

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ

ХАРКІВСЬКИЙ НАЦІОНАЛЬНИЙ
УНІВЕРСИТЕТ РАДІОЕЛЕКТРОНІКИ

АВТОМАТИЗОВАНІ СИСТЕМИ УПРАВЛІННЯ ТА ПРИЛАДИ АВТОМАТИКИ

Всеукраїнський міжвідомчий
науково-технічний збірник

Заснований у 1965 р.

Випуск 186

Харків
2025

У збірнику наведено результати досліджень, що стосуються моделювання інфокомунікаційних систем і технологій, розробки моделей та методів формування пояснень в інтелектуальних системах, моделювання знань в системах підтримки прийняття рішень. Запропоновано нові та вдосконалені моделі, методи та інформаційні технології в галузі розробки та вдосконалення баз даних і спеціалізованих датасетів, створення програмно-апаратних систем виявлення камуфляжу та криптобезпеки.

Для викладачів університетів, науковців, фахівців, аспірантів.

The collection presents the results of research related to the modeling of infocommunication systems and technologies, the development of models and methods for generating explanations in intelligent systems, and the modeling of knowledge in decision support systems. New and improved models, methods, and information technologies are proposed in the field of developing and improving databases and specialized datasets, creating software and hardware systems for camouflage detection and cryptosecurity.

For university teachers, scientists, specialists, and graduate students.

Редакційна колегія:

І.В. Рубан, д-р техн. наук, проф. (гол. ред.), *В.М. Левикін*, д-р техн. наук, проф. (заст. гол. ред.), *М.В. Євланов*, д-р техн. наук, проф. (відпов. секр.), *В.В. Безкоровайний*, д-р техн. наук, проф., *М.О. Волк*, д-р техн. наук, проф., *І.В. Гребеннік*, д-р техн. наук, проф., *В.В. Євсєєв*, д-р техн. наук, проф., *А.Л. Єрохін*, д-р техн. наук, проф., *А.О. Каргін*, д-р техн. наук, проф., *І.Ш. Невлюдов*, д-р техн. наук, проф., *С.П. Новоселов*, д-р техн. наук, проф., *К.С. Смеляков*, д-р техн. наук, проф., *О.М. Цимбал*, д-р техн. наук, проф., *С.Г. Удовенко*, д-р техн. наук, проф., *О.Є. Федорович*, д-р техн. наук, проф., *В.О. Філатов*, д-р техн. наук, проф., *О.І. Філіпенко*, д-р техн. наук, проф.; *О.В. Чала*, д-р техн. наук, проф.

Затверджений до друку Науково-технічною Радою Харківського національного університету радіоелектроніки (протокол № 8 від 26 вересня 2025 р.).

За достовірність викладених фактів, цитат та інших відомостей відповідальність несе автор.

Унікальність текстів публікацій перевіряється за допомогою системи пошуку ознак плагіату StrikePlagiarism.

Рішення Національної ради про реєстрацію
Ідентифікатор медіа

№ 1410 від 25.04.2024 р.
R30-03874

Адреса редакційної колегії: Україна, 61166, Харків, просп. Науки, 14, Харківський національний університет радіоелектроніки, кімн. 254, тел. (057) 70-21-451

© Харківський національний університет
радіоелектроніки, 2025

ЗМІСТ

ПЕТРОВ К.Е., БОКОВ І.П., КОБЗЄВ І.В. РОЗРОБКА КОМБІНОВАНОГО МЕТОДУ АНАЛІЗУ ЕМОЦІЙНОЇ ЗАБАРВЛЕНOSTІ ТЕКСТІВ	5
МИХАЙЛІЧЕНКО І.В., ЛЯШЕНКО О.С. МОДЕЛЬ ПРОГНОЗУВАННЯ ВИКОРИСТАННЯ РЕСУРСІВ У ХМАРНИХ ОБЧИСЛЕННЯХ З ВИКОРИСТАННЯМ АРХІТЕКТУРИ INFORMER	17
ЗОЛОТУХІН О.В., КУДРЯВЦЕВА М.С., ФІЛАТОВ В.О. ПІДТРИМКА ОБМЕЖЕНЬ ЦІЛІСНОСТІ РЕЛЯЦІЙНОЇ БАЗИ ДАНИХ НА ЕТАПАХ ЕКСПЛУАТАЦІЇ ТА РЕІНЖИНІРИНГУ У ЗАДАЧАХ ІНТЕЛЕКТУАЛЬНОГО АНАЛІЗУ	29
SHTANGEY S.V., MELNIKOVA L.I., MARCHUK A.V., LYNNYK O.V., SOKOLOV O.K. MODELING AND ANALYSIS OF GRAPH NEURAL NETWORKS FOR OPTIMIZATION ROUTING IN INFOCOMMUNICATION NETWORKS	40
ПРОСОЛОВ В.В., ХАЛІМОВ Г.З., ШУЛІК П.В., СМІРНОВ А.О., В'ЮХІН Д.О. ВИКОРИСТАННЯ МЕТОДІВ МАШИННОГО НАВЧАННЯ ДЛЯ ВИЯВЛЕННЯ АТАК НА БЛОКЧЕЙН-СИСТЕМИ	55
ДЕМЕНКО Є.Є., ГРЕБЕННИК І.В., КОЛМИКОВ М.М. МЕТОД СТВОРЕННЯ ДАТАСЕТІВ ДЛЯ ОЦІНКИ АЛГОРИТМІВ РОЗПОДІЛУ ВАЛІДАТОРІВ НА ОСНОВІ МЕХАНІЗМУ PROOF OF STAKE	71
КРИЦЬКИЙ Д.М., ОНІЩУК Д.А., ВАЛЮЖЕНІЧ О.О. СИСТЕМА НА ОСНОВІ RGB-ДАТЧИКА ДЛЯ ВИЯВЛЕННЯ КАМУФЛЯЖУ У ВІЙСЬКОВОМУ СЕРЕДОВИЩІ	83
ЧАЛИЙ С.Ф., ЛЕЩИНСЬКА І.О. МЕТОД ОЦІНКИ НЕГАТИВНИХ АСПЕКТІВ РІШЕННЯ ІНТЕЛЕКТУАЛЬНОЇ СИСТЕМИ В ЗАДАЧАХ ПОБУДОВИ ПОЯСНЕНЬ	95
ЧАЛА О.В., БІТЧЕНКО О.М., САЙКІВСЬКА Л.Ф., АЛФЬОРОВ М.Є., ГАНШИН Д.Г. ПРЕСКРИПТИВНА МОДЕЛЬ ТЕМПОРАЛЬНИХ ЗНАНЬ З OWL-ІНТЕГРАЦІЄЮ ДЛЯ ПІДТРИМКИ РІШЕНЬ В ІНФОРМАЦІЙНИХ СИСТЕМАХ	103
РЕФЕРАТИ	113

CONTENT

PETROV K., BOKOV I., KOBZEV I. DEVISING A COMBINED METHOD FOR ANALYZING THE EMOTIONAL COLORING OF TEXTS	5
MYKHAILICHENKO I., LIASHENKO O. RESOURCE USAGE FORECASTING MODEL IN CLOUD COMPUTING USING INFORMER ARCHITECTURE	17
ZOLOTUKHIN O., KUDRYAVTSEVA M., FILATOV V. SUPPORTING RELATIONAL DATABASE INTEGRITY CONSTRAINTS DURING OPERATION AND REENGINEERING IN BUSINESS INTELLIGENCE TASKS	29
SHTANGEY S.V., MELNIKOVA L.I., MARCHUK A.V., LYNNYK O.V., SOKOLOV O.K. MODELING AND ANALYSIS OF GRAPH NEURAL NETWORKS FOR OPTIMIZATION ROUTING IN INFOCOMMUNICATION NETWORKS	40
PROSOLOV V., KHALIMOV H., SHULIK P., SMIRNOV A., VIUKHIN D. APPLICATION OF MACHINE LEARNING METHODS FOR DETECTING ATTACKS ON BLOCKCHAIN SYSTEMS	55
DEMENKO Y., GREBENNIK I., KOLMYKOV M. METHOD FOR CREATING DATASETS TO EVALUATE VALIDATOR ALLOCATION ALGORITHMS BASED ON THE PROOF OF STAKE MECHANISM	71
KRYTSKYI D., ONISHCHUK D., VALIUZHENYCH O. SYSTEM BASED ON RGB SENSOR FOR DETECTING CAMOUFLAGE IN MILITARY ENVIRONMENTS	83
CHALY S., LESHCHYNSKA I. METHOD FOR EVALUATING NEGATIVE ASPECTS OF INTELLIGENT SYSTEM SOLUTIONS IN EXPLANATION GENERATION TASKS	95
CHALA O., BITCHENKO O., SAYKIVSKA L., ALFEROV M., GANSHIN D. PRESCRIPTIVE MODEL OF TEMPORAL KNOWLEDGE WITH OWL INTEGRATION FOR DECISION SUPPORT IN INFORMATION SYSTEMS	103
ABSTRACTS.....	113

УДК 004.4'6

DOI: 10.30837/0135-1710.2025.186.005

*К.Е. ПЕТРОВ, І.П. БОКОВ, І.В. КОБЗЕВ***РОЗРОБКА КОМБІНОВАНОГО МЕТОДУ АНАЛІЗУ ЕМОЦІЙНОЇ ЗАБАРВЛЕНOSTІ ТЕКСТІВ**

Розглянуто психологічні й лінгвістичні основи визначення емоційної забарвленості природномовних текстів, різновиди класифікацій та роль емоцій у цифровій комунікації. Проведено порівняльний аналіз існуючих методів аналізу емоційної забарвленості текстів. Запропоновано комбінований метод, який базується на поєднанні двох векторних представлень тексту – статистичного (TF-IDF) та контекстуального (BERT). Таке поєднання дозволяє враховувати як частотні закономірності, так і глибокі семантичні залежності між словами, що в свою чергу, підвищує точність визначення емоційної забарвленості текстів. Наведено результати комп'ютерного моделювання, які демонструють працездатність та ефективність запропонованого методу.

1. Вступ

Визначення емоційної забарвленості тексту є однією з ключових задач обробки природної мови (Natural language processing, NLP) [1] та відіграє важливу роль у численних прикладних сферах, зокрема в маркетингу, соціології, психології, аналізі громадської думки та інформаційній безпеці. Суттєве зростання обсягу текстових даних у соціальних мережах, новинних ресурсах, блогах, коментарях та на форумах значно посилює потребу в автоматичному визначенні забарвленості текстів [1]. Системи аналізу емоційної забарвленості текстів дозволяють оперативно отримувати структуровану інформацію про настрої суспільства, прогнозувати реакцію на певні події, виявляти потенційні загрози або деструктивний контент, а також здійснювати моніторинг громадської думки в реальному масштабі часу.

Попри значні успіхи в цій сфері, існуючі методи аналізу забарвленості мають ряд обмежень, які знижують їхню ефективність. Зокрема, традиційні методи, що базуються на використанні словників чи статистичного аналізу, часто не враховують контекстуального значення слів, що є критично важливим для точного розпізнавання емоційної забарвленості текстів [2]. Крім того, багато методів стикаються з труднощами при аналізі багатозначних слів, сарказму, іронії, неформальних та сленгових виразів, які досить часто використовуються в сучасній комунікації. Ці нюанси роблять аналіз текстів складнішим завданням та потребують глибшого розуміння контексту й лінгвістичних особливостей мови. Крім того, тексти, що містять змішані емоції чи неоднозначні почуття, також є важкими для їх коректної класифікації за допомогою стандартних методів, які часто орієнтовані на чітко виражені емоції [1]. Для подолання цих проблем необхідно розвивати та застосовувати гібридні методи, які б ефективно поєднували переваги існуючих рішень.

Таким чином, існує потреба у подальшому вдосконаленні методів аналізу забарвленості тексту шляхом інтеграції різних підходів, зокрема з використанням сучасних моделей глибокого навчання, таких як нейронні мережі з механізмом уваги (attention mechanism), трансформерні моделі (наприклад, BERT, RoBERTa, GPT), які дозволяють краще враховувати контекстуальні та семантичні особливості текстів [3], [4]. Використання трансформерних моделей забезпечує глибше розуміння контексту через можливість моделювати семантичні взаємозв'язки між словами на різних відстанях у тексті. Водночас, поєднання нейронних мереж із механізмами уваги дозволяє акцентувати увагу моделі на найзначущіших елементах тексту, підвищуючи точність і чутливість аналізу [5]. Тому актуальною є розробка комбінованого методу визначення емоційної

забарвленості природномовних текстів [6], який би забезпечив підвищення точності визначення емоційної забарвленості за рахунок ефективного поєднання різних методів і передових моделей машинного навчання.

2. Аналіз існуючих методів аналізу емоційної забарвленості текстів та визначення проблеми дослідження

Виявлення емоцій у тексті є складним завданням, оскільки воно тісно пов'язане з особливостями контексту, багатозначністю мови, наявністю стилістичних засобів, таких як сарказм, іронія чи метафори. Через це для класифікації емоцій у текстах використовуються різноманітні методи.

Одним з найпоширеніших – є метод класифікації за полярністю емоцій, який передбачає поділ текстів на три основні категорії: позитивні, негативні та нейтральні.

Інший популярний метод базується на теорії базових емоцій Пола Екмана [7], в якій виділяються шість універсальних базових емоцій: радість (joy), смуток (sadness), гнів (anger), страх (fear), здивування (surprise) та огида (disgust)..

Подальшим розвитком концепції [6] стала модель «Колесо емоцій» («Wheel of Emotions») Роберта Плутчика [8]. Вона охоплює вісім основних емоцій: радість, довіру (trust), страх, здивування, сум, відразу, гнів та очікування (expectation). Ці емоції можуть комбінуватись між собою, формуючи складніші та тонші емоційні стани, наприклад, любов (love, поєднання радості та довіри).

Окремо можна виділити метод континуальної класифікації емоцій, в якому використовуються багатовимірні простори для опису емоційних станів. Однією з найпоширеніших є тривимірна модель, що описує емоції через три параметри [9]: валентність (ступінь позитивності чи негативності емоції), активацію (від стану спокою до інтенсивного збудження) і домінантність (ступінь контролю над ситуацією). Наприклад, щастя характеризується високою валентністю, високою активацією та значною домінантністю.

Виявлення емоцій у текстах базується на двох ключових аспектах: лінгвістичному та психологічному. Лінгвістичний аспект [10] передбачає аналіз мовних засобів, які використовуються авторами для вираження емоцій, зокрема лексики, морфології, синтаксису та стилістичних особливостей тексту. Психологічний аспект акцентує увагу на емоціях як когнітивно-емоційних станах [11], що відображають внутрішній досвід автора, проявляючись через специфічні експресивні засоби у тексті. Це передбачає моделювання емоційних станів [7], [8], [9] на основі різних психологічних теорій.

Сучасні методи аналізу емоцій у текстах намагаються інтегрувати лінгвістичні та психологічні аспекти, створюючи комплексні моделі, які здебільшого використовують машинне навчання та нейронні мережі, що дозволяє враховувати як лексичні особливості слів, так і контекстуальні та психологічні аспекти текстів. Такий інтегрований метод суттєво покращує точність автоматичного аналізу емоцій.

У галузі автоматичного аналізу емоцій у текстах використовуються різноманітні методи, які умовно можна поділити на лексиконні, статистичні, методи машинного навчання та гібридні. Кожен з них має свої переваги та обмеження.

Лексиконні методи до аналізу емоційної забарвленості тексту [1], [12] посідають ключове місце серед традиційних методів обробки природної мови. Вони базуються на гіпотезі, що емоційний стан автора можна визначити через використані ним мовні одиниці – переважно слова, що мають певне емоційне навантаження [13]. У центрі цих методів знаходяться так звані лексикони – попередньо сформовані спеціальні словники (наприклад, SentiWordNet, AFINN, NRC Emotion Lexicon), що містять слова, кожне з яких має заздалегідь визначену емоційну категорію або числову оцінку полярності (від сильно негативної до сильно позитивної) [1, 12]. Залежно від типу лексикону, емоційні характеристики можуть

бути представлені у вигляді бінарної шкали (позитивне або негативне), багаторівневої числової шкали (наприклад, від -5 до +5), або через категоріальну приналежність до базових емоцій за теоріями Плутчика [8], Екмана [7] та інших.

Серед лексиконних методів можна виділити:

- методи, що використовують сентиментні лексикони – спеціально сформовані переліки слів із визначеною емоційною характеристикою (позитивна, негативна чи нейтральна). Наприклад, до таких ресурсів належать WordNet-Affect, SentiWordNet, LIWC (Linguistic Inquiry and Word Count);

- методи, що базуються на використанні розширених лексиконів – баз даних, які окрім емоційних слів, включають також їхні синоніми, антоніми та оцінку інтенсивності емоцій. Одним із прикладів є VADER (Valence Aware Dictionary and sEntiment Reasoner), який добре підходить для аналізу текстів із соціальних мереж;

- методи на основі правил (Rule-based methods), які ґрунтуються на використанні наборів правил, що дозволяють враховувати контекст для посилення чи зменшення емоційного забарвлення фраз (наприклад, різниця між «дуже щасливий» і «трохи щасливий»).

Лексиконні методи аналізу емоційної забарвленості текстів відіграють важливу роль у задачах класифікації настроїв завдяки своїй простоті, обчислювальній ефективності та високій інтерпретованості результатів. Використання попередньо сформованих лексиконів дозволяє швидко оцінити емоційний зміст тексту без потреби в навчанні моделей на великих корпусах. Такі методи є особливо ефективними у випадках обробки коротких текстів, а також у системах, де необхідна прозорість прийнятих рішень.

Водночас обмеження лексиконних методів, зокрема нечутливість до контексту, труднощі з інтерпретацією заперечень, іронії та багатозначних слів, знижують їх точність у складних мовних конструкціях. Через це сучасні дослідницькі підходи дедалі частіше комбінують лексичний аналіз із контекстно-залежними моделями, використовуючи лексикони як модулі попередньої обробки або джерела інтерпретованих ознак. Таким чином, лексиконні методи зберігають свою актуальність у сучасному ландшафті обробки природної мови і використовуються як основа для гібридних систем емоційної класифікації.

На відміну від лексиконних методів, які ґрунтуються на заздалегідь сформованих словниках з емоційними маркерами, статистичні методи класифікації текстів використовують формальні математичні моделі для аналізу текстових даних, визначення статистичних закономірностей і подальшого віднесення тексту до певної емоційної категорії [14]. Ці методи ґрунтуються на ідеї, що текст можна представити у вигляді числових ознак, які можна обробити так само, як будь-які інші дані у задачах класифікації. Завдяки цьому емоційна забарвленість повідомлень виявляється шляхом виявлення кореляцій між статистичними характеристиками та цільовими етикетками (позитивний, негативний, нейтральний тощо).

Статистичний метод є особливо цінним у контексті автоматизації обробки великих обсягів інформації, де ручне анотоване кодування або використання фіксованих словників є малоефективним або взагалі неможливим [3]. Замість цього моделі навчаються розпізнавати шаблони у тексті на основі аналізу розподілу слів, частоти їх вживання та інших числових метрик. Це дозволяє виявити латентні (приховані) закономірності, які могли бути непомітними для людини або непридатними для лексичного аналізу. Наприклад, статистичні моделі здатні враховувати характерні поєднання слів або тенденції до використання певної лексики у різних емоційних контекстах, що особливо важливо в умовах неоднозначності та варіативності мови.

Статистичні методи ґрунтуються на аналізі частоти вживання слів, виявленні співвідношення емоційно забарвленої лексики та дослідженні словосполучень (колокацій),

які часто зустрічаються разом у тексті [14].

Найпопулярніші статистичні методи [2], [3] – Term Frequency (TF), Term Frequency–Inverse Document Frequency (TF-IDF), аналіз N-грам, Latent Semantic Analysis (LSA) – дозволяють ефективно обробляти великі корпуси текстів.

Кожен з цих методів має свої переваги й обмеження використання. TF-IDF дозволяє виокремити ключові слова, але не враховує порядок і контекст. Аналіз N-грам краще відображає локальний контекст, проте зі збільшенням N зростає обчислювальна складність. LSA виявляє приховані семантичні зв'язки, однак потребує великого обсягу даних.

