launched the Joint Cyber Defense Collaborative in 2023, which brings together government, the private sector, and international partners to counter ransomware attacks [12]. In Europe, a similar role is played by the European Cyber Crisis Liaison Organization Network (CyCLONe), created under NIS2. These initiatives demonstrate that the transformation of standards must be accompanied by practical mechanisms for their implementation, including joint training, technology sharing, and policy harmonization.

To increase the effectiveness of international cybersecurity standards, a global harmonization framework based on cooperation through platforms such as the International Telecommunication Union should be developed to facilitate the harmonization of requirements across jurisdictions. Strengthening cyber literacy programs by introducing regular trainings for organizational staff is important, as only 30% of companies in the EU, according to ENISA 2023, conduct them systematically. Establishing international cyber threat response centers, following the example of CISA's Joint Cyber Defense Collaborative, will allow for the rapid exchange of information about incidents and coordination between states and the private sector [12]. The integration of artificial intelligence into cybersecurity standards, in particular to automate threat detection and analysis, is necessary to counter sophisticated attacks such as deepfake or automated phishing. Implementation of these measures requires joint efforts of governments, technology companies and international organizations to ensure the resilience of cyberspace in the face of modern challenges.

Conclusion. The transformation of international cybersecurity standards is a necessary response to technological innovation, geopolitical tensions, and growing cyber threats. Updates to ISO/IEC 27001:2022, NIST CSF 2.0, and proposals such as NIS2 are adapting regulatory frameworks to the challenges of AI, IoT, and state-sponsored attacks, while the GDPR remains the global benchmark for data protection. However, regulatory fragmentation and policy differences, particularly between Western and Chinese approaches, complicate harmonization. The human factor that causes most data breaches underscores the importance of cyber literacy.

Cross-sectoral and international cooperation, such as the Cybersecurity Tech Accord and ENISA platforms, is critical to implementing standards and building cyber resilience. Going forward, success will depend on flexible standards, increased coordination between governments, the private sector and international organizations, and investments in technology and education to ensure the resilience of global cyberspace.

Bibliography:

1. Arcila C., Pritam N., Kepple S. Data Breach Investigations Report: Vulnerability exploitation boom threatens cybersecurity. *Verizon*, 2024. URL: <https://www.verizon.com/about/news/2024-data-breach-investigations-report-vulnerability-exploitation-boom> (Accessed at: 15.05.2025).
2. Barafort B., Mesquida A. L., Mas A. Integrated risk management process assessment model for IT organizations based on ISO 31000 in an ISO multi-standards context. *Computer Standards & Interfaces*, 2018. № 60. C. 57–66.
3. California Consumer Privacy Act of 2018. *California Legislative Information – Website*. URL: https://leginfo.ca.gov/faces/codes_displayText.xhtml?division=3.&part=4.&lawCode=CIV&title=1.81.5 (Accessed at: 11.05.2025).
4. Chang L. Y. Legislative frameworks against cybercrime: The Budapest convention and Asia. *The Palgrave Handbook of International Cybercrime and Cyberdeviance*, 2020. PP. 327–343.
5. Ciglic K., Hering J. A multi-stakeholder foundation for peace in cyberspace. *Journal of Cyber Policy*, 2021. № 6(3). PP. 360–374.
6. Cost of a Data Breach Report 2024. *IBM – Website*. URL: <https://www.ibm.com/reports/data-breach> (Accessed at: 03.05.2025).
7. Edwards M. ISO 27001:2022 Annex A 5.23 – Information Security for Use of Cloud Services. *ISMS.online*, 2025. URL: <https://www.isms.online/iso-27001/annex-a/5-23-information-security-use-of-cloud-services-2022/> (Accessed at: 03.05.2025).
8. Folorunso A., Mohammed V., Wada I., Samuel B. The impact of ISO security standards on enhancing cybersecurity posture in organizations. *World Journal of Advanced Research and Reviews*, 2024. № 24(1). PP. 2582–2595.
9. GB/T 22239-2019 Information security technology–Baseline for classified protection of cybersecurity (English Version). *Code of China*, 2019. URL: <https://www.codeofchina.com/standard/GBT22239-2019.html> (Accessed at: 11.05.2025).
10. ISO 27001:2022 Controls: Annex A list. *Scrut Automation*, 2025. URL: <https://www.scrut.io/iso-27001/iso-27001-controls/> (Accessed at: 03.05.2025).
11. ISO/IEC 27035-1:2023. *ISO – Website*, 2023. URL: <https://www.iso.org/ru/standard/78973.html> (Accessed at: 15.05.2025).
12. Joint Cyber Defense Collaborative. *CISA – Website*. URL: <https://www.cisa.gov/topics/partnerships-and-collaboration/joint-cyber-defense-collaborative> (Accessed at: 19.05.2025).
13. Microsoft Digital Defense Report. Microsoft Threat Intelligence, 2023. 131 p.
14. Morgan S. Global Ransomware Damage Costs Predicted To Hit \$57B Annually In 2025. *Elastio*, 2025. URL: <https://elastio.com/research-report/2025-ransomware-report> (Accessed at: 11.05.2025).
15. NIS 2 strengthens cybersecurity across the EU by setting higher standards for essential services. *ENISA – Website*. URL: https://www.enisa.europa.eu/topics/state-of-cybersecurity-in-the-eu/cybersecurity-policies/nis-directive-2?utm_source=chatgpt.com#contentList (Accessed at: 11.05.2025).
16. Parsons D. Sans Survey Ics 2023. *SCRIBD*, 2023. 19 p. URL: <https://ru.scribd.com/document/678301429/Sans-Survey-Ics-2023> (Accessed at: 15.05.2025).
17. Regulation (EU) 2019/881 of the European Parliament and of the Council of 17 April 2019 on ENISA (the European Union Agency for Cybersecurity) and on information and communications technology cybersecurity certification and repealing Regulation (EU) No 526/2013 (Cybersecurity Act). 2019. PP. 15–69.
18. Relekar I. NIST Cybersecurity Framework 2.0. *ACA*, 2024. URL: <https://www.acaglobal.com/industry-insights/nist-cybersecurity-framework-20/> (Accessed at: 11.05.2025).
19. Ribeiro A. ENISA Threat Landscape 2024 identifies availability, ransomware, data attacks as key cybersecurity threats. *Industrial Cyber*, 2024. URL: <https://industrialcyber.co/reports/enisa-threat-landscape-2024-identifies-availability-ransomware-data-attacks-as-key-cybersecurity-threats/> (Accessed at: 15.05.2025).
20. Tanweer A. A Reliable Communication Framework and Its Use in Internet of Things (IoT). *ResearchGate*, 2018. № 3. URL: https://www.researchgate.net/publication/325645304_A_Reliable_Communication_Framework_and_Its_Use_in_Internet_of_Things_IoT (Accessed at: 03.05.2025).
21. Venkat A. Geopolitics plays major role in cyberattacks, says EU cybersecurity agency. *CSO*, 2022. URL: <https://www.csoononline.com/article/573999/geopolitics-plays-major-role-in-cyberattacks-says-eu-cybersecurity-agency.html> (Accessed at: 03.05.2025).