Застосування статистичних методів є ефективним для аналізу великих корпусів тексту, але вони також не враховують глибокий контекст слів та їхній емоційний зміст.

Для подолання вказаних вище проблем аналізу емоційної забарвленості текстів активно використовуються методи машинного навчання та глибокі нейронні мережі.

Основні методи включають використання: класичних методів машинного навчання, векторного представлення слів (word embeddings), а також нейронних мереж, зокрема трансформерних моделей.

Методи класичного машинного навчання широко використовуються при розробці моделей, що здатні автоматично навчатися на основі розмічених текстових корпусів і прогнозувати емоційні мітки для нових текстів. Завдяки своїй обчислювальній ефективності, відносній простоті реалізації та високій швидкодії, ці методи широко застосовуються в практиці аналізу текстових даних. Найпоширенішими з них є моделі наївного байєсівського класифікатора (Naive Bayes), логістичної регресії (Logistic Regression), метод опорних векторів (SVM), ансамблевий метод випадкових лісів (Random Forest) та метод k найближчих сусідів (k-NN) [13], [15]. Кожен з цих методів має унікальні характеристики, переваги та сфери оптимального застосування, що дозволяє гнучко адаптувати їх до різних сценаріїв класифікації емоцій.

Останнім часом для інтелектуального аналізу текстів все частіше застосовують глибоке навчання, що базується на використанні, зокрема, рекурентних нейронних мереж (RNN), мереж довгої короткострокової пам'яті (LSTM), згорткових мереж (CNN) та трансформерів (BERT, GPT, RoBERTa тощо) [16]-[19]. Вони враховують контекст і семантичні зв'язки між словами, а також можуть інтерпретувати складні стилістичні прийоми, такі як іронія та сарказм.

Переваги використання глибокого навчання полягають у високій точності навіть у складних мовних випадках, здатності обробляти великі корпуси неструктурованих текстових даних та можливості донавчання на вузькоспеціалізованих доменах [12].

Здатність архітектур глибокого навчання адаптуватися до особливостей природної мови, враховувати контекст, розпізнавати приховані патерни та забезпечувати конкурентну точність дозволяє створювати системи, які перевершують можливості систем, що базуються на використанні класичних методів за результатами.

Разом з тим існують і недоліки, притаманні моделям глибокого навчання. Зокрема, їхня висока обчислювальна складність потребує наявності потужного апаратного забезпечення, що обмежує використання в реальному часі на малопотужних пристроях [12], [17]. Крім того, складність їхньої архітектури утруднює інтерпретацію результатів, що є критичним чинником у сферах, де потрібна прозорість прийняття рішень. Ще однією проблемою є потреба у великій кількості якісно анотованих текстових даних, без яких повноцінне навчання моделі стає неможливим.

Для підвищення точності аналізу в сучасних системах визначення емоційної забарвленості текстів доцільно використовувати комбіновані методи, що поєднують лексиконні методи, статистичний аналіз та глибокі нейронні мережі. Наприклад, спочатку може виконуватися

попередній аналіз тексту за допомогою лексиконних методів, а потім отримані дані подаються на вхід глибокої нейронної мережі для уточнення емоційної забарвленості.

Такі комбіновані методи дозволять враховувати як поверхневі лексичні ознаки, так і глибокі семантичні зв'язки у тексті. Застосування декількох методів у межах однієї системи сприятиме підвищенню загальної точності класифікації текстів. Тому вирішення проблеми адекватної класифікації природномовних текстів за їхньою емоційною забарвленістю на основі комбінації кількох методів є актуальним з теоретичної та прикладної точок зору.

3. Мета і задачі дослідження

Метою дослідження є підвищення точності класифікації емоційної забарвленості природномовних текстів за рахунок використання лексиконних, статистичних і контекстуальних методів, які дозволять врахувати як поверхневі лексичні ознаки, так і глибокі семантичні зв'язки у тексті.

Для досягнення поставленої мети необхідно вирішити такі основні задачі:

- розробити комбінований метод для вирішення задачі визначення емоційної забарвленості тексту;
- провести експериментальну перевірку працездатності та ефективності запропонованого методу.

4. Розробка комбінованого методу визначення емоційної забарвленості текстів

Комбінований метод визначення емоційної забарвленості текстів, що пропонується, ґрунтується на поєднанні лексиконних, статистичних і контекстуальних методів для досягнення максимальної точності при класифікації емоцій.

Основна ідея методу полягає в тому, щоб інтегрувати переваги класичних методів машинного навчання, які працюють з лексичними ознаками, із сучасними можливостями глибоких нейронних мереж, зокрема трансформерів, що здатні враховувати контекст і семантичні зв'язки між словами [4]. Архітектура методу передбачає чітко визначену послідовність етапів (рис. 1), кожен з яких відіграє ключову роль у забезпеченні високої якості результатів.



Рис. 1. Основні етапи комбінованого методу класифікації емоцій

Розглянемо основні етапи запропонованого методу детальніше.

На першому етапі здійснюється попередня обробка тексту. На цьому етапі Вхідний текст проходить стандартні процедури очищення: видалення зайвих символів, пунктуації, HTML-тегів, зниження регістру, а також лематизацію – приведення слів до їхньої базової форми. Здійснюється також видалення стоп-слів, які не несуть значущого семантичного навантаження. Цей крок сприяє зменшенню шуму в даних та покращенню якості подальших перетворень.

Другий етап – векторизація тексту, яка реалізується паралельно за двома напрямками. Перший напрямок – лексичне векторне представлення тексту, зокрема через TF-IDF.

Цей метод дозволяє відобразити статистичну значущість кожного терміну у тексті з урахуванням його частоти у документі та у всьому корпусі. У запропонованому методі TF-IDF використовується для побудови векторів фіксованої довжини (5000 найрелевантніших ознак), які передають частотну і тематичну інформацію про текст.

Другий напрямок – контекстуальне представлення тексту за допомогою трансформерної моделі BERT (Bidirectional Encoder Representations from Transformers). Замість використання BERT у режимі fine-tuning, обрано режим feature extraction: текст пропускається через попередньо натреновану модель bert-base-uncased і потім з нього видобувається векторне представлення, отримане шляхом усереднення CLS-токенів з останнього шару моделі. Це дозволяє отримати глибоко семантизоване уявлення про зміст тексту, що враховує контекстне вживання слів у реченні.

Третій етап передбачає інтеграцію ознак з обох представлень: TF-IDF та BERT-векторів. Отримані два векторні простори з'єднуються в єдиний гібридний вектор шляхом конкатенації. Це дозволяє комбінувати статистичні закономірності з семантичними зв'язками, що значно розширює інформаційне поле для класифікації текстів та знижує ймовірність втрати важливих ознак.

Четвертим етапом є класифікація текстів за їхньою емоційною забарвленістю з використанням моделі машинного навчання. У запропонованому методі використовується Random Forest Classifier – ансамблевий метод, що базується на ухваленні рішення за допомогою колективу дерев. Цей метод було обрано через його стійкість до перенавчання, здатність працювати з великими масивами ознак і високу інтерпретованість результатів. Для оптимізації продуктивності та зменшення ризику перенавчання, перед навчанням класифікатора здійснювалася стандартизація ознак [2]. Усі вектори приводилися до єдиного діапазону значень з використанням Z-score normalization. Далі здійснюється передбачення емоційної забарвленості текстів, що не входили до навчального набору. Вихідним результатом є прогноз емоційної категорії (позитивна, негативна або нейтральна) разом із ймовірнісною оцінкою, яка дає змогу враховувати рівень упевненості моделі у власному рішенні.

У структурному плані запропонований метод реалізовано як послідовність етапів, завдяки чому його можна легко адаптувати до вирішення нових завдань або замінити окремі компоненти, що використовуються на кожному окремому етапі, на досконаліші – наприклад, використати замість TF-IDF сучасніші методи векторизації, як-от BM25, або змінити базову трансформерну модель BERT на RoBERTa чи DistilBERT [3], [4]. Це забезпечує аналітичне підґрунтя для подальшої оптимізації або модифікації окремих складових методу, що є особливо важливим при масштабуванні системи або адаптації до нових наборів даних та доменно-специфічного контенту.

Для реалізації методу було обрано мову програмування Python, яка є найпоширенішою у сферах обробки природної мови і машинного навчання [17] та має розвинену екосистему бібліотек для обробки текстів, векторизації, роботи з моделями глибокого навчання та візуалізації результатів [20]. Зокрема, при розробці методу було використано такі бібліотеки [20]: Pandas і NumPy – для забезпечення попередньої обробки й агрегації даних; Scikit-learn – для реалізації ключових етапів векторизації, навчання моделі, масштабування ознак і побудови метрик оцінки ефективності; Transformers у поєднанні з PyTorch – для отримання та інтеграції трансформерного представлення тексту, зокрема BERT, що значно підвищило якість ознак за рахунок урахування контексту; Matplotlib та Seaborn – для візуалізації результатів та аналізу точності класифікації емоцій.

5. Експериментальна перевірка ефективності використання комбінованого методу

Для експериментальної перевірки працездатності та ефективності запропонованого методу було використано набір даних з емоційної класифікації англomовних твітів, що доступний на платформі Kaggle – Emotion Dataset for NLP [21]. Цей датасет містить короткі повідомлення з мітками емоцій, які відповідають певним емоційним станам автора. Набір даних включає понад 20000 рядків тексту, які містять емоційні мітки з таких категорій, як joy, anger, sadness, fear, love, surprise тощо. Для спрощення задачі та приведення її до трикласової моделі (позитивна, негативна, нейтральна), було здійснено агрегування класів: позитивні емоції – joy, love, surprise; негативні – anger, sadness, fear; нейтральні – невиражені або конфліктні випадки.

Балансування класів було здійснено шляхом підсумовування підвбірок, щоб уникнути домінування одного емоційного класу над іншими. Перед векторизацією всі тексти пройшли однакову обробку: зниження регістру, видалення пунктуації та спеціальних символів, токенизацію та видалення стоп-слів. Ці процедури дозволили уніфікувати структуру тексту, покращити якість векторизації та зменшити розмірність вхідного простору.

Як зазначено вище, комбінований метод реалізує концепцію двонаправленого опрацювання текстових даних, що передбачає паралельне використання двох незалежних, але комплементарних методів векторного представлення:

- TF-IDF – для обчислення було обрано 5000 найпоширеніших термінів. Використано TfidfVectorizer зі стандартними параметрами, включно з ngram_range=(1,2) та max_df=0.9, щоб відсіяти занадто часті слова. Результатом є розріджена матриця розміром $n \times 5000$;

- BERT (base-uncased): використано модель bert-base-uncased із Hugging Face, без fine-tuning. Отримання векторів здійснювалося через обчислення CLS-токену або середнього по всіх токенах на останньому прихованому шарі. Розмірність отриманого вектору – 768 ознак.

Обидва представлення об'єднувались у єдиний вектор ознак розмірності 5768, після чого здійснювалась стандартизація за допомогою StandardScaler.

Як класифікатор обрано Random Forest, оскільки він добре справляється з великою кількістю ознак, не потребує масштабування і стійкий до перенавчання. Було використано такі параметри моделі: n_estimators = 300; max_depth = 20; min_samples_leaf = 2; random_state = 42; n_jobs = -1 (для використання всіх CPU-ядер).

Модель навчалась на 80 % даних, інші 20 % використовувалися для тестування. Було також реалізовано п'ятикратну крос-валідацію для перевірки стабільності результатів.

Для оцінки ефективності моделі було використано стандартні метрики: accuracy – загальна точність класифікації (частка документів, за якими класифікатор прийняв правильне рішення, відносно до загальної кількості класифікованих документів); precision – точність в межах класу (частка документів, що істинно належать даному класу, відносно до загальної кількості документів які система віднесла до цього класу); recall – повнота (частка знайдених класифікатором документів, що належать класу, відносно до загальної кількості документів цього класу в тестовій вибірці); F1-score – гармонійне середнє між precision і recall [22]. Крім того, для візуалізації правильних і неправильних передбачень щодо віднесення тексту до відповідного класу в дослідженні використано Confusion Matrix (матриця похибок класифікації).

На рис. 2 представлено матрицю похибок, яка демонструє достатньо високу точність класифікації, з невеликим рівнем помилок між класами, що характерно для комбінованих моделей із сильними векторними поданнями представленнями.

Для оцінювання ефективності використання комбінованого методу було здійснено його порівняння з іншими поширеними методами, а саме TF-IDF + Random Forest (RF) та BERT + Random Forest (RF). Метою такого аналізу було визначення, наскільки запропонований

комбінований метод випереджає базові методи як за точністю, так і за стабільністю результатів.

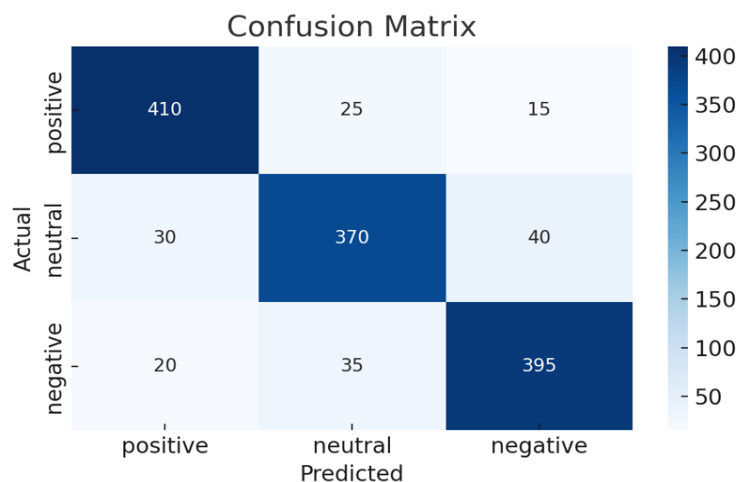


Рис. 2. Матриця похибок класифікації емоційної забарвленості текстів

На рис. 3 представлено діаграму значень точності різних моделей класифікації.

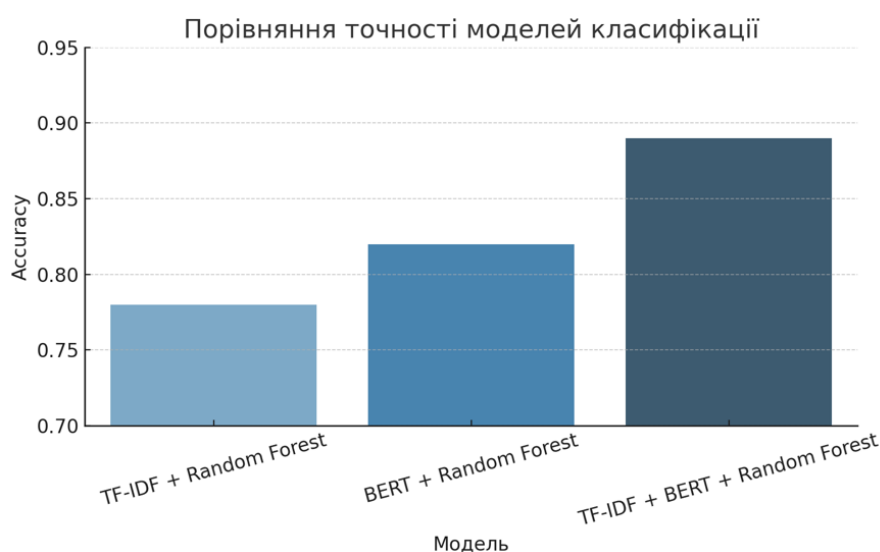


Рис. 3. Діаграма значень точності моделей класифікації

Значення F1-score представлено на рис. 4, який наочно ілюструє різницю в збалансованості точності та повноти між моделями, що базуються на різних комбінаціях ознак.

Оцінки часових витрат на тренування різних моделей, представлено на рис. 5.

З рис. 5 видно, що комбінований метод, який базується на використанні моделі TF-IDF + BERT + RF тренується найдовше, оскільки включає обидва типи векторизації та обробку великої кількості ознак.

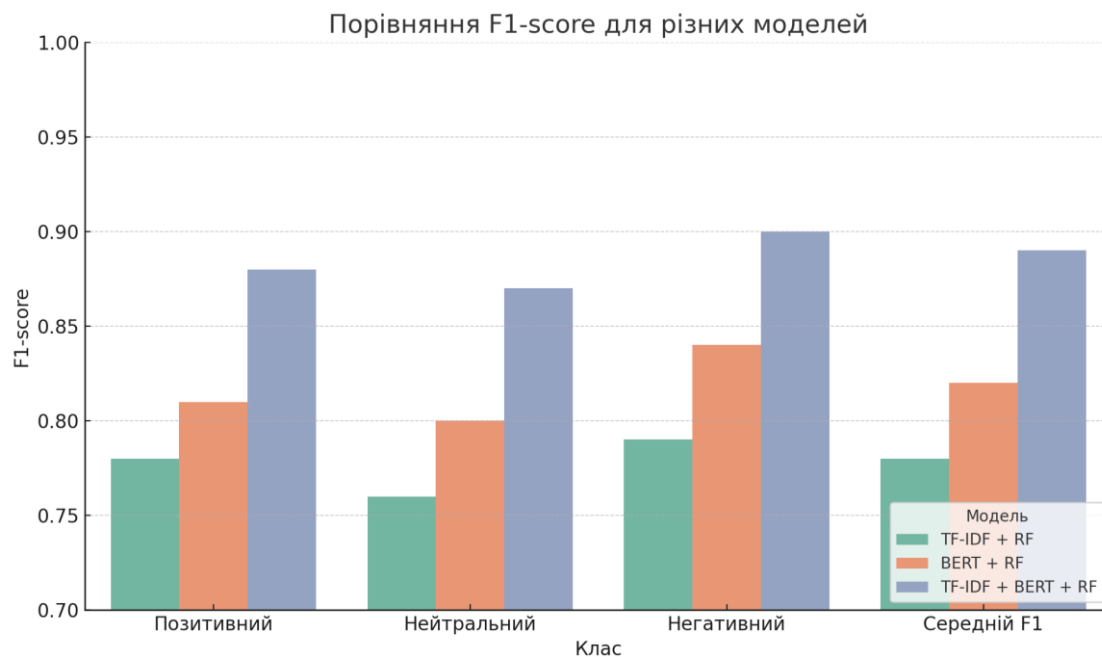


Рис. 4. Діаграма значень F1-score для різних моделей

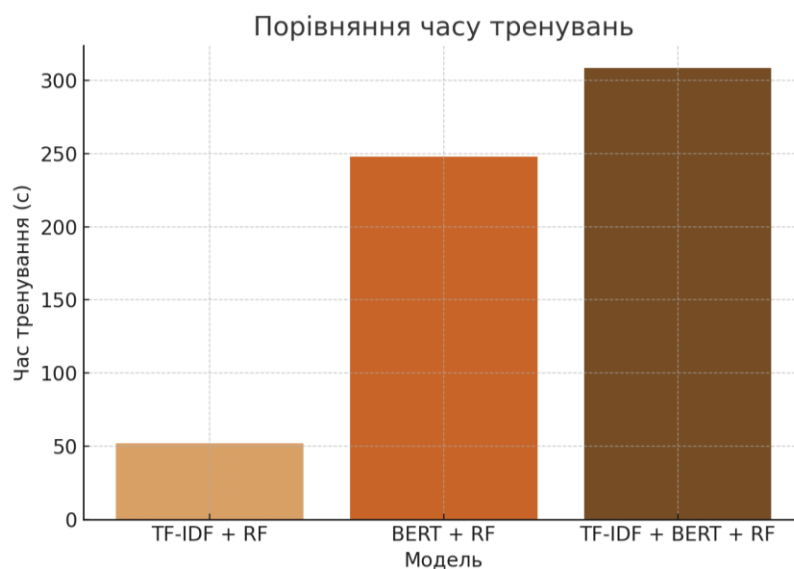


Рис. 5. Діаграма оцінок витрат часу тренування різних моделей

6. Обговорення результатів дослідження

Результати експериментальних досліджень підтвердили ефективність запропонованого комбінованого методу до визначення емоційної забарвленості текстів, що поєднує ознаки TF-IDF і контекстні вектори BERT.