Дата надходження статті: 02.06.2025

Дата прийняття статті: 30.06.2025

Опубліковано: 23.09.2025

UDC 502.55:504.5:576.8:631.46

DOI <https://doi.org/10.32689/maup.it.2025.2.32>

Sofia SHVETS

Master of Science, Junior Research Scientist (Computer Systems),

Applied Programmer, Data Scientist, NinjaTech AI

ORCID: 0009-0000-7896-6479

INTEGRATION OF LARGE LANGUAGE MODELS INTO CHATBOT-BASED CUSTOMER CALL PROCESSING SYSTEMS

Abstract. Purpose of the work. To propose a hybrid model of automated text customer service that combines semantic classification with the generative capabilities of LLM to improve the accuracy, relevance, and naturalness of responses.

Methodology. A system has been developed that selects one of three response mechanisms depending on the classification of the intent and emotional tone of the query: template rules, search matching, or LLM generation. Experimental verification was performed on a corpus of 7,500 queries (authentic and synthetic); evaluation was conducted using BLEU, ROUGE-L, and expert criteria for comprehensibility, naturalness, and user trust.

Scientific novelty. For the first time, it has been shown that semantic routing in conjunction with LLM forms a more robust and adaptive system capable of correctly processing complex or emotionally charged queries. The proposed model outperformed baseline approaches in terms of accuracy (92.1%), BLEU 0.78, ROUGE-L 0.81, and the lowest failure rate, and also received the highest expert ratings.

Conclusions. The hybrid model reduces operator workload, increases user satisfaction, and easily adapts to customer behavior dynamics, providing empathetic and effective responses. Its practical value has been confirmed by examples from e-commerce, banking, healthcare, and public services. Implementation challenges include integration with legacy systems, regular knowledge base updates, and moderation of generated content. Further research is focused on personalization, multimodal interaction, active learning, and optimization of computing resources, laying the foundation for the development of advanced chatbots in areas with a critical need for high-quality automated support.

Key words: customer service processing, hybrid model, semantic classification, large language models, BLEU Score, ROUGE-L Score, chatbot, customer service optimization.

Софія ШВЕЦЬ. ІНТЕГРАЦІЯ ВЕЛИКИХ МОВНИХ МОДЕЛЕЙ У СИСТЕМИ ОБРОБКИ КЛІЄНТСЬКИХ ЗАПИТІВ НА ОСНОВІ ЧАТ-БОТІВ

Анотація. Мета роботи. Запропонувати гібридну модель автоматизованого текстового обслуговування клієнтів, що поєднує семантичну класифікацію з генеративними можливостями LLM для підвищення точності, релевантності та природності відповідей.