З рис. 3 видно, що більшість прикладів правильно класифіковані, зокрема позитивні й негативні емоції мають найвищий рівень точності. Найбільше похибок спостерігається між класами «нейтральний» і «позитивний», що є типовим для коротких текстів, у яких емоційна забарвленість виражена неявно. Незначна кількість хибних спрацьовувань для негативного класу вказує на високу чутливість моделі до емоцій негативного спектра.

Загалом отримана матриця підтверджує ефективність комбінованого методу:

поєднання статистичних і контекстуальних ознак дозволило краще відрізнити тонко виражені емоційні стани, порівняно з класичними методами.

Порівняння з базовими методами показав, що саме запропонований метод демонструє найвищі значення асигасу та F1-score, перевершуючи як моделі на основі лише лексичних ознаках, так і моделі, що використовують лише BERT.

Комбінована модель TF-IDF + BERT + Random Forest забезпечує найвищий рівень точності (рис. 3), яка досягає значення близько 89 %. Це підтверджує ефективність інтеграції двох типів ознак – статистичних (TF-IDF) і контекстуальних (BERT) – у рамках одного векторного простору. Модель BERT + Random Forest показала середню точність – приблизно 82 %, що вказує на добру здатність BERT моделювати контекст, однак без доповнення частотними ознаками вона поступається гібридному методу. Найнижчу точність продемонструвала модель TF-IDF + Random Forest – близько 78 %, що свідчить про її обмежену здатність враховувати глибокі семантичні зв'язки в тексті.

Отримані значення F1-score (рис. 4) свідчать, що гібридна модель TF-IDF + BERT + RF демонструє найкращі результати для всіх трьох класів – позитивного, нейтрального та негативного. Середній F1 для цієї моделі становить близько 88 %, що помітно перевищує показники інших моделей. Модель BERT + RF посіла друге місце за ефективністю, особливо добре впоравшись із класифікацією негативних повідомлень. Її результати свідчать про високу цінність контекстуального представлення, хоча в сукупності ця модель поступається комбінованій моделі. Найгірші результати показала модель TF-IDF + RF, яка не змогла досягти високої точності в жодному з класів. Це ще раз підкреслює обмеженість використання виключно частотного методу в задачах, де важливе розуміння семантики.

Водночас, комбінована модель потребує більше часу на тренування, що було продемонстровано на відповідній діаграмі (рис. 5). Однак ці витрати компенсуються суттєвим зростанням якості результатів.

На основі аналізу результатів, отриманих в ході проведення експериментальних досліджень можна сформулювати низку рекомендацій щодо підвищення точності аналізу емоційної забарвленості природномовних текстів. По-перше, доцільно комбінувати різні типи ознак – частотні, контекстуальні та синтаксичні – для отримання повнішого представлення тексту. По-друге, варто проводити розширену попередню обробку, включно з врахуванням емодзі, синонімів, синонімічних перетворень і фразеологізмів, що є носіями прихованих емоцій. Крім того, покращення може бути досягнуто за рахунок тонкого налаштування трансформерної моделі під конкретний домен або використання ансамблевих методів, що поєднують прогнози декількох моделей. У сукупності ці заходи здатні забезпечити вищу чутливість моделі до нюансів емоційної виразності, особливо в коротких повідомленнях.

Подальші дослідження можуть бути зосереджені на оптимізації швидкодії моделі, зокрема шляхом використання полегшених варіантів BERT, таких як DistilBERT або TinyBERT, а також застосування методів квантування або дистиляції знань. Перспективним напрямом є також дослідження впливу різних схем агрегування ознак, включно з механізмами уваги (attention mechanism) [23] або autoencoder-об'єднанням векторів. У майбутньому доцільним є адаптація методу для багатомовного середовища та доменно-специфічного аналізу, що дозволить ширше застосовувати його для вирішення реальних завдань.

Крім того, доцільним є дослідження інтеграції зовнішніх знань, наприклад, у вигляді семантичних мереж або емоційних онтологій, які можуть підсилити якість інтерпретації тексту. Застосування багаторівневого ансамблю моделей може забезпечити гнучкість та адаптивність системи до різних джерел текстових даних. Перспективним є також включення модулів самооцінювання моделі, які дозволять оцінювати рівень впевненості у класифікаційних рішеннях. Це особливо важливо для критичних застосувань, де недостовірна

емоційна інтерпретація може призвести до хибних висновків. У напрямі адаптації метода до використання в системах реального часу доцільно розглянути гібридні архітектури, що поєднують глибоке навчання з класичними методами для досягнення компромісу між точністю та швидкістю. Важливо також досліджувати вплив різних стратегій попередньої обробки, зокрема нормалізації тексту, фільтрації шуму та виявлення іронії, які суттєво впливають на результати класифікації. Нарешті, інтеграція з інтерфейсами користувача (наприклад, візуалізація емоційної динаміки) може розширити функціональність систем та зробити їх зручнішими й інформативнішими для кінцевих користувачів.

7. Висновки

В рамках проведеного дослідження було розроблено комбінований метод визначення емоційної забарвленості текстів, що поєднує його статистичні й семантичні представлення для досягнення підвищеної точності класифікації емоцій.

Запропонований у роботі комбінований метод є прикладом гібридного методу до аналізу емоційної забарвленості тексту, в якому поєднуються два паралельні векторні представлення тексту: статистичне (TF-IDF) та контекстуальне (BERT). Таке поєднання дозволило враховувати як частотні закономірності, так і глибокі семантичні залежності між словами. У поєднанні з ансамблевим класифікатором Random Forest вдалося побудувати стійку модель, яка здатна ефективно класифікувати короткі англійські тексти з високим рівнем точності.

Для оцінки ефективності запропонованого методу було проведено його порівняння з базовими методами, що базуються на використанні моделей TF-IDF + RF та BERT + RF. Експерименти показали, що комбінована модель досягає точності класифікації до 89%, що перевищує результати альтернативних моделей (78 % і 82 % відповідно). Середній F1-score також виявився найвищим у комбінованій моделі – близько 88 %. Такий результат підтверджує доцільність інтеграції різних типів ознак. Було також розглянуто часові витрати на тренування моделей. Комбінована модель очікувано виявилася найресурсовитратнішою, що пов'язано з необхідністю одночасного обчислення двох типів векторів. Однак ці витрати повністю компенсуються зростанням точності та стабільності результатів.

Подальші дослідження можуть бути зосереджені на оптимізації архітектури моделі для скорочення часу її тренування, зокрема шляхом використання квантованих або дистильованих трансформерів. Крім того, гібридна архітектура може бути розширена на багатокласову класифікацію емоцій, що дозволить досягти детальнішого аналізу психологічного стану автора повідомлення. Не менш важливою є розробка explainable AI-механізмів для пояснення прийнятих рішень, що особливо актуально у таких чутливих сферах, як медицина, освіта та юриспруденція. Запропонований метод доцільно інтегрувати у веб-сервіси, CRM-платформи, чат-боти і аналітичні системи соціальних мереж та медіа, що забезпечить його практичну реалізацію.

Перелік посилань:

1. Goldberg Y. *Neural Network Methods in Natural Language Processing*. San Rafael: Morgan & Claypool, 2017. 287 p. URL: <https://doi.org/10.1007/978-3-031-02165-7>
2. Joshi D. Emotion Detection using Transformer Model with Deep Learning. *International Journal of Engineering Research and Technology*. 2025. Vol. 14 (3). P. 1–7 p. URL: <https://doi.org/10.17577/IJERTV14IS030116>
3. Acheampong F., Wenyu C., Nunoo-Mensah H. Text-based emotion detection: Advances, challenges, and opportunities. *Engineering Reports*. 2020. Vol. 2 (7). P. 1–24. URL: <https://doi.org/10.1002/eng2.12189>
4. Bing L. *Sentiment Analysis: Mining Opinions, Sentiments, and Emotions*. Cambridge : Cambridge University Press, 2020, 448 p.
5. Vanshika, Rani N., Walia R. A Comprehensive Review of Sentiment Analysis: Techniques, Datasets, Limitations, and Future Scope. *2024 Sixth International Conference on Computational Intelligence and*

Communication Technologies (CCICT), Sonapat, India, April 19-20, 2024. P. 403-409, URL: <https://doi.org/10.1109/CCICT62777.2024.00072>.

6. Боков І. П., Петров К. Е. Дослідження методів аналізу емоційного забарвлення текстів. *Радіoeлектроніка та молодь у ХХІ столітті: матеріали ХХVIII Міжнар. молодіж. форуму*, 16–18 квіт. 2025 р. Харків, 2025, Т. 6. С. 10–11.

7. Ekman P. Telling lies: Clues to deceit in the marketplace, politics, and marriage. W. W. Norton & Company, 2009. 400 p.

8. Emotion: Theory, Research and Experience. Volume 1. Theories of Emotion. Edited by R. Plutchik and H. Kellerman. Academic Press: London, 1980. 399 p. URL: <https://doi.org/10.1017/S0033291700053769>

9. Russell S., Norvig P. Artificial Intelligence: A Modern Approach. 4th ed. Boston : Pearson, 2020. 1136 p.

10. Jurafsky D., Martin J. Speech and Language Processing. 3rd ed. New Jersey : Pearson, 2008. 1024 p.

11. Sailunaz K., Alhajj R. Emotion and Sentiment Analysis from Twitter Text. *Journal of Computational Science*. 2019. Vol. 36. P. 437–448. URL: <https://doi.org/10.1016/j.jocs.2019.05.009>.

12. Dang N., Moreno-García M., De la Prieta F. Sentiment analysis based on deep learning: a comparative study. *Electronics*. 2020. Vol. 9(3). P. 1–29. URL: <https://doi.org/10.3390/electronics9030483>

13. Ghafoor Y., Jinping S., Calderon F., Huang Y., Chen K., Chen Y. TERMS: textual emotion recognition in multidimensional space. *Applied Intelligence*. 2023. Vol. 53(3). P. 2673–2693. URL: <https://doi.org/10.1007/s10489-022-03567-4>

14. Manning C., Schütze H. Foundations of Statistical Natural Language Processing. Cambridge: MIT Press, 1999. 718 p.

15. Minaee S., Kalchbrenner N., Cambria E., Nikzad N., Chenaghlu M., Gao J. Deep learning based text classification: A comprehensive review. *ACM Computing Surveys*. 2022. Vol. 54(3). P. 1–40. URL: <https://doi.org/10.1145/3439726>

16. Otter D., Medina J., Kalita J. A Survey of the Usages of Deep Learning for Natural Language Processing. *IEEE Transactions on Neural Networks and Learning Systems*. 2020. Vol. 32(2). P. 604–624. URL: <https://doi.org/10.1109/TNNLS.2020.2979670>

17. Chen S., Zhang Y., Yang Q. Multi-Task Learning in Natural Language Processing: An Overview. *ACM Computing Surveys*. 2024. Vol. 56(12). P. 1–32. URL: <https://doi.org/10.1145/3663363>

18. Pimpalkar P., Ingle G., Sonewane R., Lad V., Bangre R. Deep Learning for Sentiment Analysis: A Survey. *International Journal on Advanced Computer Theory and Engineering*. 2025. Vol. 14(1). P. 347–351. URL: <https://journals.mriindia.com/index.php/ijacte/article/view/479>

19. Петров К. Е., Воробйов Є. К., Кобзев І. В. Синтез моделі класифікації діалогових актів на основі використання рекурентних нейронних мереж. *АСУ та прилади автоматики*. 2022. Вип. 178. С. 4–12. URL: <https://doi.org/10.30837/0135-1710.2022.178.004>

20. Howard J., Gugger S. Deep Learning for Coders with Fastai and PyTorch. Sebastopol: O'Reilly Media, 2020. 624 p.

21. Emotions dataset for NLP classification tasks. <https://www.kaggle.com/datasets/praveengovi/emotions-dataset-for-nlp> (дата звернення: 05.05.2025).

22. Understanding Precision, Recall, and F1 Score Metrics. <https://medium.com/@piyushkashyap045/understanding-precision-recall-and-f1-score-metrics-ea219b908093> (дата звернення: 05.06.2025).

23. Yang C., Cao J. Interpretable Sentiment Analysis Using the Attention-Based Multiple Instance Classification Model: An Application to Wine Reviews. *Harvard Data Science Review*. 2025. Vol. 7(2). P. 1–11. URL: <https://doi.org/10.1162/99608f92.caab9466>

Надійшла до редколегії 13.07.2025 р.

Петров Костянтин Едуардович, доктор технічних наук, професор, завідувач кафедри ІУС ХНУРЕ, м. Харків, Україна, e-mail: kostiantyn.petrov@nure.ua, ORCID: <https://orcid.org/0000-0003-1973-711X>.

Боков Ігор Петрович, здобувач вищої освіти, група СШМ-23-2, факультет комп'ютерних наук ХНУРЕ, м. Харків, Україна, e-mail: ihor.bokov@nure.ua

Кобзев Ігор Володимирович, кандидат технічних наук, доцент, доцент кафедри мультимедійних систем і технологій ХНЕУ ім. Семе́на Кузне́ця, м. Харків, Україна, e-mail: ikobzev12@gmail.com, ORCID: <https://orcid.org/0000-0002-7182-5814>

УДК 681.5:004.89

DOI: 10.30837/0135-1710.2025.186.017

*І.В. МИХАЙЛІЧЕНКО, О.С. ЛЯШЕНКО***МОДЕЛЬ ПРОГНОЗУВАННЯ ВИКОРИСТАННЯ РЕСУРСІВ У ХМАРНИХ ОБЧИСЛЕННЯХ З ВИКОРИСТАННЯМ АРХІТЕКТУРИ INFORMER**

Представлено модифікований підхід до прогнозування та моніторингу навантаження хмарної інфраструктури з використанням нейронних мереж трансформерного типу. На основі підходу реалізовано модель з архітектурою типу Informer, адаптовану для багатовимірних часових рядів телеметрії. Модель протестовано на системі, яка реалізує замкнену архітектуру типу MAPE (Monitor – Analyze – Plan – Execute, Моніторинг – Аналіз – Планування – Виконання). Модель виконує багатопараметричне прогнозування динаміки навантаження на CPU, пам'ять, мережеві та дискові операції з урахуванням сезонних закономірностей. Проведено порівняльний аналіз полегшеної модифікованої версії з класичною архітектурою моделі Informer. Експериментальні дослідження показали, що полегшена версія забезпечує близьку до класичної точність прогнозування, при швидкості навчання в 8 разів більше та зменшеній кількості необхідних параметрів, що робить її придатною для розгортання в існуючих системах моніторингу та управління хмарними ресурсами.

1. Вступ

Хмарні платформи з мікросервісною архітектурою та оркестраторами на кшталт Kubernetes генерують великі обсяги багатовимірної телеметрії (CPU, пам'ять, дискові та мережеві метрики, показники сервісів і база даних). Класичні системи моніторингу, що спираються на порогові правила або спрощені статистичні моделі, погано масштабуються для подібних даних і не забезпечують своєчасного виявлення відхилень у стані інфраструктури [1]. Основна проблема дослідження полягає у побудові методів довгострокового прогнозування та раннього виявлення аномалій у багатовимірних нестационарних часових рядах хмарної телеметрії за умов обмежень на затримку обробки та обчислювальні ресурси.

Ці невирішені задачі роблять зазначену проблему особливо складною та відкритою. По-перше, суттєву роль відіграє нестационарність даних і так званий дрейф концепції: робочі навантаження змінюються внаслідок оновлень програмного забезпечення, міграцій сервісів, релізів нових функцій та сезонно-добових циклів, що призводить до нестабільності статистичних властивостей часових рядів. По-друге, дані мають високу розмірність і складні міжметричні залежності: кореляції можуть одночасно проявлятися між десятками метрик, охоплюючи як окремі сервіси, так і різні рівні архітектури – від прикладного до інфраструктурного. Додатково завдання ускладнюється наявністю багомасштабних закономірностей, коли короткочасні сплески поєднуються з довготривалими трендами та регулярними сезонними компонентами добового або тижневого характеру. Нарешті, у промислових сценаріях неминуче постає питання економічних обмежень, що вимагає пошуку балансу між якістю прогнозів та вартістю обчислювальних ресурсів і енергоспоживанням під час розгортання системи у виробничому середовищі.

Науковою задачею пов'язаною з розвитком трансформерних архітектур для довгих послідовностей, здатних моделювати довготривалі залежності у багатовимірних часових рядах при зниженій обчислювальній складності [1]. Серед актуальних напрямів варто відзначити використання сезонно-позиційних вбудовувань для явного кодування циклічних компонентів, застосування розріджених та ефективних механізмів уваги, інтеграцію спільного навчання прогнозних значень і сигналів аномалій, а також методи оцінювання невизначеності прогнозів для підтримки процесів прийняття рішень.

Практично успішне розв'язання цих задач забезпечує підвищення рівня SLA/SLO-дотримання завдяки проактивному масштабуванню, зменшення кількості хибних спрацьовувань і «хвилювання» автомасштабування, оптимізацію витрат на ресурси й енергоспоживання та підвищення стабільності роботи сервісів у пікові періоди.

Поєднання швидкозмінних робочих навантажень, високої розмірності телеметрії та обмежень на латентність обробки робить класичні підходи недостатніми. Необхідні ефективні трансформерні моделі, які одночасно підтримують довгі горизонти прогнозу, явну сезонність, спільне формування сигналів аномалій і придатність до розгортання в реальних MAPE-циклах (Monitor – Analyze – Plan – Execute, Моніторинг – Аналіз – Планування – Виконання).

2. Аналіз літературних джерел та визначення проблеми дослідження

Ознайомлення із сучасними науковими публікаціями дозволило встановити, що вирішенню проблем побудови методів довгострокового прогнозування та раннього виявлення аномалій у багатовимірних нестационарних часових рядах хмарної телеметрії присвячено значну кількість робіт. Так, у роботі [2] запропоновано використання трансформерних моделей для прогнозування багатовимірних часових рядів у складних технічних системах. Автори демонструють переваги механізму уваги для виявлення довгострокових залежностей та підвищення точності прогнозу. Проте відзначено значні обчислювальні витрати, що обмежує застосування моделей у сценаріях реального часу.

У публікації [3] розглянуто підхід до оптимізації багаточасових прогнозів для управління хмарними ресурсами. Показано, що поєднання кількох часових горизонтів дозволяє зменшити похибки при різних режимах навантаження. Разом із тим відсутня інтеграція механізмів виявлення аномалій, що знижує практичну цінність підходу для автономних систем управління.

У роботі [4] досліджено ресурсне планування в середовищі Kubernetes з використанням машинного навчання. Запропоновано схеми динамічного розподілу ресурсів залежно від прогнозованого навантаження. Недоліком є те, що аналіз виконувався на обмеженій кількості метрик без урахування складних міжметричних залежностей.

У дослідженні [5] наведено метод прогнозування завантаження кластерів Kubernetes за допомогою моделей машинного навчання. Розроблений підхід дозволяє зменшити кількість SLA-порушень за рахунок проактивного масштабування. Водночас використано прості моделі без урахування сезонності та багатомасштабних трендів, характерних для реальних систем.

Робота [6] пропонує багатоскальне моделювання часових рядів для виявлення аномалій у хмарних сервісах. Автори показують, що поєднання кількох часових масштабів дозволяє краще виявляти відхилення у поведінці системи. Недоліком підходу є його висока складність, що ускладнює інтеграцію в онлайн-системи моніторингу.

У статті [7] представлено модель DTAAD, що поєднує згорткові та трансформерні шари для виявлення аномалій у багатовимірних часових рядах. Модель демонструє високу точність, однак не передбачає спільного прогнозування та виявлення аномалій, що обмежує її використання в системах проактивного управління.

У дослідженні [8] запропоновано модифікацію трансформерної архітектури з використанням механізму нечіткої уваги для покращення прогнозування. Показано, що такий підхід дозволяє підвищити стійкість моделі до шуму в даних. Разом із тим не розглядаються обмеження на час обчислень, критичні для промислових застосувань.

У роботі [9] представлено глибоку трансформерну мережу TranAD для виявлення аномалій у багатовимірних часових рядах. Модель досягає високої точності в порівнянні з класичними методами, але потребує значних обчислювальних ресурсів, що ускладнює її використання у режимі реального часу.

Проведений аналіз показав, що сучасні підходи забезпечують високу точність прогнозування та виявлення аномалій, проте мають обмеження щодо обчислювальної ефективності, інтеграції сезонних компонентів і роботи з багатовимірними даними у режимі онлайн.

Невирішеною частиною загальної проблеми залишається розробка полегшеної трансформерної моделі, здатної одночасно виконувати прогнозування та виявлення аномалій у багатовимірних часових рядах із урахуванням сезонності та обмежень на обчислювальні ресурси. У даному дослідженні розглядається підхід, спрямований на усунення зазначених недоліків і забезпечення можливості інтеграції моделі в контур MAPE для управління хмарними ресурсами в режимі реального часу.

3. Мета і задачі дослідження

Метою дослідження є розробка та експериментальна оцінка моделі прогнозування використання ресурсів у хмарній інфраструктурі на базі архітектури Informer, модифікованої для багатовимірних часових рядів і придатної для інтеграції в замкнений контур MAPE з урахуванням ресурсних обмежень.

Поставлену мету заплановано досягти в результаті вирішення таких задач дослідження:

- розробка структури автономної системи моніторингу та оптимізації хмарної інфраструктури, розгорнутої в середовищі Kubernetes;
- модифікація архітектури Informer з урахуванням багатовимірності даних і фактору сезонності;
- проведення експериментального порівняльного аналізу класичної та модифікованої моделей Informer за точністю й ефективністю.

4. Матеріали і методи дослідження

Об'єктом дослідження є хмарна інфраструктура з мікросервісною архітектурою під керуванням оркестратора Kubernetes, що генерує багатовимірні часові ряди телеметрії (CPU, пам'ять, диск, мережа, затримки сервісів). Предмет дослідження – методи прогнозування навантаження та виявлення аномалій у таких часових рядах з використанням трансформерних моделей.

Головна гіпотеза – використання модифікованої архітектури Informer із сезонними та позиційними вбудовуваннями, а також двома вихідними «головами» (прогноз та аномалії) дозволить забезпечити баланс між точністю прогнозування та обчислювальною ефективністю, придатний для інтеграції у контур MAPE для проактивного управління ресурсами.

Дані, використовувані під час дослідження, охоплюють часовий період у 30 днів і включають ключові показники продуктивності (KPI) системи:

- CPU Utilization (%);
- Memory Utilization (MB);
- Disk I/O (MB/s)
- Network In/Out (MB);
- Task Count (кількість активних процесів);
- сезонні фактори часу (sin/cos кодування).

Загальний обсяг даних склав близько 9 млн. записів, що перетворює задачу прогнозування на аналіз багатовимірного високорозмірного часового ряду. Кожен запис представляє стан кластера в дискретний момент часу з періодом вибірки $\Delta t = 1$ сек.

Початкові дані містили як інформативні, так і константні ознаки. Для зниження розмірності та шуму було виконано:

- видалення константних ознак $Var(feature_i) = 0 \Rightarrow feature_i$, що призвело до

коротшання кількості ознак із 12 до 10, без втрати інформації;

- нормалізацію даних – для стабільності градієнтного спуску застосовано Z-score масштабування;

- формування часових вікон – для навчання моделі використано ковзне вікно розміром $T = 30$ кроків, що дозволило враховувати часові залежності.

Дані було розподілено у пропорції – 80 % на навчання, 10 % на валідацію та 10 % для тестування задля уникнення інформаційного витоку та наближення до умов реального прогнозування майбутніх станів системи.

Було реалізовано та протестовано дві моделі для вирішення задачі:

- модель на основі стандартної архітектури Informer із параметрами: $d_{\text{model}} = 256$, $L_{\text{enc}} = 4$, та 8 головами уваги, ProbSparse механізмом для вибору ключів у матриці уваги;

- запропонована модель на основі модифікованої полегшеної архітектури Optimized Informer з параметрами: $d_{\text{model}} = 128$, $L_{\text{enc}} = 2$, кількість параметрів – 2М. Було використано Mixed Precision Training для прискорення на GPU.

5. Результати розробки структури автономної системи моніторингу та оптимізації хмарної інфраструктури

Розроблена автономна система реалізує замкнений контур управління за моделлю MAPE (Monitor–Analyze–Plan–Execute) [10] для моніторингу та оптимізації хмарної інфраструктури, розгорнутої в середовищі Kubernetes. Структуру автономної системи показано на рис. 1.

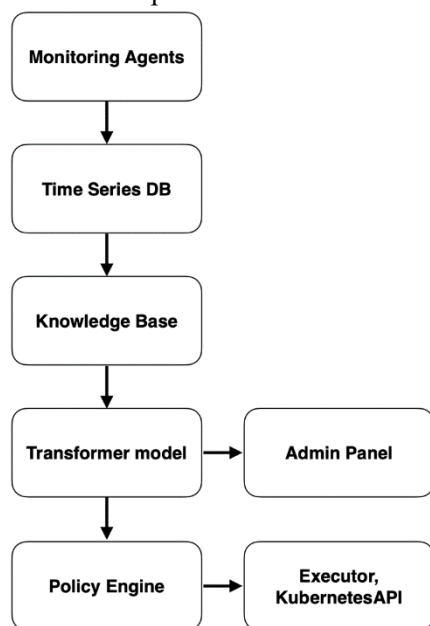


Рис. 1. Схема структури автономної системи моніторингу та оптимізації хмарної інфраструктури

Агенти моніторингу (Monitoring Agents) розгортаються на кожному вузлі кластера та в кожному контейнеризованому сервісі, збираючи телеметричні дані в режимі реального часу. Ці дані охоплюють ключові показники продуктивності: завантаження процесора (CPU), використання оперативної пам'яті (RAM), мережеву активність, операції з диском, затримки бази даних, частоту помилок API та інші метрики.

Усі метрики агрегуються та зберігаються в сховищі, що забезпечує ефективний доступ як до поточних, так і до історичних даних. Це сховище слугує джерелом даних для подальшого аналізу.

Центральний модуль інтелектуального аналізу побудований на основі модифікованої нейронної мережі трансформерного типу, адаптованої для обробки багатовимірних часових рядів. Модель отримує на вхід $X(t) \in R^n$ – (набір метрик на поточний момент часу), і прогнозує майбутні значення $X(t) \in R^n$, а також обчислює відхилення

$\delta X(t)$ від нормальної поведінки. Механізм self-attention дозволяє виявляти як короткострокові відхилення, так і довгострокові тренди та сезонні коливання [11].

На основі передбачених значень і виявлених аномалій модуль планування приймає рішення $u(t) \in U$ щодо необхідних дій: масштабування подів, перезапуск сервісів, перерозподіл навантаження. Планувальник реалізує або заздалегідь визначені правила, або навчувані евристичні. Модуль підтримує розширення політик на основі даних із бази знань.

Планувальник взаємодіє з API Kubernetes для виконання операцій масштабування, перезапуску контейнерів, перерозподілу ресурсів та інших дій в інфраструктурі [12]. Усі

дії, метрики до й після втручання, а також результати виконання зберігаються в базі знань. Ця інформація використовується для ретроспективного аналізу, навчання моделі (онлайн-навчання), побудови політик і пояснюваності прийнятих рішень.

Незважаючи на автономний характер системи, інтерфейс адміністратора забезпечує візуалізацію метрик, виявлених аномалій і рекомендованих дій. У такому режимі оператор може схвалити, відхилити або змінити оптимізаційні рішення, запропоновані системою.

6. Результати модифікації архітектури Informer

Запропонована в даному дослідженні модель реалізує інтелектуальний модуль моніторингу, розгорнутий в середовищі Kubernetes. Вхідними даними є багатовимірний ряд $X(t) \in R^n$, де компоненти $x_i(t)$ відповідають різним метрикам: завантаженню CPU, затримкам, трафіку тощо. Для кожного моменту часу t до вхідних векторів додаються два типи вбудовувань-представлень: позиційні вбудовування, що кодують відносну позицію в часовій послідовності, та сезонні вбудовування, які відображають періодичність вимірюваних метрик. Позиційні вбудовування реалізовані стандартним способом, а сезонні – як гармоніки відомих періодів. Сумарне вбудовування-представлення подається на розділену архітектуру кодер-декодер трансформера.

Кодер складається з N блоків самоуваги (self-attention), кожен з яких включає багатоголову самоувагу (multi-headed self-attention) і наступні шари позиційної нормалізації та FFN (feed-forward network) [13].

Декодер також містить кілька блоків: на етапі декодування використовується маскована багатоголова самоувага (masked multi-headed self-attention), яка враховує лише попередні моменти часу під час побудови прогнозу, а також механізм перехресної уваги (cross-attention) до виходу кодера. Специфічною модифікацією є використання ефективної самоуваги, що зменшує обчислювальну складність та зберігає домінуючі зв'язки, а також генеративного декодера, орієнтованого на довгострокове прогнозування послідовності.

Вихід моделі поділяється на дві частини: прогнозовані значення $\hat{y}(t)$ та оцінка аномалії $\delta(t)$. Загальну схему архітектури кодер-декодер наведено на рис. 2 [14].

Вхідні значення можуть бути представлені у вигляді часових рядів [15]:

$$X(t) = [x_1(t), x_2(t), \dots, x_n(t)]^T \in R^n, (1)$$

де n – кількість метрик. Додатково в модель вводяться позиційні вбудовування $P(t) \in R^d$, які, наприклад, реалізуються за допомогою синусоїдальних та косинусоїдальних функцій. Для каналів $j = 0, 1, \dots, \frac{d}{2} - 1$:

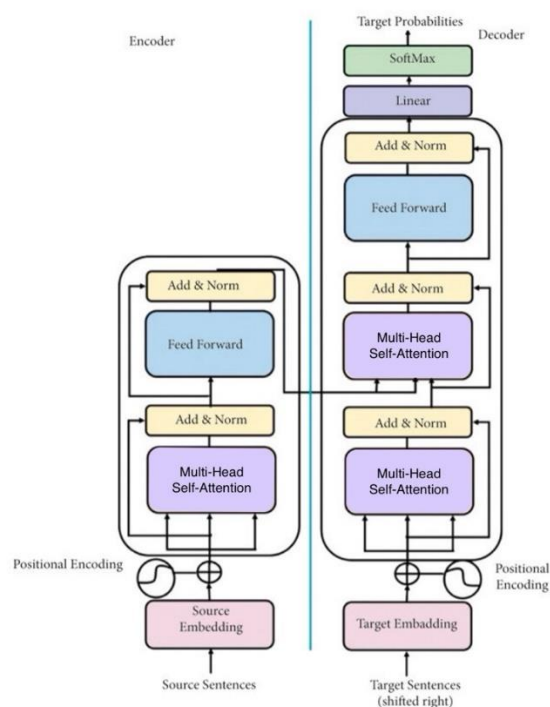


Рис.2. Архітектура моделі кодер-декодер

$$P_{2j}(t) = \sin\left(\frac{t}{C^{2j/d}}\right), P_{2j+1}(t) = \cos\left(\frac{t}{C^{2j/d}}\right), \quad (2)$$

де $C > 1$ – коефіцієнт масштабування. Сезонні вбудовування $S(t) \in R^k$ використовуються для кодування календарних циклів [16]. Для кожного відомого періоду T_p до вхідних значень додаються компоненти:

$$S_{2j}(t) = \sin\left(\frac{2\pi t}{T_p}\right), S_{2j+1}(t) = \cos\left(\frac{2\pi t}{T_p}\right). \quad (3)$$

Підсумковий вектор ознак для кожного моменту часу формується шляхом додавання або конкатенації вхідних значень із відповідними позиційними та сезонними вбудовуваннями:

$$Z(t) = X(t) + P(t) + S(t). \quad (4)$$

Отримана послідовність $Z(t)$ подається на вхід першого шару трансформера. Така форма представлення дозволяє моделі враховувати не лише абсолютні значення метрик, а й їхню позицію в часі, включаючи періодичні закономірності.

Модель навчається з використанням комбінованої функції втрат, яка одночасно враховує точність прогнозу та стійкість до сезонних коливань. Цільова функція мінімізується за такою формулою:

$$F = \frac{1}{T} \sum_t (|\hat{y}(t) - y(t)|^2 + \lambda |\hat{s}(t) - s(t)|^2). \quad (5)$$

Перший доданок відповідає середньоквадратичній похибці прогнозу (MSE), а другий – штрафу за відхилення у сезонній структурі, що дозволяє моделі залишатися стійкою до очікуваних циклічних коливань.

Для виявлення аномалій використовується надійна евристика, що ґрунтується на відхиленні фактичного значення від прогнозованого. Сигнал аномалії визначається таким чином:

$$\delta(t) = \begin{cases} 1, & |y(t) - \hat{y}(t)| > k \cdot \sigma_{resid}(t) \\ 0, & \text{інакше} \end{cases} \quad (6)$$

Модуль прогнозування реалізовано у двох варіантах:

- класичний *Informer* – трансформерна модель із багатоголовою самоувагою, ймовірнісною матрицею розрідженої уваги та шаром дистилляції, що знижує складність обчислень із $O(Y^2)$ до $O(Y * \log Y)$, де Y – довжина часової послідовності [17];

- модифікований *Informer* – запропонована в цій роботі модифікація з параметрами розмірного простору $d_{model} = 128$, $L_{enc} = 2$, яка зменшує кількість параметрів на 90% і прискорює навчання на GPU у 8 разів.

В модифікованій моделі додано сезонні вбудовування $S(t)$, що явно кодують добові та тижневі цикли в телеметрії Kubernetes. Створено також подвійний вихід моделі – окремі голови для прогнозу та оцінки аномалії, що дозволяє інтегрувати сигнал безпосередньо у MAPE-контур. Використовується комбінована функція втрат, яка одночасно мінімізує похибку прогнозу та відхилення у сезонних компонентах, що знижує кількість хибних спрацьовувань.

7. Результати проведення експериментального порівняльного аналізу класичної та модифікованої моделей *Informer*

Сформовано вимоги до моделі прогнозування з урахуванням виробничих обмежень: горизонт прогнозування не менше 200 кроків; латентність інференсу не більше 100 мс для батча з 32 послідовностей; споживання пам'яті не більше 2 ГБ на GPU або 1 ГБ на CPU. Сформовані цільові метрики якості:

- MAE (Mean Absolute Error) – середня абсолютна похибка прогнозу;
- RMSE (Root Mean Squared Error) – квадратична похибка;
- R^2 – коефіцієнт детермінації, який показує частку дисперсії даних;
- MAPE (Mean Absolute Percentage Error) – відносна похибка у відсотках;

Встановлено необхідність підтримки багатовимірного вводу щонайменше з десяти метрик та явного урахування добових і тижневих сезонних компонентів.

Розроблено трансформерну модель з позиційними та сезонними вбудовуваннями і двома вихідними головами: для прогнозування багатовимірного ряду та сигналу аномалії. Вхідне представлення формується як $Z(t) = X(t) + P(t) + S(t)$ і подається на кодер–декодер із ефективною самоувагою. Для навчання застосовано комбіновану функцію втрат, що поєднує похибку прогнозу та відхилення сезонної складової. Сигнал аномалії визначається порогоуванням за модулем залишку відносно ковзної оцінки стандартного відхилення. Окремо реалізовано дві конфігурації: базову $d_{\text{model}} = 256$, $L_{\text{enc}} = 4$, та 8 головами уваги, ProbSparse механізмом для вибору ключів у матриці уваги та модифіковану з параметрами: $d_{\text{model}} = 128$, $L_{\text{enc}} = 2$, кількість параметрів – 2М.

Проведено препроцесинг даних за 30 днів телеметрії з дискретизацією 1 с (приблизно 9 млн записів): видалено константні ознаки (зменшення з 12 до 10), виконано нормалізацію Z-score, сформовано ковзні вікна довжиною $T=30$. Розподіл вибірки становив 80 % для навчання, 10 % для валідації та 10 % для тестування, масштабування параметрів здійснювалося за навчальною підвибіркою. Отримано криві навчання для обох конфігурацій.

Для класичного Informer отримані середні показники – $\text{MAE}=5.2$, $\text{RMSE}=7.8$, $R^2=0.92$, $\text{MAPE}=6.5\%$. Для модифікованого Informer – $\text{MAE}=6.9$, $\text{RMSE}=10.4$, $R^2=0.88$, $\text{MAPE}=7.8\%$.

Виконано інтеграцію модуля прогнозування та індикації аномалій у контур MAPE прототипу керування ресурсами Kubernetes через відповідні виклики API. Реалізовано формування керувальних дій масштабування подів, перезапусків і перерозподілу ресурсів на підставі прогнозних значень і бінарного індикатора аномалій.

На рис. 3 наведено криві навчання двох моделей. Класичний Informer демонструє значно швидшу і глибшу збіжність: початкове значення втрат становило близько 0.38, і вже після другої епохи воно зменшилося удвічі, до 0.19.

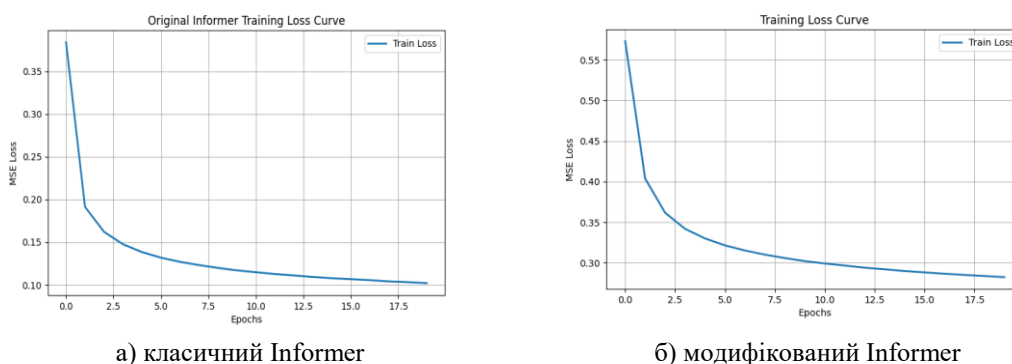


Рис. 3. Порівняння функції втрат

До завершення навчання, на 20-й епосі модифікована модель досягла втрат приблизно 0.10. Крива має чітко виражений експоненціальний спад, що вказує на стабільний процес навчання та високу здатність моделі виявляти і узагальнювати залежності в багатовимірних часових рядах.

На наведених на рис. 4 графіках показано якість прогнозування метрики CPU Utilization двома моделями: класичним Informer (графік а) та модифікованим Informer (графік б). На рис. 4 Metric_0 означає навантаження CPU. Обидві моделі порівнюються з фактичними значеннями навантаження на процесор у вибірці з 200 послідовних вимірювань.

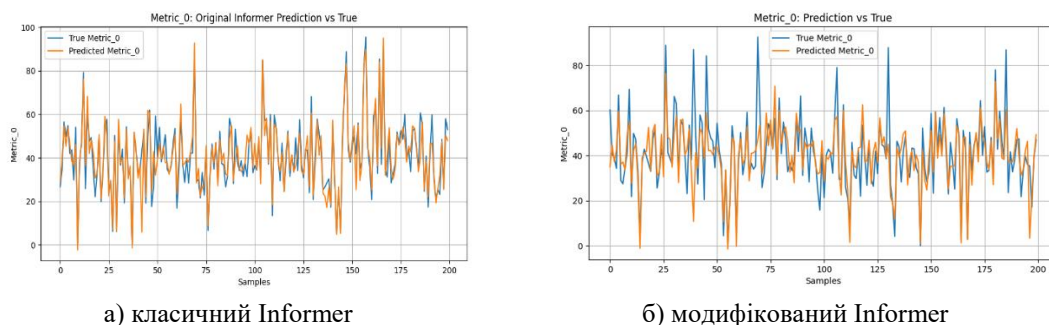


Рис. 4. Порівняння прогнозів CPU навантаження

Порівняння прогнозів навантаження пам'яті, розроблених із використанням класичного та модифікованого Informer, наведено на рис. 5. На рис. 5 Metric_1 означає навантаження пам'яті.