Методологія. Розроблено систему, що залежно від класифікації інтенції та емоційного тону запиту вибирає один із трьох механізмів відповіді: шаблонні правила, пошукове зіставлення або генерацію LLM. Експериментальна перевірка виконана на корпусі 7 500 запитів (автентичних і синтетичних); оцінка проведена за BLEU, ROUGE-L та експертними критеріями зрозумілості, природності й довіри користувачів.

Наукова новизна. Вперше показано, що семантична маршрутизація у зв'язці з LLM формує більш стійку й адаптивну систему, здатну коректно обробляти складні або емоційно тоновані звернення. Запропонована модель перевищила базові підходи за точністю (92,1 %), BLEU 0,78, ROUGE-L 0,81 та найменшою частотою збоїв, а також отримала найвищі експертні оцінки.

Висновки. Гібридна модель зменшує навантаження операторів, підвищує задоволеність користувачів і легко адаптується до динаміки поведінки клієнтів, пропонуючи одночасно емпатичні та дієві відповіді. Практичну цінність підтверджено на прикладах з електронної комерції, банківської сфери, охорони здоров'я та публічних сервісів. Виклики впровадження охоплюють інтеграцію з легасі-системами, регулярне оновлення баз знань і модерацію згенерованого контенту. Подальші дослідження спрямовано на персоналізацію, мультимодальну взаємодію, активне навчання та оптимізацію обчислювальних ресурсів, що закладає основу для розвитку передових чат-ботів у сферах з критичною потребою у високоякісній автоматизованій підтримці.

Ключові слова: обробка клієнтських запитів, гібридна модель, семантична класифікація, великі мовні моделі, BLEU, ROUGE-L, чат-бот, оптимізація обслуговування клієнтів.

Problem statement. Most customers face the chatbots' inability in support services to understand atypical or emotional requests, which increases the number of unresolved requests and worsens the user experience. Traditional rule-based or scripted solutions are losing relevance because of limited flexibility, the inability to adequately respond to atypical or emotionally coloured requests, as well as the difficulty in scaling multilingual support. These limitations reduce customer satisfaction and create additional workload for agents in cases where automated systems cannot cope.

© S. Shvets, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

LLMs, such as GPT-4, show potential in solving these problems because of their ability to generate natural and contextually-based responses. However, integrating LLMs into support systems has a number of challenges. These include choosing the right response mode (LLM, rule-based, retrieval) and building an architecture that ensures speed, accuracy, and secure generation. Furthermore, most current studies focus either on the isolated use of LLM or on classic chatbots without a hybrid approach. This creates space for the development of flexible, adaptive architectures that combine the strengths of different approaches.

Review of recent studies and publications. Recently, a significant amount of research has been devoted to the integration of large language models (LLMs) into customer service systems. These models show the potential to improve the quality of dialogues, personalization and efficiency of query processing. For example, Xu et al. in their study [1] presented a new customer service method that combines search-augmented generation with a knowledge graph. This method was applied in the LinkedIn customer service team for six months and reduced the average time to resolve the problem by 28.6%.

Jo & Seo [3] analysed the role of LLMs in modelling the emotional tone of the response, focusing on emotional aspects, but not covering technical or informational requests. Wulf & Meierhofer [11] studied automatic text correction, customer query generalization, and question answering using LLMs. Rüdél & Leidner [10] considered the integration of rule-based systems with neural chatbots, emphasizing the importance of controlling and avoiding “hallucinations” of the models. Kruk et al. [5] presented BanglAssist, a multilingual customer service chatbot capable of handling code-switching and dialect variations, which emphasizes the need for systems to adapt to linguistic diversity.

Hong et al. [2] proposed the Similar Question Generation (SQG) approach based on LLMs. It enables creating a significant number of diverse questions while maintaining semantic consistency with the original question-answer (QA) pair. This is achieved by using the natural language understanding (NLU) capabilities of the master through fine-tuning using specially designed prompts.

Practical implementations also demonstrate the effectiveness of LLM in support services. For example, Klarna implemented AI agents that perform the work of 700 support staff, which indicates significant automation potential. Shopify also used LLM to process basic queries, reducing the workload on operators and improving the speed of responses. In 2024, Duolingo integrated GPT-4 into its platform to improve the user experience. In particular, the Role Play and Explain My Answer functions enable users to have dialogues with AI characters and receive detailed explanations for their answers. This brought the platform closer to the level of a personal tutor. As Bicknell, the project manager, notes: “We’ve really come close to our ideal of being a personal tutor for every user” [7].

Despite these achievements, there are gaps in research, including:

- Lack of adaptive query processing strategies that take into account the type, complexity, and emotional context of the request.
- Limited attention to multilingualism and dialect variation processing.
- Insufficient implementation of hybrid architectures that combine rule-based, retrieval, and LLM approaches with dynamic switching between them.

This creates space for the development of a new flexible, adaptive chatbot architecture that combines the strengths of different approaches, provides multilingual support, and takes into account the emotional context of requests. This is what this study deals with.

Problem statement. The aim of this study is to develop and experimentally test a hybrid chatbot architecture. It integrates LLMs into the client request processing system, taking into account the type, complexity, and emotional colouring of requests.