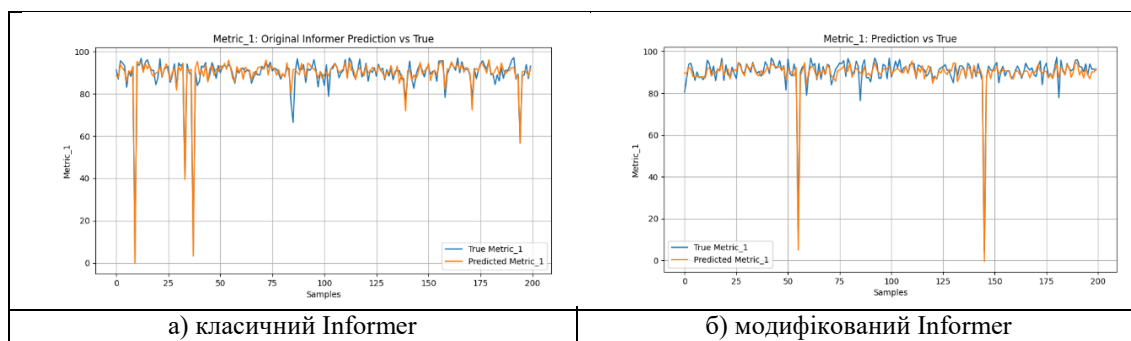


Рис. 5. Порівняння прогнозів навантаження пам'яті

Порівняння прогнозів часових ознак, розроблених із використанням класичного та модифікованого Informer, наведено на рис. 6. На рис. 6 Metric_9 означає прогнозування часових ознак.

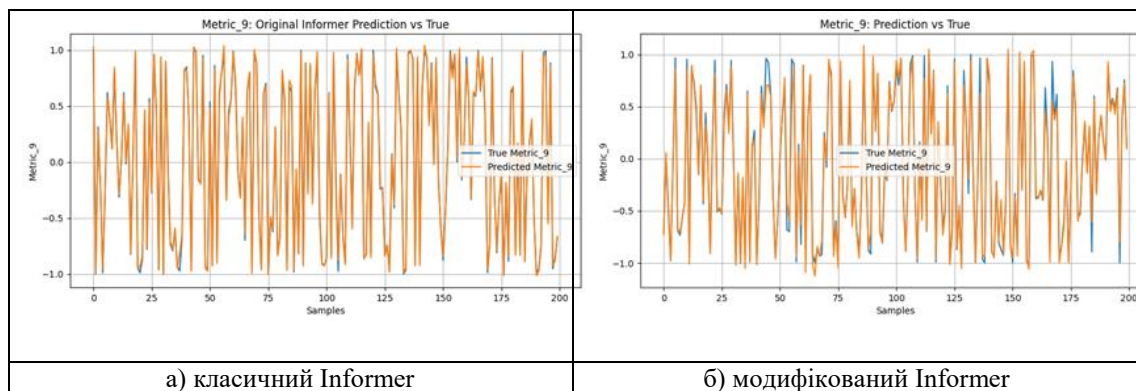


Рис. 6. Порівняння прогнозування часових ознак

Порівняння прогнозів вхідного трафіку, розроблених із використанням класичного та модифікованого Informer, наведено на рис. 7. На рис. 7 Metric_3 означає прогнозування вхідного трафіку.

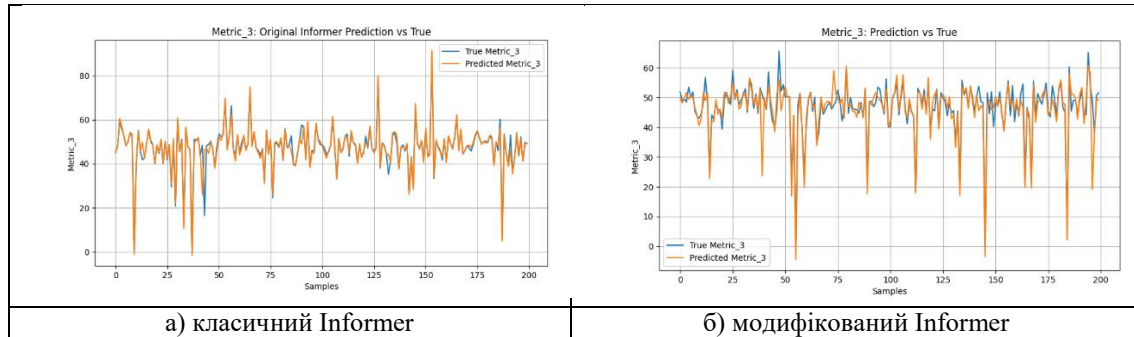


Рис. 7. Порівняння прогнозувань вхідного трафіку

8. Обговорення результатів дослідження

В дослідженні було запропоновано структуру автономної системи моніторингу та оптимізації хмарної інфраструктури. Для цієї системи було модифіковано архітектуру Informer, в якій:

- додано сезонні вбудовування $S(t)$, що явно кодують добові та тижневі цикли в телеметрії Kubernetes;
- створено окремі голови для прогнозу та оцінки аномалії, що дозволяє інтегрувати сигнал безпосередньо у MAPE-контур;
- використовується комбінована функція втрат, яка одночасно мінімізує похибку прогнозу та відхилення у сезонних компонентах, що знижує кількість хибних спрацьовувань.

Проведено експериментальний порівняльний аналіз класичної та модифікованої моделей Informer. Результати порівняння наведено на рис. 3 – рис. 7.

Результати порівняння функцій втрат (див. рис. 3) дають можливість зробити висновок, що модифікований Informer починає навчання з більшого значення втрат – близько 0.57 – і зменшує його повільніше. Попри це, модифікований варіант моделі має значну перевагу у швидкості навчання: одна епоха триває близько 2 хвилин проти 18 хвилин для класичної архітектури. Загалом це забезпечує приблизно восьмиразове прискорення та зниження споживання пам'яті у 10 разів.

Результати порівняння прогнозів CPU навантаження (див. рис. 4) дають можливість зробити висновок, що класичний Informer добре відслідковує як середній рівень завантаження CPU, так і локальні пікові значення. Модифікована модель демонструє здатність відтворювати основні тренди та загальний характер динаміки метрики, зберігаючи при цьому високу стабільність прогнозів. Завдяки згладженню короточасних флуктуацій, модель мінімізує ризик надмірного реагування на шум або незначні коливання.

Таке згладжування є важливою перевагою в умовах автономних систем моніторингу, де надмірна чутливість до піків може призводити до хибних спрацьовувань. Зменшення розміру моделі, кількості параметрів і часу навчання дозволяє застосовувати її у сценаріях з обмеженими ресурсами або у випадках, коли критичною є швидкість прийняття рішень. Таким чином, модифікований Informer поєднує ефективність обчислень із прийнятною точністю, що робить його придатним для реального часу та хмарних середовищ із

високою інтенсивністю телеметричних даних.

Результати порівняння прогнозів навантаження пам'яті (див. рис. 5) дають можливість зробити наступний висновок. Модифікована архітектура Informer демонструє здатність ефективно відтворювати загальний рівень використання пам'яті, зберігаючи основні сезонні та трендові характеристики метрики. Прогнозна крива, побудована за допомогою моделі, відображає ключові зміни в динаміці ресурсу, зокрема типові зниження та відновлення рівня завантаження. Завдяки меншій складності моделі та зменшенню кількості параметрів, результати залишаються стабільними, навіть за наявності незначних локальних коливань. Такий тип поведінки є цінним у випадках, коли система повинна швидко реагувати на загальні тенденції навантаження, не відволікаючись на незначні флуктуації.

Результати порівняння прогнозів часових ознак (див. рис. 6) дають можливість зробити наступний висновок. У випадку класичного Informer видно, що прогнозні значення практично повністю збігаються з фактичними даними. Помаранчева крива накладається на синю майже по всій довжині графіка. Модифікований Informer також демонструє високу точність для цієї метрики. Помаранчева лінія лише зрідка відхиляється від синьої, але загалом відтворює структуру даних із незначними похибками. Це свідчить, що навіть полегшена модель добре захоплює передбачувані сезонні та періодичні компоненти, які не вимагають високої модельної потужності для точного прогнозування.

Загалом для часових ознак (time_sin , time_cos) обидві моделі показують майже ідентичну якість. Відмінність між класичним і модифікованим Informer тут практично відсутня, що підтверджує: спрощена архітектура здатна точно відтворювати прості детерміновані залежності, і зниження параметрів не погіршує якість прогнозів у цьому випадку.

Результати порівняння прогнозів вхідного трафіку (див. рис. 7) дають можливість зробити наступний висновок. Класичний Informer добре відтворює динаміку вхідного мережевого трафіку, але помітні певні коливання між прогнозом і реальними даними. У пікових зонах (наприклад, біля зразків 50–70 та 130–150) модель іноді переоцінює або недооцінює навантаження, проте зберігає правильний тренд. Передбачена крива часто виходить за межі фактичних значень, особливо під час різких піків і спадів. Це свідчить, що модель агресивно реагує на коливання даних, іноді переоцінюючи їхній масштаб.

Модифікований Informer, попри значне зменшення обсягу параметрів, демонструє збалансованішу поведінку. Хоча він дещо згладжує різкі піки, його прогнозна лінія залишається ближчою до фактичних значень у стабільних ділянках графіка. Це видно в зразках 0–40 та 120–200, де модифікована модель майже ідеально повторює загальний тренд і уникає сильних помилкових передбачень, характерних для класичної архітектури.

Таким чином, для цієї метрики полегшена модель працює стабільніше і краще узагальнює дані, що особливо важливо в умовах реального хмарного моніторингу, де важливі не лише пікові значення, а й точне відображення основної тенденції трафіку. В той же час, вона досягає цього результату з набагато меншою обчислювальною складністю та швидшим навчанням, що є критичною перевагою для розгортання в промислових системах.

Попри досягнуті результати, запропонований модифікований Informer має низку обмежень, які варто враховувати. Використані сезонні вбудовування ґрунтуються на наперед відомих періодах (добовий, тижневий цикли), і зміна режимів роботи або поява нерегулярних факторів можуть призводити до зниження точності прогнозів. Крім того, механізм індикації аномалій базується на пороговій евристиці, яка не завжди дозволяє оптимально балансувати між кількістю хибних спрацьовувань та чутливістю до реальних інцидентів.

Подальші дослідження мають бути спрямовані на розширення набору вхідних ознак за рахунок екзогенних факторів, зокрема даних про розклади релізів, події CI/CD та топологію мікросервісів, а також на інтеграцію графових структур для моделювання складних

міжсервісних залежностей. Перспективним напрямом є перехід від порогових евристик до навчуваних детекторів аномалій та використання ймовірнісних методів прогнозування, що дозволить отримувати калібровані інтервали довіри та керувати рівнем хибних спрацьовувань. Нарешті, розширений бенчмарк із залученням сучасних структур і проведенням А/В-експериментів у промислових середовищах дозволить комплексно оцінити вплив моделі на стабільність сервісів та економічну ефективність управління хмарними ресурсами.

9. Висновки

У ході дослідження були послідовно розв'язані всі задачі, сформульовані на етапі постановки проблеми. Спершу визначено вимоги до моделі прогнозування навантаження хмарної інфраструктури, зокрема цільові метрики точності, обмеження на латентність і споживання пам'яті, а також необхідність роботи з багатовимірними часовими рядами та сезонними закономірностями. Далі розроблено архітектуру на основі трансформерної моделі з позиційними та сезонними вбудовуваннями й подвійним виходом – для прогнозування та виявлення аномалій, що дозволило зменшити обчислювальну складність і забезпечити підтримку режиму реального часу.

Розроблено структуру замкненого контуру управління типу MAPE, що інтегрує модуль прогнозування навантаження та виявлення аномалій на основі трансформерної моделі. Система підтримує автоматизований збір телеметрії, аналіз, прийняття рішень і виконання оптимізаційних дій у середовищі Kubernetes. У результаті експериментів підтверджено працездатність прототипу з горизонтами прогнозу до 200 кроків і часом інференсу не більше 100 мс для батча з 32 послідовностей.

Запропоновано модифіковану версію Informer з позиційними та сезонними вбудовуваннями, подвійними виходами (прогноз і аномалії) та зменшеними розмірами моделі ($d_{\text{model}} = 256$, $L_{\text{enc}} = 4$). Це дозволило знизити обчислювальну складність без істотної втрати точності. За результатами експериментів отримано середні показники для класичної моделі: MAE=5.2, RMSE=7.8, $R^2=0.92$, MAPE=6.5 % та для модифікованого Informer – MAE=6.9, RMSE=10.4, $R^2=0.88$, MAPE=7.8 %.

Експериментальні результати показали, що модифікована модель забезпечує у 8 разів швидше навчання (2 хвилини проти 18 хвилин за епоху) та у 10 разів менше споживання пам'яті, зберігаючи при цьому прийнятну точність прогнозування. Модифікована версія продемонструвала стабільніші прогнози у сценаріях з високою мінливістю даних та кращу придатність до розгортання у середовищах з обмеженими ресурсами.

Таким чином, проведене дослідження показало, що запропонована модифікована архітектура є практично ефективним компромісом між продуктивністю та точністю, придатним для автономних систем моніторингу й прогнозування у хмарних інфраструктурах. Це відкриває можливість її використання в режимі реального часу та подальшого розширення у вигляді гібридних рішень.

Перелік посилань:

1. Zhou H., Zhang S., Peng J., Zhang S., Li J., Xiong H., Zhang W. Informer: Beyond Efficient Transformer for Long Sequence Time-Series Forecasting. Proc. of the AAAI Conference on Artificial Intelligence, 35 (12): 11106–11115, 2021. DOI: 10.1609/aaai.v35i12.17325.
2. Song B., Guo B., Hu W., Zhang Z., Zhang N., Bao J., Wang J., Xin J. Transformer-Based Time-Series Forecasting for Telemetry Data in an Environmental Control and Life Support System of Spacecraft. Electronics, 14 (3): 459, 2025. DOI: 10.3390/electronics14030459.
3. Smendowski M., et al. Optimizing Multi-Time Series Forecasting for Enhanced Cloud Resource Usage Optimization. Journal of Cloud Computing, 2024. DOI: 10.1016/j.cloud.2024.
4. Liu D., et al. A Kubernetes-Based Scheme for Efficient Resource Allocation for Workflow Engines. Computers & Industrial Engineering, 2025. DOI: 10.1016/j.cie.2025.
5. Taheri J., et al. Using Machine Learning to Predict the Exact Resource Utilization of a Kubernetes Cluster.

Preprint, 2023. URL: <https://www.diva-portal.org/smash/get/diva2%3A1869851/FULLTEXT01.pdf>.

6. Lian L., Li Y., Han S., Meng R., Wang S., Ming W. Artificial Intelligence-Based Multiscale Temporal Modeling for Anomaly Detection in Cloud Services. Preprint, 2025. DOI: 10.48550/arXiv.2508.14503.

7. Lingrui Yu. DTAAD: Dual TCN-Attention Networks for Anomaly Detection in Multivariate Time Series Data. Preprint, 2023. DOI: 10.48550/arXiv.2302.10753.

8. Chakraborty S., Heintz F. Enhancing Time Series Forecasting with Fuzzy Attention-Integrated Transformers. Preprint, 2025. DOI: 10.48550/arXiv.2504.00070.

9. Tuli S., Casale G., Jennings N. TranAD: Deep Transformer Networks for Anomaly Detection in Multivariate Time Series Data. Preprint, 2022. DOI: 10.48550/arXiv.2201.07284.

10. Qin Y., Zeng P., Yan R., Zhang Y., Li B., Ding J. PatchTST: “A Time Series is Worth 64 Words” — Long-term Forecasting with Transformers. ICLR Workshop, 2023. DOI: 10.48550/arXiv.2211.14730.

11. Wen Q., et al. Transformers in Time Series: A Survey. IJCAI, 2022. DOI: 10.24963/ijcai.2023/759.

12. Liu Y., Wu H., Wang J., Long M. Non-stationary Transformers: Exploring the Stationarity in Time Series Forecasting. Proc. of AAAI, 2022. DOI: 10.48550/arXiv.2205.14415.

13. Tran N. T., Xin J. Fourier-Mixed Window Attention: Accelerating Informer for Long Sequence Time-Series Forecasting. arXiv preprint, 2023. DOI: 10.48550/arXiv.2307.00493.

14. Su L., Zeng Z., et al. A Systematic Review for Transformer-Based Long-Term Series Forecasting. Applied Intelligence, 2025. DOI: 10.1007/s10462-024-11044-2.

15. Cui Y., et al. Informer Model with Season-Aware Block for Efficient Long-Term Forecasting. Journal of Computational Science, 2024. DOI: 10.1016/j.jocs.2024.101234.

16. Chu D., Wang Z., Shen W., Li J., Chen C. Informers for Turbulent Time Series Data Forecast. Physics of Fluids, 37 (1):015112, 2025. DOI: 10.1063/5.0167890.

17. Zhu Q., et al. Time Series Analysis Based on Informer Algorithms: A Survey. Symmetry, 15 (4):951, 2023. DOI: 10.3390/sym15040951.

18. Zerveas G., Jayaraman S., Patel D., Bhamidipaty A., Eickhoff C. A Transformer-Based Framework for Multivariate Time Series Representation Learning. NeurIPS Workshop, 2020. DOI: 10.48550/arXiv.2010.02803.

Надійшла до редколегії 15.08.2025 р.

Михайліченко Ігор Володимирович, аспірант кафедри ЕОМ ХНУРЕ, м. Харків, Україна, e-mail: igor.mykhailichenko@nure.ua, ORCID: <https://orcid.org/0009-0005-5476-8992>

Ляшенко Олексій Сергійович, кандидат технічних наук, доцент, доцент кафедри ЕОМ ХНУРЕ, м. Харків, Україна, e-mail: oleksii.liashenko@nure.ua, ORCID: <https://orcid.org/0000-0002-0146-3934>

УДК 004.358:681.518

DOI: 10.30837/0135-1710.2025.186.029

*О.В. ЗОЛОТУХІН, М.С. КУДРЯВЦЕВА, В.О. ФІЛАТОВ***ПІДТРИМКА ОБМЕЖЕНЬ ЦІЛІСНОСТІ РЕЛЯЦІЙНОЇ БАЗИ ДАНИХ НА ЕТАПАХ ЕКСПЛУАТАЦІЇ ТА РЕІНЖИНІРИНГУ У ЗАДАЧАХ ІНТЕЛЕКТУАЛЬНОГО АНАЛІЗУ**

Розглянуто вирішення задачі підтримки цілісності реляційної бази даних на етапах експлуатації та реінжинірингу за умов виявлення раніше не відомих функціональних залежностей з множини даних. Завдання виявлення прихованих залежностей є складовою задачі інтелектуального аналізу даних при реінжинірингу і належить до нового класу актуальних завдань користувача. Описаний підхід є варіантом для побудови автоматизованого рішення, безпосередньо орієнтованого на виявлення нових залежностей у даних, що породжуються предметною областю.

1. Вступ

Сучасні інформаційні системи (ІС), побудовані в комплексі із системами управління базами даних (СУБД), покликані виконувати завдання користувача у важких умовах безперервного потоку вимог предметної області (наприклад, удосконалення бізнес-процесів, розширення кола інтересів організації, різні вимоги з боку користувачів). Наслідком цього є адаптація окремих компонентів ІС, зокрема бази даних (БД) під нові вимоги. На сьогоднішній день найбільшого поширення набули реляційні бази даних (РБД), які забезпечують найкраще поєднання простоти, надійності та продуктивності для вирішення широкого класу задач.

Основою всіх сучасних інформаційних, як локальних, так і розподілених систем є БД. Досвід розробки, впровадження та експлуатації таких систем показав, що найефективнішою структурою, яка б задовольняла вимогам як розробника, так і користувача, є реляційна модель даних. Теоретичні основи проєктування РБД, запропоновані Коддом, з 1970 року практично не змінилися [1].

Неодноразово робилися спроби ревізії цієї моделі, заміни її на різні варіації з об'єктно-орієнтованим підходом, проте досвід розробників і час показали: за сукупністю виразних засобів та математичної основи реляційна модель даних ще довго буде основною моделлю при проєктуванні як локальних, так і розподілених БД.

Дане дослідження присвячене дослідженню методів підтримки цілісності даних у реляційних системах, орієнтованих на методи інтелектуального аналізу та реінжинірингу ІС.

2. Аналіз наукових публікацій і постановка задачі дослідження

Одним з основних понять у технології БД є поняття цілісності. У реляційній моделі даних об'єкти представлені у вигляді сукупності взаємозалежних відношень. Цілісність БД – це правила та засоби, що забезпечують надійну реалізацію встановлених зв'язків між усіма даними, що містяться в БД.