Hybrid chatbot architecture. LLMs [8] are a modern advancement in Natural Language Processing (NLP). They are an important component in creating intelligent chatbots, virtual assistance systems, and automated customer support services.

A language model is a statistical or neural system that studies probabilistic patterns of sequences of words or tokens in text. It predicts the next word or generates a text response based on these patterns. The main architectural basis of most modern LLMs are Transformers [9], a type of neural network. They enable efficient processing of large amounts of text data thanks to the Self-Attention mechanism, which ensures that dependencies between all words in the input sequence are taken into account, regardless of their position.

LLMs are trained by supervised learning or self-supervised learning. In the case of self-supervised learning, the model is tasked with predicting hidden or next elements in text data without explicitly labelling examples. The process involves processing large text corpora that can contain millions of tokens. Intermediate steps teach the model to recognize language structure, logical connections between sentences, contextual dependencies, and even elements of general erudition.

Typical tasks for pre-training are word masking tasks (Masked Language Modelling, MLM) or text autogeneration (Causal Language Modelling, CLM). Despite their results in natural language generation, LLMs have the following limitations [10–11]:

- hallucinations, i.e. fictional facts or generation of incorrect information;
- difficulties in domain adaptation – LLMs may not cope well with specific topics without additional training;
- high computational costs both at the training stage and during model deployment.

The purpose of integrating chatbots with integrated LLMs into customer service systems is to provide automated support with minimal human intervention. However, the specifics of real user requests – including short phrases, unstructured queries, a large number of errors or slang – pose serious challenges to the models. Besides, in a business environment, it is important to ensure the controllability of responses, maintain service quality standards and ensure that responses comply with corporate policies.

The application of LLMs for chatbots covers a variety of industries, including e-commerce, healthcare, banking, education, government services, and technical support. At the same time, the integration of LLMs into such systems requires solving the problems of routing requests, filtering responses, a built-in fact validation mechanism, and optimizing the generation speed.

The creation of hybrid architectures is emphasized among the possible ways to improve the quality of LLM-based chatbots. They combine the capabilities of generative models with pre-processing of incoming requests, semantic classification and the use of specialized rules or knowledge bases. Such solutions increase the accuracy, relevance, and controllability of automated responses.

This study proposes a hybrid model of processing customer requests, which combines several approaches to analysing and generating responses in chatbots. The model is based on the integration of NLP methods, machine learning (ML), and current LLMs. This approach enables its adaptation to a wide range of requests – from typical to non-standard or complex. The model's function consists of several main stages (Figure 1):

1. Receiving a user request (Input): The user sends a text message in a natural language. The system accepts this message as input for further analysis.
2. Request classification (Intent Recognition): at this stage, the system classifies the user's intention. The NLU component – a software module that analyses the grammatical structure of the sentence, keywords and context – is used for this purpose.
3. Response strategy selection: After classifying the request, the system selects the optimal response strategy. This hybrid model provides three main mechanisms: – template responses (Rule-Based); – retrieval from the knowledge base (Retrieval-Based); – LLM-based generation (Generative LLM-Based).
4. Formation of the final response depending on the selected strategy (Response Generation).
5. Sending the generated response to the user in the chat (Output).

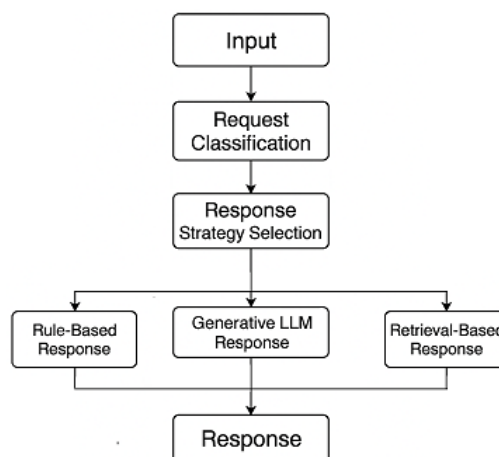


Fig. 1. Model operation diagram

Source: created by the author

This model has several advantages. The combination of three approaches enables its adaptation to different request types. The model can also serve thousands of users simultaneously, changing only the generative component as needed. LLM is involved only when other approaches do not provide sufficient quality of the

response, which helps to save resources. Besides, LLM creates more natural, polite, and contextually relevant responses.

The developed model was compared with the following models:

1. Rule-Based Model [12] This model operates on the basis of hard-coded rules and templates. It relies on predefined keywords, phrases and logical conditions, based on which the chatbot selects the answer.

Advantages: speed, ease of implementation, stability of the response.

Disadvantages: low flexibility, inability to handle complex or new requests, the need to manually update the rules. Usually used in small projects or at the initial stages of customer support automation.

2. Retrieval-Based Model [13] The model is based on the principle of finding the most relevant answer among the previously prepared ones. The user's request is converted into a vector representation (for example, via TF-IDF, FastText or Sentence-BERT), after which it is compared with the vectorized answers in the knowledge base. The most similar answer is returned to the user.