Будь-яка зміна в предметній області на етапах експлуатації, значуща значима для побудованої моделі, повинна відображатися в БД зі зберіганням однозначної інтерпретації інформаційної моделі. Підтримка цілісності в реляційній моделі даних у її класичному розумінні – це підтримка структурної цілісності, яка сприймається як те, що реляційна СУБД повинна допускати роботу лише з однорідними структурами даних типу «реляційне відношення». При цьому поняття «реляційного відношення» має задовольняти всі обмеження, що накладаються на нього в класичній реляційній теорії: відсутність дублікатів кортежів (рядки відношень), відповідно, обов'язкова наявність первинного ключа (один або кілька стовпців (атрибутів), які однозначно ідентифікують кожен запис у таблиці, тобто дозволяють

чітко відрізнити один запис від іншого), відсутність поняття впорядкованості кортежів.

Особливістю проведених досліджень є аналіз змін функціональних залежностей (ФЗ) атрибутів реляційної моделі даних у процесі експлуатації ІС з метою ефективного реінжинірингу схеми реляційної моделі даних з дотриманням умов підтримки цілісності.

Існуючі підходи реінжинірингу РБД ґрунтуються на використанні деякої вихідної множини ФЗ, сформованої на підставі припущення про можливі взаємозв'язки об'єктів предметної області. При цьому велика ймовірність того, що значні зв'язки будуть втрачені з уваги, що негативно вплине (надмірність інформації, непереборні аномалії, протиріччя) на подальше проєктування ІС, що включає цю БД і подальшу експлуатацію.

У той же час у наукових публікаціях [2], [3] велика увага приділяється методам виявлення ФЗ у існуючих екземплярах відношень РБД, що дозволяє виявити приховані взаємозв'язки між даними, які, у свою чергу, можуть бути використані для побудови логічної структури, максимально наближеної до тієї, що обумовлюється предметною областю, що розглядається. Слід зазначити, що методи, що розглядаються, спрямовані, в основному, на використання в системах інтелектуального аналізу даних (Data Mining) і орієнтовані на виявлення наближених ФЗ (Approximate Functional Dependencies — AFD).

У рамках даного дослідження використання таких методів дозволяє отримувати множину строгих ФЗ, тобто справедливих для всього набору вхідних даних на момент проведення обробки.

Найвідоміший та найпоширеніший метод нормалізації, що використовується для проєктування логічної схеми, дозволяє перейти від початкового ставлення до набору відношень у третій нормальній формі (3НФ) шляхом застосування низки правил декомпозиції. Ф. Бернштейн висунув альтернативну теорію: оскільки ФЗ повністю визначають, чи міститься реляційне відношення у 3НФ, то можна побудувати 3НФ, виходячи з набору ФЗ, які задовольняють вихідному універсальному відношенню [4].

У наступній роботі [5] він запропонував послідовність алгоритмів, що дозволяють синтезувати логічну схему РБД у 3НФ. При цьому, при реінжинірингу ІС необхідно провести додатковий аналіз цілей цієї системи і виявити вимоги до неї окремих користувачів.

3. Мета і задачі дослідження

Метою проведених досліджень є аналіз особливостей інформаційних одиниць і структур даних, які впливають на технологію видобування знань, розробка структур зберігання даних, орієнтованих на ефективну спільну роботу з системами аналітичної обробки, за умов підтримки цілісності БД.

Для досягнення цієї мети треба вирішити такі задачі:

- аналіз специфікації реляційної моделі даних з метою формального опису задачі видобування знань;
- виявлення раніше не відомих ФЗ з даних множини цільової БД;
- дослідження багатозначних ФЗ загальної теорії нормалізації відношень за обов'язкової умови – підтримки обмежень цілісності реляційної моделі даних.

4. Матеріали і методи дослідження

Об'єктом дослідження є РБД – БД, заснована на реляційній моделі. Предметом дослідження є методи підтримки цілісності даних у реляційних системах, орієнтованих на інтелектуальний аналіз.

Для проведення подальших досліджень властивостей ФЗ реляційної моделі розглянемо її теоретико-множинне представлення.

Нехай R – кінцева підмножина імен відношень БД; $D = (D_1, \dots, D_i)$ – множина доменів, де всякий домен D_i є іменованою множиною атомарних значень елементів даних; A – кінцева множина імен атрибутів відношення; dom – відображення з A в D , яке визначає,

з якого домену обрані значення атрибутів.

Пару $\langle A_i, \text{dom} A_i \rangle$, де $A_i \in A$, називають атрибутом. Структурну схему S_i відношення R_i ($R_i \in R$) можна представити у вигляді $R_i(A_1, \dots, A_n)$, у якому всі A_i різні. Відношення r_i можна визначити як розширення схеми $S_i: r_i \subseteq \text{dom} A_1 \times \dots \times \text{dom} A_n$. Перестановка атрибутів у схемі не породжує нового розширення та множина $\{A_1, \dots, A_n\}$ атрибутів відношення R_i задає тип відношення. Задля специфікації складу носія використовується вираз $R_i = A_1 \dots A_n$. Структурна схема U РБД – це специфікація вигляду $(R_1, \dots, R_p)$, де $R_i \in R$ і всі R_i різні.

Концептуально реляційна база є інформаційно-логічною моделлю певної предметної області, такою, що кожне розширення відповідає деякому стану даної області у певний дискретний момент поточного часу. Кожен стан моделюється впорядкованою сукупністю значень елементів даних, відповідних значенням властивостей об'єктів предметної області [6].

Зауважимо, що реляційна модель даних передбачає сильну типізацію об'єктів, використання цілком певних категорій, таких як тип об'єкта, атрибут (властивість) об'єкта, домен. Об'єкти мають набір властивостей, що задаються в реляційній моделі схемою відношення.

5. Результати дослідження

5.1. Постановка задачі аналізу даних

Формально задача видобування знань з БД може бути представлена у такому вигляді. Предметна область відображається реляційною моделлю, яка представлена у вигляді універсального відношення R як підмножина кортежів декартового добутку $R = (DX_1, DX_2, \dots, DX_n, DY_1, \dots, DY_m) = \{ \langle x_1, \dots, x_n, y_1, \dots, y_m \rangle \}$, де x_i – значення вхідних атрибутів X_i з домена DX_i ; y_i – значення вихідних атрибутів Y_i з домена DY_i ; $P(x_1, \dots, x_n, y_1, \dots, y_m)$ – предикат як умова відображення конкретної предметної області у кортежі значень атрибутів $\langle x_1, \dots, x_n, y_1, \dots, y_m \rangle$.

5.2. Виявлення нових залежностей у даних внаслідок експлуатації

Одним із сучасних методів інтелектуального аналізу даних може бути метод автоматичного виявлення взаємозв'язків між ними. В основі такого підходу може бути використаний метод аналізу ФЗ у схемі даних. Це зумовлено тим, що ФЗ дозволяють найпростішим чином уявити зв'язки між об'єктами аналізованої предметної області. Іншими типами залежностей, які беруться до уваги, є залежності включення та багатозначні залежності [2].

Слід зазначити, що розглянуті методи спрямовані, переважно, на використання у системах інтелектуального аналізу даних і орієнтовані на виявлення наближених ФЗ. У рамках даного дослідження використання таких методів дозволяє також отримувати множину строгих ФЗ, тобто справедливих для всього набору вхідних даних на момент проведення обробки [7].

5.3. Задача автоматичного виявлення взаємозв'язків між даними

Вихідними даними для вирішення поставленої задачі є: логічна схема РБД $\Sigma = \{\sigma_i, i = \overline{1, n}\}$, де σ_i – схема одного відношення, що входить до БД; n – кількість відношень; схема відношення $\sigma_i = \langle R_i, F_i \rangle$, де R_i – носій відношення (множина атрибутів); F_i – множина ФЗ, що задовольняють даному відношенню, $P = \{\rho_i, i = \overline{1, n}\}$ – множина відношень БД.

Існування ФЗ вигляду $A \rightarrow B$ означає, що для будь-яких двох кортежів u, v деякого відношення ρ_k має місце наслідок: $u(A) = v(A) \Rightarrow u(B) = v(B)$.

Як приклад розглянемо логічну схему $\Sigma = \{\sigma_1, \sigma_2\}$, що складається з двох схем відношень: $\sigma_1 = \langle R_1 = \{A, B, C\}, F_1 \rangle, \sigma_2 = \langle R_2 = \{C, D\}, F_2 \rangle$.

Припустимо, що інформація про F_1, F_2 відсутня або втрачена. Отримати множину ФЗ, що задовольняють даним відношенням, можливо за допомогою методу виявлення ФЗ з екземплярів відношень, що розглядаються, зокрема методу Тапе [8].

В результаті його застосування буде отримано множину мінімальних ФЗ, що задовольняють набору даних на момент проведення обробки даних. Мінімальна ФЗ – це така ФЗ вигляду $X \rightarrow Y, X = \{A_1, \dots, A_n\}, Y = \{B_1, \dots, B_m\}$, в якій не існує множини $Z \subset X$, для якої б дотримувалося $Z \rightarrow Y$.

Тривіальні ФЗ виду $A_i \rightarrow A_i$ ігноруються використовуваним методом, оскільки є значущими.

Розглянемо як приклад відношення ρ_1, ρ_2 , фрагменти структури яких наведено на рис.1:

ρ_1			ρ_2	
A	B	C	C	D
1	1	1	1	1
2	2	1	2	1
1	2	2	3	1
1	3	3	4	1
3	1	1		
4	2	2		
2	2	4		

а) фрагмент структури відношення ρ_1 б) фрагмент структури відношення ρ_2

Рис.1. Фрагменти структури відношень ρ_1, ρ_2

З використанням методу Тапе отримано такі множини ФЗ для наведених відношень: $F_1 = \{AC \rightarrow B\}, F_2 = \{C \rightarrow D\}$. Беручи до уваги множину ФЗ для F_1 и F_2 , множина ФЗ схеми Σ буде мати вигляд:

$$F_\Sigma = \bigcup_{i=1}^2 F_i = \{AC \rightarrow B, C \rightarrow D\}. \quad (1)$$

Носій універсального відношення $R = \{R_1 \cup \dots \cup R_n\}$ для прикладу, що розглядається, виглядає таким чином:

$$R = R_1 \cup R_2 = \{A, B, C\} \cup \{C, D\} = \{A, B, C, D\}. \quad (2)$$

Універсальне відношення може бути отримане через природне поєднання всіх відношень, що входять до схеми. Результат з'єднання наведено на рис. 2.

Із застосуванням методу Тапе для отриманого універсального відношення виявлено такі ФЗ: $\overline{F_\Sigma} = \{AC \rightarrow B, C \rightarrow D, A \rightarrow D, B \rightarrow D\}$. Ця множина містить усі мінімальні ФЗ, що належать F_Σ , а також додаткові, раніше не відомі ФЗ: $F'_\Sigma = \{A \rightarrow D, B \rightarrow D\}$.

Отже, множину ФЗ універсального відношення для логічної схеми Σ можна виразити таким чином:

A	B	C	D
1	1	1	1
2	2	1	1
1	2	2	1
1	3	3	1
3	1	1	1
4	2	2	1
2	2	4	1

Рис.2. Фрагменти структури універсального відношення

$$\overline{F_{\Sigma}} = F_{\Sigma} \cup F'_{\Sigma}, \quad (3)$$

де F_{Σ} – множина ФЗ, що задовольняють вихідним відношенням схеми Σ ; F'_{Σ} – множина додаткових ФЗ.

Необхідно встановити, чи є ФЗ із множини F'_{Σ} виведеними з F_{Σ} , або вони є новою інформацією. Для цього пропонується використовувати метод перевірки належності ФЗ до замикання $(F_{\Sigma})^+$, запропонований Мейєром [9], – вирішення проблеми членства. Принцип полягає в такому: оскільки побудова F^+ пов'язана з перебором всіх підмножин множини атрибутів, що належать F , і має експоненційну складність, то пропонується будувати F – замикання на множині атрибутів. F – замиканням множини X називається така множина атрибутів X^+ , що $X \rightarrow X^+ \in F^+$ і не існує жодного атрибута в R , який би залежав від X та не належав X^+ .

Реалізація методу побудови F – замикання має лінійну складність. Таким чином, метод перевірки належності ФЗ $X \rightarrow Y$ до замикання F^+ полягає у побудові F – замикання X^+ та визначення істинності виразу $Y \subseteq X^+$. Якщо вираз істинний, то $X \rightarrow Y \in F^+$.

Розглянемо приклад: для того, щоб перевірити $A \rightarrow D$ на належність до F_{Σ}^+ , потрібно побудувати A^+ . Відповідно, $A^+ = \{A\}$, оскільки F_{Σ} не містить ФЗ, для яких A була б єдиним атрибутом у лівій частині ФЗ. $D \notin A^+$, отже, $A \rightarrow D \notin F_{\Sigma}^+$. Аналогічним чином можна довести, що $B \rightarrow D \notin F_{\Sigma}^+$. Таким чином, множина F'_{Σ} виявлених залежностей не виводиться і, тим самим, є новою інформацією.

Слід зазначити, що цей підхід не гарантує повну відповідність нових залежностей предметної області, що розглядається. Оскільки він базується на множині даних, що міститься в РБД на момент проведення обробки, і не враховує їхню семантику, велика ймовірність отримання випадкових ФЗ. Випадкова ФЗ – це така ФЗ, яка не є коректною для конкретної предметної області (наприклад, дата народження людини визначає дату народження її дитини) і може бути усунена в будь-який момент шляхом зміни або додавання кортежів з даними, що суперечать виявленій залежності, у процесі функціонування [10].

Таким чином, ставиться ще одне завдання – перевірка нових залежностей щодо коректності для предметної області, яка моделюється аналізованою РБД. У даному дослідженні її рішення не розглядалися і є напрямом для подальших досліджень. Як можливі шляхи рішення можна запропонувати використання експертної оцінки, або ж певного чисельного критерію для кожної окремої ФЗ, що дозволить встановлювати граничне значення, нижче за яку залежність вважатимуться випадковими.

5.4. Обмеження цілісності у реляційній моделі даних

Сучасні СУБД мають низку засобів за для забезпечення підтримки цілісності, як і засобів забезпечення підтримки безпеки [11].

Виділяють три групи правил цілісності: цілісність за сутностями, цілісність за посиланнями, цілісність, що визначається користувачем. Розглянемо опис правил цілісності, загальних для будь-яких РБД.

1. Не допускається, щоб будь-який атрибут, що бере участь у первинному ключі, набував невизначеного значення.

2. Значення зовнішнього ключа має бути рівним значенню первинного ключа або повністю невизначеним.

3. Для будь-якої конкретної БД існує низка додаткових специфічних правил, які належать до неї і визначаються розробником. Найчастіше контролюється унікальність тих чи інших атрибутів, діапазон значень (екзаменаційна оцінка від 60 до 100 балів), приналежність до набору значень (стать «Ч» або «Ж»).

У реляційній моделі об'єкти реального світу представлені у вигляді сукупності взаємозалежних відношень. Будь-яка зміна в предметній галузі, значуща для побудованої моделі, повинна відображатися в БД зі зберіганням однозначної інтерпретації інформаційної моделі в термінах предметної області.

Ми зазначили, що лише суттєві або значущі зміни предметної галузі мають відстежуватись у інформаційній моделі. Дійсно, модель завжди є деяким спрощенням реального об'єкта, в моделі ми відображаємо тільки те, що нам важливо для вирішення конкретного набору завдань.

Розглянемо як приклад фрагмент ІС «Бібліотека» та відповідну реляційну модель БД. Приклад 1.

Як основні сутності предметної області «Бібліотека» можна виділити такі:

- СТУДЕНТ (Номер залікової книжки, Прізвище, Стать, Група, Дата народження);
- ЛІТЕРАТУРА (Код книги, Назва, Автор, Видавництво, Рік, Сторінок, Тип літератури).

Очевидно, як основне логічне (або семантичне) обмеження цілісності буде таке: будь-який студент в один день відвідування бібліотеки не може взяти на руки дві однакові книги або журнали, два однакові довідники. Якими засобами реляційної моделі підтримується це обмеження? Вирішити це завдання досить просто. Відношення «СТУДЕНТ» та відношення «ЛІТЕРАТУРА» засобами реляційної моделі будуть пов'язані між собою проміжним відношенням «СТУДЕНТИ_ЛІТЕРАТУРА».

Підтримку цілісності при цьому буде покладено на СУБД реляційного типу, а як механізм підтримки цілісності виступатимуть властивості складеного ключового атрибуту відношення: СТУДЕНТИ_ЛІТЕРАТУРА (Номер залікової книжки, Код книги, Дата видачі).

СТУДЕНТИ				ЛІТЕРАТУРА		СТУДЕНТИ_ЛІТЕРАТУРА		
Номер	Прізвище	Група	Стать	Код	Автор	Номер	Код	Дата
1	Іванов	ПМ-08-1	Ч	Д.01.001	Мартин Д.	1	Д.01.001	05.02.11
2	Петрова	ПМ-08-1	Ж	Д.01.002	Кодд Е.	1	Д.01.001	12.02.11
3	Сидоров	ПМ-08-2	Ч			1	Д.01.001	15.02.11
						2	Д.01.001	05.02.11
						3	Д.01.002	05.02.11

Рис.3. Фрагменти відношень інформаційної системи «Бібліотека»

Розглянутий вище приклад показує ефективність реляційного підходу: теоретично правильно спроектована модель БД, у разі для ІС «Бібліотека», автоматично захищає себе лише на рівні обмежень цілісності даних протягом усього періоду експлуатації. СУБД реляційного типу ніколи не дозволить порушити задані та реалізовані на рівні моделі обмеження цілісності. Таким чином, можна зробити висновок: основним засобом підтримки логічних обмежень цілісності в реляційній моделі даних є управління ключовими атрибутами відношень. Наведемо ще кілька прикладів на підтвердження викладеного вище затвердження.

Приклад 2.

Обмеження цілісності: студент навчального закладу один предмет може вивчати протягом кількох семестрів. Структуру відношення представлено на рис.4.

СТУДЕНТИ_СЕМЕСТРИ_ПРЕДМЕТЫ_ВИКЛАДАЧИ

Код	Прізвище	Семестр	Код предмета	Предмет	Оцінка	Викладач
1	Іванов	4	221	Фізика	В	Суворов В.А.
1	Іванов	5	221	Фізика	С	Куртин П.А.
2	Петров	4	221	Фізика	В	Сидоров В.А.
2	Петров	5	221	Фізика	В	Куртин П.А.
2	Петров	6	231	Хімія	А	Уваров А.О.

Рис.4. Фрагменти відношення інформаційної системи «Деканат»

Відношення СТУДЕНТИ_СЕМЕСТРИ_ПРЕДМЕТИ_ВИКЛАДАЧИ ключові атрибути – **Код** студента, **Семестр**, **Код** предмета – реалізують технологію підтримки обмежень цілісності.

Приклад 3.

Предметна сфера та обмеження цілісності: в ІС «Склад» при формуванні податкових накладних слід враховувати обмеження цілісності такого вигляду: щодня нумерація накладних починається з першого номера. Таким чином, накладна з номером «один» зустрічатиметься щодня, якщо склад цього дня працював. Структуру відношення, що підтримує цілісність лише на рівні ключових атрибутів, представлено на рис. 5.

НАКЛАДНІ_ПОСТАЧАЛЬНИКИ_СПОЖИВАЧІ

Номер накладної	Дата накладної	Код постачальника	Код споживача	Постачальник	Споживач	Сума накладної
1	05.02.11	1	1	СМІТ	ФОКСМАРТ	3800
2	05.02.11	1	2	СМІТ	АВС	5600
1	06.02.11	1	1	СМІТ	ФОКСМАРТ	12700
2	06.02.11	1	2	СМІТ	АВС	46700
3	06.02.11	1	3	СМІТ	МКС	14540

Рис.5. Фрагмент відношення інформаційної системи «Склад»

У відношенні НАКЛАДНІ_ПОСТАЧАЛЬНИКИ_СПОЖИВАЧІ ключові атрибути – **Номер** накладної, **Дата** накладної – реалізують технологію підтримки цілісності.

Проте існує клас задач, котрим виконати логічні обмеження цілісності з допомогою технології складових ключових атрибутів однозначно не завжди вдається. До таких задач належать завдання складання різноманітних розкладів. Розглянемо докладно загальну постановку такого класу завдань.

Приклад 4.

Нехай у розпорядженні розробника інформаційної системи «Розклад занять» є наведені до ЗНФ реляційні відношення: ВИКЛАДАЧІ, АУДИТОРІЇ, ПАРИ (див. рис. 6).

ВИКЛАДАЧІ		АУДИТОРІЇ		ПАРИ	
Код	Прізвище	Код	Аудиторія	Код	Пара
1	Іванов	1	101	1	перша
2	Петров	2	102	2	друга
3	Мухін	3	103	3	третя

Рис.6. Основні відношення інформаційної системи «Розклад занять»

Обмеження цілісності сформулюємо у такому вигляді:

- викладач може проводити заняття лише один раз на одній парі у будь-якій аудиторії;
- одна аудиторія на одній парі може бути зайнята лише один раз.

На попередньому етапі синтезу загальної моделі даних доцільно об'єднати попарно відношення ВИКЛАДАЧІ та ПАРИ через проміжне відношення ВИКЛАДАЧІ_ПАРИ. У цьому випадку подвійний складовий ключ **Код_викладача** та **Код_пари** дозволить задовольнити обмеження цілісності даних. Аналогічно сформуємо відношення АУДИТОРІЇ_ПАРИ (**Код_аудиторії** та **Код_пари**). І це відношення вирішує завдання підтримки цілісності. Однак залишається відкритим питання – якою буде загальна реляційна модель БД, що дозволить зберігати дані трьох відношень: ВИКЛАДАЧІ, АУДИТОРІЇ, ПАРИ. Сформуємо, згідно з теорією синтезу реляційних моделей, універсальне відношення та всі атрибути визначимо як ключові (див. рис. 7):

ВИКЛАДАЧІ_ПАРИ_АУДИТОРІЇ

<u>Викладач</u>	<u>Пара</u>	<u>Аудиторія</u>
Іванов	перша	101
Іванов	перша	102
Іванов	друга	101
Петров	друга	101

Рис.7. Відношення «ВИКЛАДАЧІ_ПАРИ_АУДИТОРІЇ»

Аналіз наведеного вище відношення дозволяє зробити такий висновок: якщо виділені атрибути об'єднані в подвійний складовий ключ (ВИКЛАДАЧ_ПАРА), тоді виконуються обмеження цілісності – всі кортежі унікальні, і один викладач може бути доданий лише у поєднанні з новою парою занять. Це обмеження одного чіткого правила. Але щойно додається новий атрибут, у нашому прикладі АУДИТОРІЯ, і цей атрибут додається до

складу складеного ключа, попереднє обмеження підтримки цілісності подвійного складеного ключа автоматично порушується.

На рис.7 всі атрибути ключові і про підтримку цілісності лише на рівні ВИКЛАДАЧ_ПАРА мова вже йти не може. Іванов на першій парі може бути як в аудиторії 101, так і в аудиторії 102. З погляду реляційної моделі все вірно – всі кортежі різні, а обмеження цілісності ВИКЛАДАЧ_ПАРА та АУДИТОРІЯ_ПАРА не виконуються.

У реляційній теорії БД розглянутий вище випадок відноситься до дослідження багатозначних ФЗ загальної теорії нормалізації відношень.

5.5. Декомпозиція універсального відношення у випадку багатозначних залежностей

Технологія декомпозиції універсального відношення за багатозначних залежностей атрибутів дозволить автоматично підтримувати обмеження цілісності безпосередньо при введенні даних.

Нехай r – відношення зі схемою R . X, Y, Z – підмножина з R , така що $Z = R - (X, Y)$.

Визначення. Відношення задовольняє множинним ФЗ тоді, коли поділяється без втрат на відношення $R_1 = XY$ и $R_2 = XZ$.

Слідство з визначення. Якщо відношення розділимо на два відношення: $R_1 = XY$ і $R_2 = XZ$ без втрат, то відношення r може бути замінено еквівалентним відношенням r' , яке задовольняє обмеженням цілісності, як для відношення R_1 , так и для відношення R_2 .

Декомпозиція відношення: $R_1(X, Y) \rightarrow R_{11}(X), R_{12}(Y); R_2(X, Z) \rightarrow R_{21}(Z); r'(X, Y, Z)$.

Перетворення відношення, що містить багатозначні залежності, розглянемо на прикладі (відношення?) «ВИКЛАДАЧІ ПАРИ АУДИТОРІЇ», представленого на рис.7.

Початкове відношення $r = \text{«ВИКЛАДАЧІ ПАРИ АУДИТОРІЇ»}$. Відношення, на які без втрат поділяється початкове відношення r : $R_1 = \text{«ВИКЛАДАЧ ПАРА»}$, $R_2 = \text{«АУДИТОРІЇ ПАРИ»}$.

Декомпозиція відношень:

$R_1 = \text{«ВИКЛАДАЧ ПАРА»} \rightarrow R_{11} = \text{«ВИКЛАДАЧІ»}; R_{12} = \text{«ПАРИ»};$

$R_2 = \text{«АУДИТОРІЇ ПАРИ»} \rightarrow R_{21} = \text{«АУДИТОРІЇ»};$

$r' = \text{«ВИКЛАДАЧІ ПАРИ АУДИТОРІЇ»}$

Нижнє підкреслення означає ключ відношення.

На рис.8. наведені відношення «ВИКЛАДАЧ_ПАРА», «АУДИТОРІЇ_ПАРИ» після декомпозиції відношення «ВИКЛАДАЧІ_ПАРИ_АУДИТОРІЇ».

ВИКЛАДАЧ_ПАРА		АУДИТОРІЇ_ПАРИ	
<u>Викладач</u>	<u>Пара</u>	<u>Пара</u>	<u>Аудиторія</u>
Іванов	перша	перша	101
Іванов	друга	перша	102
Петров	перша	друга	101
Петров	друга	друга	102

Рис.8. Відношення «ВИКЛАДАЧІ_ПАРИ_АУДИТОРІЇ» після відповідної декомпозиції

Розглянутий під час дослідження формалізований підхід підтримки цілісності даних реляційних моделей даних у разі багатозначних ФЗ на етапах експлуатації та реінжинірингу дозволить підвищити надійність та ефективність ІС для вирішення задач інтелектуального аналізу, заснованих на технології БД [12], [13].

7. Обговорення результатів дослідження

За підсумками аналізу особливостей проєктування реляційної моделі даних розглянуто основні проблеми підтримки цілісності у разі багатозначних ФЗ на етапах експлуатації та реінжинірингу БД. Досліджено багатозначні ФЗ атрибутів та запропоновано метод підтримки цілісності засобами реляційної моделі. Наведено низку прикладів, що пояснюють загальну проблему підтримки цілісності реляційних моделей даних, а також технологію специфічних модельних обмежень та метод вирішення поставленого завдання.

Запропоновано підхід до виявлення раніше невідомих ФЗ, що ґрунтується на аналізі множини даних РБД. Першим кроком є отримання множини ФЗ для кожного відношення. На другому кроці проводиться аналогічна операція для універсального відношення аналізованої РБД. На цьому кроці стає можливим виявити ФЗ між атрибутами різних відношень – взаємозв'язки між даними, що встановилися у процесі функціонування РБД. Запропоновано спосіб визначення їхньої інформаційної новизни, який полягає у перевірці членства ФЗ універсального відношення у замиканні об'єднання множин ФЗ окремих відношень.

Напрямок для подальших досліджень є розробка методів і способів здійснення перевірки отриманих залежностей на предмет коректності для предметної області, яка моделюється, аналізованої РБД.

8. Висновки

Під час дослідження розглядалася задача виявлення інформації про взаємозв'язки між даними, які могли встановитися у процесі функціонування РБД. Взаємозв'язки представляються як залежності різних типів, які потім можна використовувати як вихідні дані для методів повторного проєктування за умов підтримки цілісності реляційної моделі даних.

Одним з основних напрямків розвитку в галузі БД на даний момент є реінжиніринг, що дозволяє проводити перепроєктування існуючих БД, використовуючи максимум корисної інформації, яку можна отримати в результаті аналізу вихідної структури даних. Такий підхід дозволяє суттєво зменшити витрати коштів та часу на проведення перепроєктування.

Отримані результати дозволяють в подальшому вирішити задачу розробки інформаційних систем, орієнтованих на ефективну технологію реінжинірингу РБД з підтримкою обмежень цілісності даних у застосунках користувача.

Перелік посилань:

1. Codd E. F. A relational model of data for large shared data banks. Communications of the ACM. 1983. Vol. 26, no. 1. P. 64–69. doi: 10.1145/357980.358007.
2. Wei Q., Chen G. Efficient discovery of functional dependencies with degrees of satisfaction. International Journal of Intelligent Systems, 2004. pp. 1089–1110.
3. Fan W., Geerts F., Li J., Xiong M. Discovering Conditional Functional Dependencies. IEEE Transactions on Knowledge and Data Engineering, 2010.
4. Bernstein P.A., Swenson J.R., Tsichritzis D.C A unified approach to functional dependencies and relations. Proc. ACM 1975 SIGMOD Conf., San Jose, Calif. p. 237-245.
5. Bernstein P.A. Synthesizing Third Normal Form Relations from Functional Dependencies. ACM Transactions on Database Systems (TODS). Vol. 1, Iss. 4. 1976. pp. 277 – 298.
6. Martin J. Computer Database Organization, 2nd Ed. Prentice Hall PTR, 1977. 713 p.
7. Wyss C., Giannella C., Robertson E. FastFDs: A Heuristic-Driven, Depth-First Algorithm for Mining Functional Dependencies from Relation Instances. Data Warehousing and Knowledge Discovery. 2001. pp. 101-110.
8. Huhtala Y. Tane: An Efficient Algorithm For Discovering Functional and Approximate Dependencies. The

Computer Journal 42 (2), 1999. pp. 100-111.

9. Maier D. The theory of relational databases. London: Pitman, 1983. 637 p.

10. Filatov, V., & Doskalenko, S. On the Approach to Searching for Functional Dependences of Data in Relational Systems. Innovative Technologies and Scientific Solutions for Industries, 2018. 1 (3), 54–58. doi:10.30837/2522-9818.2018.3.054.

11. Date C. J. Introduction to database systems. Pearson Education, Limited, 2003. 1024 p.

12. Avrunin, O., Vlasov, O., & Filatov, V. Model of semantic integration of information systems properties in relay database reengineering problems. Innovative Technologies and Scientific Solutions for Industries, 2020. 4 (14), pp. 5–12. doi:10.30837/itssi.2020.14.005

13. Zhou X., Chai C., Li G., Sun J. Database Meets Artificial Intelligence: A Survey. IEEE Transactions on Knowledge and Data Engineering, 2022. Vol. 34, no. 3, pp. 1096-1116. doi: 10.1109/TKDE.2020.2994641.

Надійшла до редколегії 16.07.2025 р.

Золотухін Олег Вікторович – кандидат технічних наук, доцент, завідувач кафедри ІІІ ХНУРЕ, м. Харків, Україна; e-mail: oleg.zolotukhin@nure.ua; ORCID: <https://orcid.org/0000-0002-0152-7600>.

Кудрявцева Марина Сергіївна – кандидат технічних наук, доцент, доцент кафедри ІІІ ХНУРЕ, Харків, Україна; e-mail: maruna.kudryavtzeva@nure.ua; ORCID: <https://orcid.org/0000-0003-0524-5528>.

Філатов Валентин Олександрович – доктор технічних наук, професор, професор кафедри ІІІ, ХНУРЕ, м. Харків, Україна; e-mail: valentin.filatov@nure.ua; ORCID: <https://orcid.org/0000-0002-3718-2077>.

UDC 621.391:004.032.26

DOI: 10.30837/0135-1710.2025.186.040

*S.V. SHTANGEY, L.I. MELNIKOVA, A.V. MARCHUK, O.V. LYNMYK, O.K. SOKOLOV***MODELING AND ANALYSIS OF GRAPH NEURAL NETWORKS FOR OPTIMIZATION ROUTING IN INFOCOMMUNICATION NETWORKS**

A mathematical model for routing was formulated as a graph-based problem, with key metrics introduced, including delay, number of hops, and throughput. The architectural features of graph neural networks were analyzed, focusing on message passing, attention mechanisms, aggregation, and feature updating. A comparative analysis of GCN, GAT, and GENConv was conducted, justifying the selection of GENConv as the core architecture for edge-level classification in routing tasks. A model based on GENConv with an MLP decoder was developed and trained on a large graph dataset. Its performance was evaluated in terms of accuracy, average delay, and the success rate of route construction. A comparison with a classical algorithm was also performed regarding solution quality and execution time.

1. Introduction

In modern infocommunication networks, characterized by high dynamics, complex topology, and increasing demands for transmission speed and reliability, the problem of optimal routing remains critically unresolved. Traditional routing algorithms often fail to adapt to real-time changes. They disregard contextual network features and have limited self-organizing capabilities [1]. This results in several negative theoretical consequences. There is no universal mathematical model that combines the topological flexibility of graph structures with the adaptability of neural networks. The potential of Graph Neural Networks (GNNs) in the context of routing is not sufficiently understood, especially their ability to generalize to new topologies. There is also a lack of research that integrates machine learning with classical theories of network management.

The unresolved nature of the optimal routing problem leads to decreased efficiency of data transmission in large and complex networks (such as IoT, 5G, and SDN). It results in increased delays, packet losses, and node overload caused by non-adaptive routes. It also makes it impossible to scale existing solutions to heterogeneous or dynamic network environments [2].

At the same time, modern machine learning approaches, particularly GNNs, open new possibilities for modeling network behavior, accounting for traffic characteristics, and predicting optimal routes in context. The distinctive feature of using GNNs in routing tasks lies in their ability to work directly with the graph structure of the network, where nodes represent network devices and edges represent physical or logical connections with specific metrics, such as delay. Unlike traditional methods, GNNs can learn from historical examples and generate routes that align with the global structure of the network. Although the resulting paths are not always strictly optimal in terms of the shortest route, their computation is significantly faster, which is critical for large, heavily loaded, or rapidly changing topologies. In such conditions, the priority shifts from absolute optimality to the ability to respond quickly and consistently to changes. A particularly promising approach is edge-level classification, where the neural network decides whether each edge should be included in the route between a pair of nodes.

The theoretical significance of the study lies in the development of new routing models based on Graph Neural Networks, which are capable of accounting for the structural properties of the network, the contextual relationships between nodes, and the dynamics of change. The proposed approach makes it possible to rethink classical notions of routing as a static process and transform it into an adaptive, learning system.

The applied significance is determined by the need for intelligent solutions in modern

infocommunication systems, where traditional methods no longer provide the required level of performance. The results of the study can be applied to optimize routing in SDN networks, mobile MANET networks, and IoT networks. They can also be used to improve the efficiency of real-time traffic management and to develop self-learning routing systems capable of adapting to changes without external intervention.

2. Analysis of related work and research problem formulation

In recent years, GNNs have seen active development in the field of infocommunications. In [3], the RouteNet model is examined, which employs GNNs to predict delay, jitter, and packet loss in SDN networks. It demonstrates the ability to generalize to new topologies and traffic patterns not represented in the training data. However, the analysis of scientific sources indicates a number of limitations that hinder their practical application.

In [4], the GraphSAGE model is proposed for predicting network load. Despite achieving high accuracy in static topologies, the model does not adapt to dynamic changes, which limits its applicability in mobile networks.

The authors of [5] investigate the use of Graph Attention Networks (GAT) for routing in SDN. Although the architecture is flexible, it requires significant computational resources, making real-time deployment challenging.

In [6], GNNs are applied in IoT networks. The model improves delay metrics but does not take energy consumption into account, which is a critical parameter in IoT environments.

An approach that combines GNNs with reinforcement learning demonstrates adaptability but becomes unstable when the topology changes, reducing its reliability [7].

For route classification, the authors of [8] employ graph autoencoders. The model is effective on simulation data but has not been tested on real networks.

In [9], GNNs are proposed for detecting high-traffic nodes. However, the model disregards contextual relationships between nodes, which leads to incorrect routing.

GNN-based routing in MANET networks is analyzed in [10]. Although performance improves, the model suffers from delays caused by weight coefficient calculations.

The authors of [7] apply clustering before routing based on GNNs. The clustering is performed without considering traffic dynamics, which reduces accuracy.

For traffic congestion prediction in urban networks, GNNs are used in [11]. The model is not scalable to global infocommunication systems.

An adaptive GNN model is proposed for SD-WAN in [12]. Although the performance is high, the model does not consider security aspects.

In [13], GNNs are explored for energy saving. The model does not adapt to topology changes, which limits its applicability.

In [14], a GNN-based model is developed for routing in hybrid networks. Despite its effectiveness in a test environment, the model does not account for delays caused by route changes.

The use of GNNs in the context of QoS-oriented routing is analyzed in [15]. The model does not provide sufficient accuracy under high traffic variability.

Thus, most existing solutions fail to consider the dynamics of topology changes, have limited scalability, are not adapted to real-time operation, and do not account for energy efficiency or security.

The relevant scientific problem is the development of an adaptive, scalable, and energy-efficient routing model based on Graph Neural Networks, capable of operating in real time, accounting for topology dynamics, and ensuring secure data transmission. This study is aimed at addressing this problem.

3. Research aim and objectives

The aim of this study is to model and analyze Graph Neural Networks for optimizing routing in infocommunication networks.

Achieving this goal will enable the use of Graph Neural Networks in real time within dynamic networks, where decision-making speed is critical, particularly under conditions of high topological complexity, unstable channels, and strict performance requirements. The proposed approach opens prospects for further research, especially in the areas of multi-criteria routing, adaptation to real-time metric changes, and integration with other components of network intelligence.

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

- construction and generation of graph data for training;
- formulation of features and target labels;
- development of the model architecture and training;
- analysis of performance and time characteristics.

4. Materials and Methods

The subject of this research is the application of Graph Neural Networks to the routing problem in infocommunication networks.

The object of the study is the process of route construction within an infocommunication network.

The research methods include the analysis of existing routing algorithms, a theoretical examination of the operating principles of Graph Neural Networks, the design of a model architecture based on the GENConv framework, as well as its experimental training and evaluation on a generated dataset of network topologies.

GNNs are a class of models capable of efficiently processing data represented as graphs, preserving both the structural information about relationships between entities and the local features of those entities. GNNs have gained widespread adoption in tasks where topology plays a critical role – particularly in routing, where the graph reflects the physical or logical structure of a network [16].

In GNNs, a graph is typically represented as a 5-tuple $G = (V, E, X, F, A)$, where V is the set of nodes, E is the set of edges, $X \in R^{|V| \times d}$ is the node feature matrix, $F \in R^{|E| \times k}$ is the edge feature matrix, and $A \in \{0,1\}^{|V| \times |V|}$ is the adjacency matrix defining the graph structure. Here, d denotes the node feature dimension, and k denotes the edge feature dimension. In the case of a weighted graph, the binary adjacency matrix A is replaced by a weight matrix $W \in R^{|V| \times |V|}$, where each element w_{uv} corresponds to a routing metric (e.g., delay) between nodes u and v [17].

Unlike classical neural networks, where each input sample is treated independently, Graph Neural Networks (GNNs) update the representation of each node by considering not only its own features but also the information received from its neighboring nodes in the graph.

This process is known as message passing, and it involves iterative exchange of information between adjacent nodes. At each layer of the model, every node $v \in V$ updates its state h_v using the aggregated information from its neighbors $u \in N(v)$:

$$h_v^{(l)} = \text{UPDATE}^{(l)}(h_v^{(l-1)}, m_v^{(l)}), \quad (1)$$

$$m_v^{(l)} = \text{AGGREGATE}^{(l)}(\{\phi(h_u^{(l-1)}, f_{uv}) : u = \overline{1, |V|}, u \in N(v)\}), \quad (2)$$

where $h_v^{(l)}$ is the hidden state of node v at layer l , $l = \overline{1, L}$, with L being the number of layers in the model, $v = \overline{1, |V|}$, ϕ is a message generation function that takes into account the features of

the neighbor u and the edge (u, v) , *AGGREGATE* and *UPDATE* are differentiable functions, which can be implemented, for instance, as mean, sum, or learnable transformations.