Advantages: good quality for repeated requests, maintaining control over the answers.

Disadvantages: lack of flexibility in formulating queries, poor adaptation to new situations.

This approach is widely used in FAQ systems, banking chatbots, technical support.

3. Generative LLM-based model (LLM-Only) In this model, LLM only is used to generate the answer, such as GPT-3.5, GPT-4 or similar. All logic, classification and formulation of the answer are delegated to the model, which, based on training on huge arrays of text, generates answers that are stylistically and emotionally close to human communication.

Advantages: high flexibility, naturalness of language, ability to solve non-standard situations.

Disadvantages: high consumption of computing resources, slowness in some cases, possible generation of inaccurate or unwanted answers.

Increasingly used in corporate solutions but requires control over the quality of the answer and validation of the content.

The proposed hybrid model combines the strengths of all three approaches: template response efficiency, extraction accuracy, and generation flexibility. Comparison with these models allows for a comprehensive assessment of response quality, performance, stability, and resource consumption.

Course of the experiment. The experiment was conducted on the basis of an internal simulation platform that models the support service of a hypothetical telecommunications provider. The database of requests was formed by combining real user requests – an anonymous history of support service requests (6,000 dialogues). This was supplemented by artificially generated requests created using generative models to cover rare situations, as well as a set of test scenarios developed manually by customer service experts (150 dialogues covering typical cases – password loss, service connection, complaints, billing issues, etc.). In total, the experiment involved a database of 7,500 requests covering 35 main request types.

Some of the requests were simulated through real users for obtaining a qualitative assessment. In particular, the study involved 50 volunteers, who played the role of users with predefined scenarios. This included unexpected wording of requests, language errors, sarcasm, etc. Ten support agents acted as a reference response, answering the same queries without the chatbot. The experiment lasted for two weeks. During this time, each model (Rule-Based, Retrieval-Based, LLM-Only, hybrid) answered the same set of queries.

Generative models were run on a cloud infrastructure with GPU (NVIDIA A100), template and extraction models on CPU servers, chatbots were integrated via API into a simulated support platform. The logging system provided a full record of dialogues and reactions. All requests were fixed in advance (i.e., not dynamically generated). Each model had the same time limit (3 seconds) for a response. LLM responses were additionally filtered for failures.

Methods for evaluating results. The evaluation was carried out using a combination of metrics and expert evaluation. Automatic metrics:

Accuracy – the accuracy of hitting the correct request template. Shows how well the system finds the correct answer immediately (without clarification).

$$Accuracy = \frac{N_r}{N} 100\%,$$

where N_r – the number of correct first answers, N – the total number of requests.

BLEU score – determines the coincidence of text fragments of a certain length (n-grams) between the generated and reference answer.

$$BLEU = BP \cdot \exp\left(\sum_{n=1}^N w_n \log(p_n)\right),$$

where BP – penalty for short answer:

$$BP = \begin{cases} 1, & \text{if } l \geq l_e \\ e^{\frac{1-l_e}{l}} & \text{else} \end{cases},$$

where l_e – length of the reference answer, l – length of the answer.

pn – accuracy of text fragments, $wn = 0,25$ – weight of each text fragment.

ROUGE-L Score – determines the length of the longest common substring between the answer and the reference answer:

$1 - pn - wn = 0.25 - \text{ROUGE-L Score}$ –

$$\text{ROUGE-L} = \frac{(1 + \beta^2) \cdot P \cdot R}{R + \beta^2 \cdot P},$$

where P – the proportion of words in the response that are present in the reference. It is defined as the ratio of the number of common words to the number of words in the response;

R – the proportion of words in the reference response that are present in the model response. It is defined as the ratio of the number of common words to the number of words in the reference response;

β – a coefficient that is the ratio of P to R .

Latency – the average response generation time:

$$\text{Latency} = \frac{T}{N},$$

where T – the sum of all response times.

Failure Rate – the proportion of requests for which the model did not provide any meaningful response:

$$\text{Failure Rate} = \frac{N_e}{N} 100\%,$$

where N_e – number of incorrect answers.

Expert evaluation was carried out on a scale from 1 to 5 according to the following criteria:

- quality of the answer (relevance, accuracy),
- clarity (simplicity of formulation),
- naturalness (is the answer like a human answer),
- trust in the answer (does the user believe that the information is correct).

The evaluation was carried out by 5 independent reviewers without technical training and 2 specialists in the field of customer support. Table 1 shows the values of the metrics for evaluating the results of the experiment. Figures 2–3 show histograms of automatic and expert indicators, respectively.

Table 1

Values of metrics for evaluating the models' performance

Metrics	Rule-Based	Retrieval-Based	LLM-Only	Hybrid model
Accuracy (%)	62.3	74.8	87.5	92.1
BLEU Score	0.45	0.61	0.72	0.78
ROUGE-L Score	0.51	0.65	0.76	0.81
Latency (c)	0.8	1.2	2.4	1.6
Failure Rate (%)	18.2	11.5	6.1	3.7
Answer quality (1–5)	3.2	3.9	4.4	4.7
Intelligibility (1–5)	3.1	3.8	4.5	4.8
Naturalness of language (1–5)	2.7	3.5	4.8	4.9
User trust (1–5)	2.9	3.6	4.3	4.6

Source: created by the author

The conducted experiment made it possible to obtain quantitative assessments of the quality of the hybrid model and three basic models of processing customer requests in chatbots. The hybrid model demonstrated the highest average BLEU Score among all participants in the experiment. This indicates a high similarity between the generated responses and reference texts by n-grams. In particular, the rule-based model had the

lowest BLEU, which was expected because of the limited number of predefined scenarios. According to the ROUGE-L indicator, the proposed hybrid model also outperformed other solutions, demonstrating a better ability to preserve the main content of reference responses.

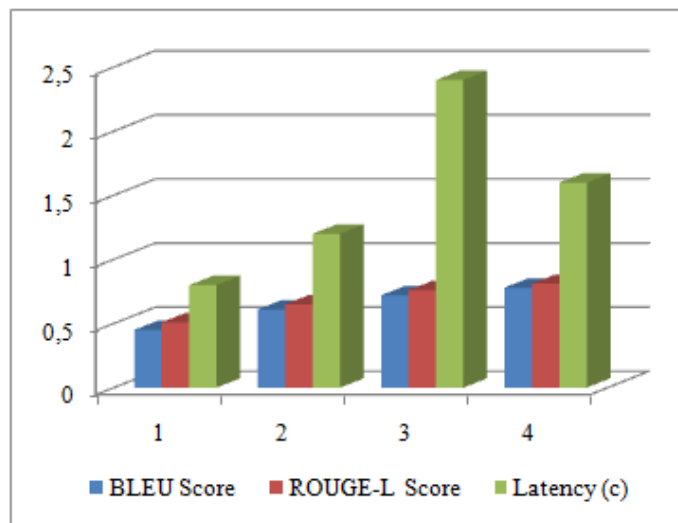


Fig. 2. Histogram of BLEU, ROUGE-L, Latency metrics for the studied models:
1 – Rule-Based, 2 – Retrieval-Based, 3 – LLM-Only, 4 – hybrid model

Source: created by the author

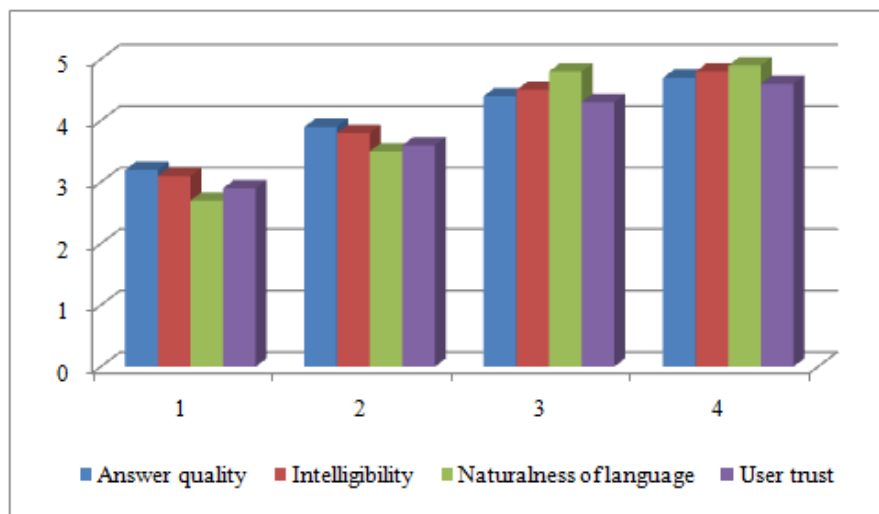


Fig. 3. Histogram of expert metrics for evaluating results for the studied models:
1 – Rule-Based, 2 – Retrieval-Based, 3 – LLM-Only, 4 – hybrid model

Source: created by the author

Fine-tuned LLM showed a medium result, while the rule-based approach significantly lost ground, probably because of the lack of adaptation to the specifics of the requests. The hybrid model received the highest scores for the Intelligibility criterion, which indicates the intuitiveness and ease of perception of its texts by users. Particularly large differences were observed between the hybrid model and the rule-based bot, which sometimes generated too formal or limited in content answers.

The naturalness of the hybrid model's answers was rated higher than that of other models, indicating a better imitation of the human communication style. The rule-based bot showed the lowest results because of rigid template constructions in the answers. The hybrid system ranked took first place in terms of the accuracy of the provided answers. This emphasizes its ability not only to create understandable and natural phrases, but also to adequately resolve users' information requests.

All results showed a low level of variance of estimates for the hybrid model, which indicates its stability on different subsets of requests. Integration of a specialized semantic pre-processor with generative capabilities of LLM significantly improves both formal metrics of correspondence (BLEU, ROUGE-L) and human impressions of interaction (Intelligibility, Naturalness).

Basic approaches have certain advantages in specific cases, but lose efficiency when the variety of user requests increases. The results of the conducted research have high practical significance for a wide range of applications in the field of automation of processing of client requests. The proposed hybrid model, which combines pre-processing of requests using semantic classifiers with the generative power of LLMs, demonstrates a significant improvement in the quality of interaction with users. An example of a request and responses is given below:

Example of user request 1:

"Good afternoon! I haven't received my bills in the mail for two months. What's the problem?"

Answer of the Rule-Based template system:

"You can download the bill in your personal account at the link..."

LLM answer:

"Sorry for the inconvenience. Try checking your spam folder or updating your e-mail in your personal account."

Hybrid model answer:

"There was probably a failure in sending bills. Check the e-mail in the user account or contact the operator. If the e-mail is correct, I will resend the bill."

Example of user request 2:

"I can no longer pay for the Internet at my tariff. It is now expensive for me. How can I suspend payment for the service?"

Answer of the Rule-Based template system: "Payment information is available in the Tariffs section of the website."

LLM answer:

"I understand that you are in a difficult situation right now. You may want to contact our support team to review your plan."

Hybrid model answer:

"We are working with a social support programme. Call 0-800-xxx or fill out a form to temporarily suspend payments."

Both examples demonstrate the advantage of the hybrid model – contextual sensitivity and a specific, actionable response that goes beyond a template approach. The rule-based system provides general information, LLM shows empathy, but does not always offer a practical solution. Instead, the hybrid model combines an emotional response with a real next step, which is important in a service environment.

The proposed hybrid model will contribute to:

- Increasing customer satisfaction due to high indicators of clarity, naturalness, and accuracy of answers. The system generates responses that are stylistically and emotionally close to human communication, which has a positive effect on the perception of the service by customers.
- Reducing the workload on contact centre operators. More accurate and relevant automatic answers significantly reduce the number of cases of escalation of appeals to people.
- Adaptability to new scenarios. Unlike classic rule-based systems, the hybrid model can quickly integrate new types of requests without significant costs for rewriting scripts.
- Reducing response time and increasing service efficiency, which directly affects the companies' business indicators.

The hybrid approach to processing appeals combines semantic classification with a generative language model. This is an improvement on existing architectures designed to achieve a better balance between the structuring and flexibility of answers. The proposed methodology for assessing the quality of the request processing system with a combination of automatic metrics (BLEU, ROUGE-L) and subjective human assessments (Intelligibility, Naturalness, Accuracy) provides a comprehensive approach to analysing effectiveness. Clarification of the characteristics of the BLEU and ROUGE metrics specifically for short dialogic responses, which was not always taken into account in similar studies before.

So, the study contributes both to the theory of LLM integration into application service systems and to practical methods for building more flexible and effective new generation chatbots. Thanks to its architecture, the proposed model can be applied in various areas, such as banking, e-commerce, technical support, health-care, government services, and others. It is effective in those industries where high-quality customer service is required 24/7. The conducted research opens up a number of promising areas for further improvement. For

example, the introduction of active learning methods, which will allow the system to independently identify the most problematic types of requests and request human assistance for clarification. This will increase accuracy without the need for massive manual data labelling.

The hybrid architecture can also be further expanded by personalizing responses for a specific user, using his history of requests and behavioural patterns. In particular, the next step may be to integrate the processing of not only text, but also voice, visual or structured customer requests, which will significantly expand the capabilities of the system. Besides, the analysis of the security of hybrid systems will be an important direction: studying their behaviour in response to aggressive, manipulative, or unethical user requests. It is also worth researching options for using less resource-intensive models or knowledge distillation methods to build compact, energy-efficient solutions.

Research limitations. Despite the positive results, the study has a number of limitations that should be taken into account when interpreting the obtained findings. The models were tested on a limited corpus of requests, which, although covering the main types of requests, do not fully reflect the full diversity of speech behaviour of real users. It should also be noted that LLM was run in a relatively stable, no-load environment during the experiment, while performance may be affected by external API delays or resource constraints in real-world applications. Furthermore, the hybrid model requires maintaining up-to-date data (contacts, policies, programme terms), which poses scalability and maintenance challenges. Assessing the clarity and relevance of responses also remains somewhat subjective.

Implementation prospects. The integration of LLMs into customer service systems opens up new opportunities for business and the public sector. In technical support services, LLMs can automatically respond to user requests, reducing the load on operators and response times. Online stores can use LLMs for personalized recommendations, order processing, and resolving customer issues in real time. Automating responses to typical citizen requests, such as processing documents or receiving social benefits, can increase the efficiency and accessibility of government services.

Implementation barriers. Despite significant advantages, there are certain challenges:

1. Integration with existing systems. Many organizations use legacy systems with which it is difficult to integrate modern LLMs.
2. The need for moderation. LLMs can generate incorrect or unacceptable answers, so a mechanism for checking and moderating content is required.
3. Updating knowledge bases. It is necessary to regularly update the databases from which LLM receives data to ensure that the answers are up-to-date.

Conclusions. A new hybrid architecture for automatic customer request processing was developed and tested in the study, which significantly improves the quality of interaction with the user compared to existing approaches. The proposed model demonstrates high results according to the selected metrics for assessing the quality of responses and confirms the effectiveness of combining semantic routing of requests with the capabilities of generative language models. The results of the study have significant practical value for implementation in various areas of business and open up new opportunities for the further development of intelligent customer service systems. The proposed approaches also form the basis for further research aimed at personalization, multimodal request processing, increasing the dialogue system security, and optimizing resource costs.

Bibliography:

1. Asai A., Min S., Zhong Z., Chen D. Retrieval-based language models and applications. Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 6: Tutorial Abstracts), July 2023. P. 41–46. URL: <https://aclanthology.org/2023.acl-tutorials.6/> (date of access: 1.05.2025).
2. Hong M., Song Y., Jiang D., Wang L., Guo Z., Zhang C. J. Expanding chatbot knowledge in customer service: Context-aware similar question generation using large language models. arXiv preprint arXiv:2410.12444. URL: <https://arxiv.org/abs/2410.12444> (date of access: 1.05.2025).
3. Jo S., Seo J. ProxyLLM: LLM-driven framework for customer support through text-style transfer. arXiv preprint arXiv:2412.09916. URL: <https://arxiv.org/abs/2412.09916> (date of access: 1.05.2025).
4. Kaddour J., Harris J., Mozes M., Bradley H., Raileanu R., McHardy R. Challenges and applications of large language models. arXiv preprint arXiv:2307.10169. URL: <https://arxiv.org/abs/2307.10169> (date of access: 1.05.2025).
5. Kruk F., Herath S., Choudhury P. BanglAssist: A Bengali-English generative AI chatbot for code-switching and dialect-handling in customer service. arXiv preprint arXiv:2503.22283. URL: <https://arxiv.org/abs/2503.22283> (date of access: 1.05.2025).
6. Li X., Gao M., Zhang Z., Yue C., Hu H. Rule-based data selection for large language models. arXiv preprint arXiv:2410.04715. URL: <https://arxiv.org/abs/2410.04715> (date of access: 1.05.2025).
7. Marr B. The amazing ways Duolingo is using AI and GPT-4. URL: https://bernardmarr.com/the-amazing-ways-duolingo-is-using-ai-and-gpt-4/?utm_source=chatgpt.com (date of access: 1.05.2025).

8. Naveed H., Khan A. U., Qiu S., Saqib M., Anwar S., Usman M. та ін. A comprehensive overview of large language models. arXiv preprint arXiv:2307.06435. URL: <https://arxiv.org/abs/2307.06435> (date of access: 1.05.2025).
9. Nerella S., Bandyopadhyay S., Zhang J., Contreras M., Siegel S., Bumin A. та ін. Transformers and large language models in healthcare: A review. *Artificial Intelligence in Medicine*, 2024. Vol. 106. Art. 102900. DOI: <https://doi.org/10.1016/j.artmed.2024.102900>.
10. Rüdell T., Leidner J. L. Control in hybrid chatbots. arXiv preprint arXiv:2311.11701. URL: <https://arxiv.org/abs/2311.11701> (date of access: 1.05.2025).
11. Wulf J., Meierhofer J. Utilizing large language models for automating technical customer support. arXiv preprint arXiv:2406.01407. URL: <https://arxiv.org/abs/2406.01407> (date of access: 1.05.2025).
12. Xu Z., Cruz M. J., Guevara M., Wang T., Deshpande M., Wang X., Li Z. Retrieval-augmented generation with knowledge graphs for customer service question answering. *Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval*, July 2024. P. 2905–2909. DOI: <https://doi.org/10.1145/3626772.3661370>.
13. Yao Y., Wang P., Tian B., Cheng S., Li Z., Deng S. та ін. Editing large language models: Problems, methods, and opportunities. arXiv preprint arXiv:2305.13172. URL: <https://arxiv.org/abs/2305.13172> (date of access: 1.05.2025).

Дата надходження статті: 23.06.2025

Дата прийняття статті: 04.07.2025

Опубліковано: 23.09.2025

НОТАТКИ

НАУКОВЕ ВИДАННЯ

**ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ
ТА СУСПІЛЬСТВО**

**INFORMATION TECHNOLOGY
AND SOCIETY**

**ВИПУСК 2 (17)
ISSUE 2 (17)**

2025

Коректура
Ірина Чудеснова

Комп'ютерна верстка
Наталія Кузнєцова

Формат 60x84/8. Гарнітура Cambria.
Папір офсет. Цифровий друк.
Підписано до друку 01.09.2025.
Ум. друк. арк. 33,95. Замов. № 0625/518. Наклад 300 прим.

Видавництво і друкарня – Видавничий дім «Гельветика»
65101, Україна, м. Одеса, вул. Інглєзі, 6/1
Телефон +38 (095) 934 48 28, +38 (097) 723 06 08
E-mail: mailbox@helvetica.ua
Свідоцтво суб'єкта видавничої справи
ДК No 7623 від 22.06.2022 р.