In addition to node features, edge features f_{uv} play a crucial role in GNNs, as they can encode routing metrics such as delay or bandwidth. These values can directly influence the message passing process, enabling the model to make more informed decisions about the relevance of information received from specific neighbors. These values can directly influence the message passing process, allowing the model to make more informed and selective decisions about the relevance of information received from individual neighbors.

Overall, GNNs provide a flexible framework for modeling dependencies in graphs, enabling each graph element (node or edge) to construct its representation based on both structure and local context.

In routing tasks, this allows models to learn from real traffic patterns or network topologies and to uncover patterns that are beyond the reach of traditional algorithms.

The first generations of Graph Neural Networks, such as Graph Convolutional Networks (GCN) and Graph Attention Networks (GAT), have served as the foundation for numerous models in graph-related tasks, including the processing of network topologies.

However, both architectures have limitations that significantly affect their effectiveness in complex, deep, or highly dynamic graphs – such as those found in telecommunications networks [18, 19].

GCN performs neighborhood feature aggregation using averaging weighted by the normalized adjacency matrix. The model updates are defined as follows [3]:

$$H^{(l+1)} = \sigma(\tilde{D}^{-1/2} \tilde{A} \tilde{D}^{-1/2} H^{(l)} W^{(l)}), \quad (3)$$

where $\tilde{A} = A + I$ is the adjacency matrix with added self-loops, $\tilde{D}$ is the diagonal degree matrix, $H^{(l)}$ are the hidden features, and $W^{(l)}$ are the layer- l learnable parameters.

Although this model is simple and effective, it does not account for differences between neighboring nodes. Moreover, deep architectures quickly encounter the problem of oversmoothing, where node representations become indistinguishable and lose their expressive power.

GAT addresses this limitation by introducing an attention mechanism – a dynamic approach that allows the model to independently determine which neighboring nodes are more important. For each edge (u, v) , an attention coefficient α_{uv} is computed, which reflects the importance of the message sent from node u to node v . These coefficients are normalized via the softmax function, enabling their interpretation as weights in the aggregation process [9]:

$$\alpha_{uv} = \frac{\exp(\text{LeakyReLU}(a^T [Wh_u || Wh_v]))}{\sum_{k \in N(v)} \exp(\text{LeakyReLU}(a^T [Wh_k || Wh_v]))}, \quad (4)$$

where *LeakyReLU* is a modified version of the standard *ReLU* activation function that allows a small negative gradient for negative inputs, typically with a slope coefficient of $\alpha \approx 0.2$. This helps to avoid the "dying neuron" problem where units stop learning, which is particularly important for computing attention coefficients that may take both positive and negative values.

The GENConv (Generalized Aggregation Graph Convolution) architecture, in turn, represents a more modern approach that combines the flexibility and expressiveness of attention mechanisms with the scalability and stability of classical models.

A key feature of GENConv is learnable aggregation – an aggregation function that is not fixed but is learned jointly with the rest of the model. For example, it can behave like a max

aggregator or a softmax-weighted average. This is achieved through a parameterized softmax function with temperature scaling:

$$m_{uv} = \frac{\exp(h_u/\tau)}{\sum_{k \in N(v)} \exp(h_k/\tau)}, \quad (5)$$

where τ is a learnable temperature parameter. A small τ results in behavior similar to max aggregation (where a single message dominates), while a large τ leads to uniform averaging.

Thus, GENConv enables the model to autonomously adapt to the local graph structure and determine the most appropriate aggregation strategy depending on the number of neighbors, feature variability, and topological structure.

The feature update formula in GENConv is defined as:

$$h_v^{(l)} = MLP^{(l)} \left(h_v^{(l-1)} + AGG \left(\left\{ \frac{h_u^{(l-1)}}{\tau} : u \in N(v) \right\} \right) \right), \quad (6)$$

where MLP denotes a nonlinear transformation with normalization and dropout, AGG is the learnable aggregator, and the addition implements a residual connection.

GENConv demonstrates stable performance even with a large number of layers (30+), whereas models like GCN and GAT typically suffer from performance degradation. This is particularly important in the context of routing tasks in telecommunication networks, where relevant information about the optimal path may be located at a considerable distance within the topology, and the model must be able to accurately capture it [18].

In classical tasks addressed by Graph Neural Networks, the primary objective is often node classification or link prediction. However, in the task of constructing a route between a given source-target node pair, edge classification becomes central – specifically, determining which edges should be included in the route.

Unlike the node-level approach, where the model generates a vector for each node and makes decisions based on it, in edge-level classification, vectors are constructed for edges. In this case, each edge $(u, v) \in E$ is represented as a combination of node features h_u, h_v , edge features f_{uv} , and additional information that may be specific to the given source–target pair. The input to the edge classifier takes the form of a concatenation:

$$z_{uv} = [h_u \parallel h_v \parallel f_{uv} \parallel d_u^s \parallel d_v^t], \quad (7)$$

where h_u, h_v are the hidden representations of the nodes after passing through the GNN, f_{uv} are the edge features, and d_u^s, d_v^t are heuristic distances to the source and target nodes, respectively, e.g., precomputed shortest-path distances, which add routing-specific context to the classification task.

The edge-level approach allows the model to frame the routing task as a binary classification problem, where each edge is evaluated to determine whether it belongs to the path between a given source s and target t . Based on these predictions, the model can construct a route that aligns with the graph context, taking into account both local features and the global structure.

In the next step, the vector is passed to a decoder implemented as a Multi-Layer Perceptron (MLP). The perceptron consists of several fully connected layers, each applying a linear transformation, an activation function (e.g., ReLU), as well as dropout or normalization. The final layer typically contains a single neuron with a sigmoid activation function, which outputs the probability that the edge is part of the route between s and t .

The edge classification approach using an MLP enables the model to simultaneously account for the topological context, connection characteristics, and relevance to the routing task. Unlike

models that attempt to navigate step-by-step from node to node, the edge-level approach allows for constructing a complete view of the route in the form of either a binary edge mask or a probabilistic field.

An advantage of the edge-level approach is its flexibility and accuracy compared to the node-level paradigm. The model can take into account the specific characteristics of each edge, as well as the interaction between the connected nodes.

5. Research results

5.1. Graph data construction and generation for training

To train the Graph Neural Network (GNN) tailored for solving routing tasks in complex and dynamic telecommunication networks, a custom synthetic dataset was constructed. In total, 3000 graphs were generated, each simulating the topology of a medium- or large-scale network with a variable number of nodes, structural heterogeneity, and diverse connection characteristics.

The graphs were generated using the Python programming language and the NetworkX library [19,20,21], which provides tools for creating random graphs and working with topological structures. The number of nodes in each graph was randomly selected within the range of 50 to 100, allowing for both simple and structurally rich scenarios. Graph construction employed the `networkx.gnm_random_graph` generator, which creates a graph based on a fixed number of nodes and a randomly chosen number of edges in the range from $N+10$ to $3N$, where N is the number of nodes. Connectivity was enforced as a mandatory condition and verified using `nx.is_connected`.

Each edge in the graph was assigned a transmission delay, represented as a random value in the range from 0.001 to 2.0 seconds. This was implemented using the `random.uniform` function, which ensured a uniform distribution of delays within the specified interval. Such modeling makes it possible to simulate both high-speed communication links and slow or congested lines, thereby providing sufficient diversity in training conditions.

For each graph, a random pair of nodes was selected to serve as the source and target. Based on the edge weights, the shortest path between them was computed using Dijkstra's algorithm as implemented in the NetworkX library. This path was then used to generate binary labels: edges belonging to the path were marked as positive (label 1), while all others were marked as negative (label 0). In this way, the task was formulated as an edge classification problem in the context of route construction.

All graphs (Figs. 1, 2) were converted into a format compatible with the PyTorch Geometric library [21]. Each sample includes an edge index matrix (`edge_index`), edge features (`edge_attr`), node features (`x`), and labels (`edge_label`). Node features consisted of binary indicators specifying whether the node is the source or target, the normalized degree, and the number of hops to the target node (Table 1). Edge features were represented by the delays assigned during graph generation.

Table 1

Structure in .pt for PyTorch Geometric

Field	Description	Dimension
<code>x</code>	Node features	$[N,d]$
<code>edge_index</code>	Indices of node pairs for each edge	$[2,E]$
<code>edge_attr</code>	Edge features (delay)	$[E,k]$
<code>edge_label</code>	Binary labels: whether the edge belongs to the route	$[E]$
<code>source</code>	Index of the source node	$[1]$
<code>target</code>	Index of the target node	$[1]$

Graph 12 — Source: 6, Target: 91

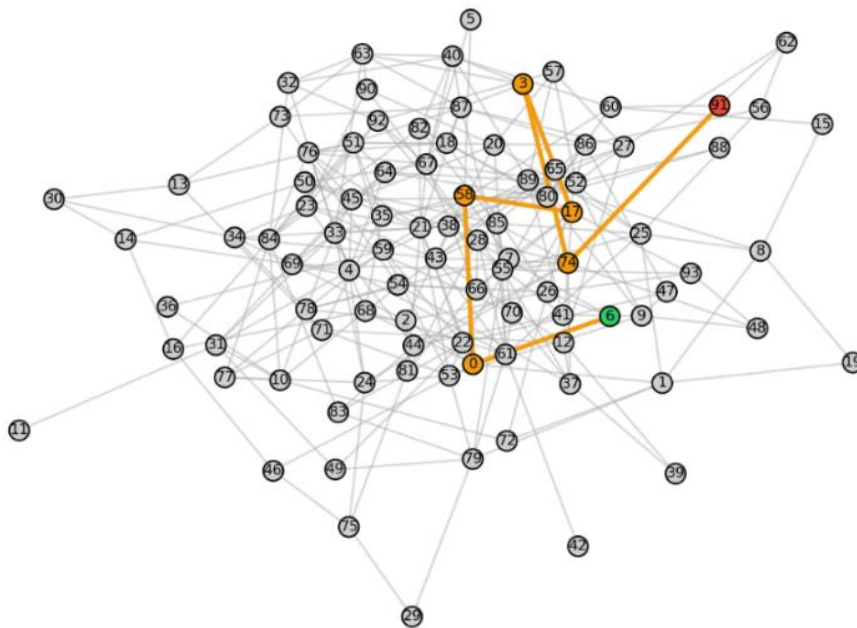


Fig. 1. Example of generated graph 1

Graph 7 — Source: 4, Target: 20

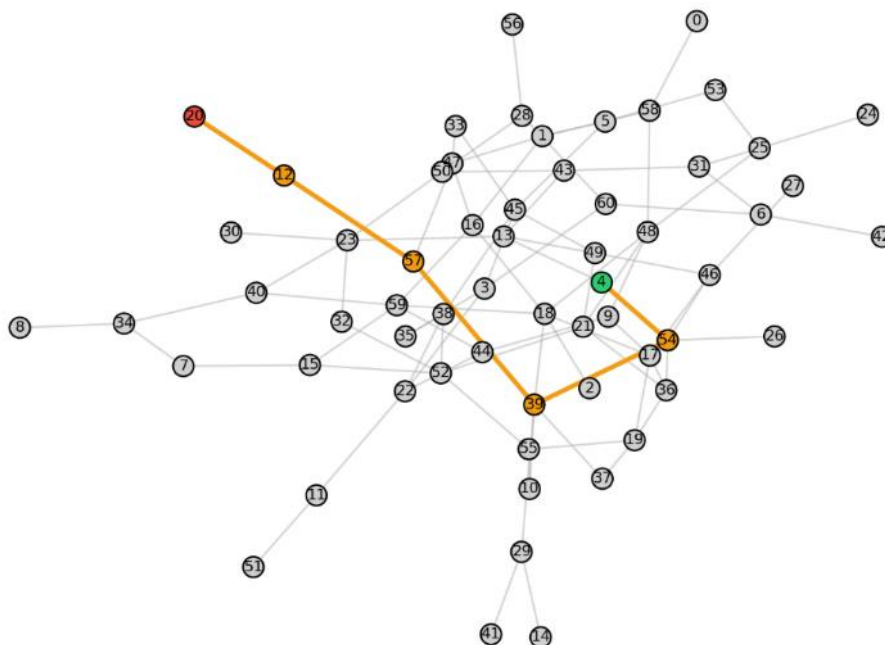


Fig. 2. Example of generated graph 2

5.2. Feature and target label definition

Since the model is tasked with classifying edges in the context of a route between a given pair of nodes, it was critically important to design input features that not only capture local properties of graph elements but also provide context specific to the routing task.

To this end, several types of features were selected for both nodes and edges, enabling the model to effectively localize the source and target, infer the routing direction, and account for the quality of connections.

Node features include binary indicators `is_source` and `is_target`, which are essential for enabling the model to distinguish the currently active source-target pair (s, t) . Without this information, the routing task would be ill-defined, since the model processes the graph as a whole and must explicitly know where the route starts and ends. These features, therefore, provide task-specific addressability.

The normalized node degree is a simple yet important topological feature that provides the model with insight into the local connectivity of a node. A high degree is often associated with key routing nodes, while peripheral nodes typically have limited capacity for traffic forwarding.

Additionally, heuristic meta-information is introduced in the form of the number of hops from a node to the target. This value is computed based on the shortest path without considering edge weights and provides a general notion of the node's distance from the destination.

Although this feature is not precise in a weighted graph, it helps the model orient itself toward the target by combining local decisions with global context.

As for edges, the sole but crucial feature is latency – a value that directly corresponds to the routing metric. In real-world networks, latency or its derivatives are typically used as the primary criterion for path optimization, making the inclusion of this feature essential for effective model training.

Training labels are generated based on the shortest path between nodes s and t , computed using Dijkstra's algorithm. This approach frames the learning task as binary edge classification, where the positive class corresponds to edges that belong to the optimal path, and the negative class to all others. Such formulation scales well and enables the use of standard neural network optimization techniques for a pathfinding task that would otherwise be combinatorially complex.

Thus, the constructed system of features and labels enables the model to simultaneously account for local graph parameters, the global routing objective, and the network metric to be optimized. This establishes a logical bridge between the graph's topological structure and the practical requirements of route construction in a network.

5.3. Model architecture and training

The developed model (Table 2) implements an edge-level classification approach using a Graph Neural Network based on GENConv convolutional layers. The architecture is designed to enable efficient information propagation between graph nodes and accurate evaluation of each edge in the context of a specific routing task. At the core of the model is the EdgeEncoderOptimized module, which consists of three sequential GENConv layers and a single MLP decoder that takes as input the constructed feature vector for each edge.

Each GENConv layer performs message aggregation from neighboring nodes using softmax normalization with a learnable temperature coefficient, allowing the aggregation behavior to adapt from uniform to selective. After each convolutional layer, BatchNorm1d is applied to stabilize training, and Dropout is used to prevent overfitting. All hidden layers have a dimensionality of 64.

As a result, each node in the graph obtains a generalized representation that incorporates not only its own features but also the context of its local neighborhood over multiple hops.

Table 2

Neural network architecture		
Component	Type	Parameters
conv1	GENConv	in: 4, out: 64, aggr: softmax
bn1	BatchNorm1d	64
conv2	GENConv	in: 4, out: 64, aggr: softmax
bn2	BatchNorm1d	64
conv3	GENConv	in: 4, out: 64, aggr: softmax
bn3	BatchNorm1d	64
dropout	Dropout	p = 0,3
edge_mlp	MLP decoder (3 layers)	193 → 64 → 32 → 1, ReLU + BatchNorm

As a result, a vector of size 193 is formed and passed to the MLP decoder, which consists of three fully connected layers:

- First layer: Linear(193 → 64) with ReLU activation and BatchNorm;
- Second layer: Linear(64 → 32) with ReLU;
- Third layer: Linear(32 → 1) without activation (output is a logit for subsequent sigmoid processing).

ReLU activations enable the model to learn nonlinear dependencies between structural and contextual parameters, while BatchNorm after the first layer ensures stability of the activation distribution and accelerates convergence.

For training, the generated data was divided into a training set (80% of the dataset) and a validation set (20% of the dataset). Binary cross-entropy with logits was used as the loss function during training, incorporating class imbalance through weighted scaling. Optimization was performed using the Adam algorithm with a fixed learning rate of 10^{-3} . The batch size was set to 1 (one graph per iteration), and the model was trained for 35 epochs. Performance evaluation relied on metrics such as loss, accuracy, and the success rate of correctly reconstructed routes.

Figure 3 shows the loss function, which steadily decreases on the training data and remains well-controlled on the validation set.

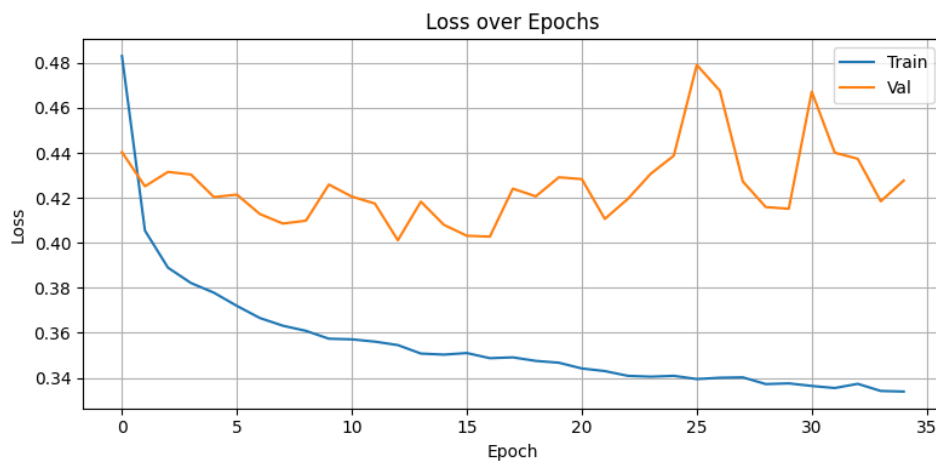


Fig. 3. Loss function plot on training and validation data

Figure 4 shows the edge classification accuracy, where the validation accuracy consistently exceeds 94%, occasionally reaching 95-96%, indicating reliable generalization. Figure 5 presents detailed results from the final epochs, including all key metrics. The most indicative are the high route construction success rate and stable accuracy.

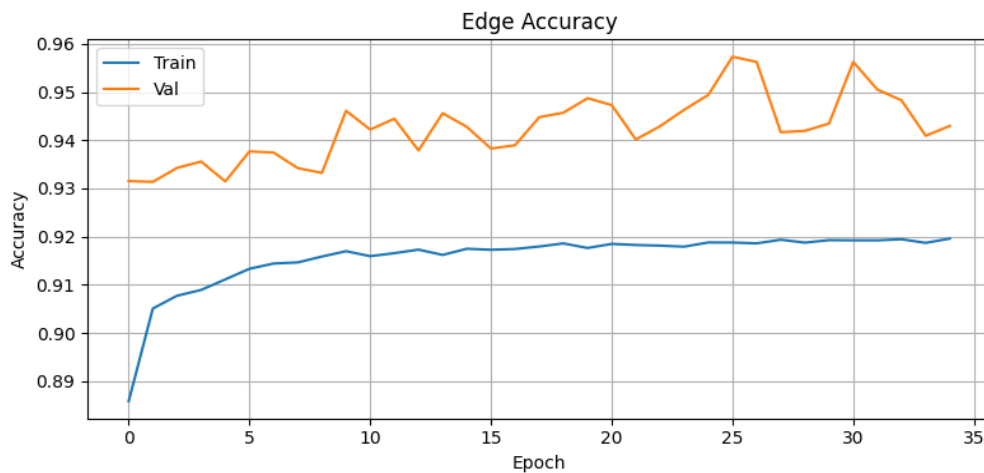


Fig. 4. Accuracy plot on training and validation data

Epoch 29	Train Loss: 0.3372	Acc: 0.919	Lat: 2.6501	Success: 99.6%
	Val Loss: 0.4159	Acc: 0.942	Lat: 2.3910	Success: 91.2%
Epoch 30	Train Loss: 0.3375	Acc: 0.919	Lat: 2.6550	Success: 99.7%
	Val Loss: 0.4152	Acc: 0.943	Lat: 2.3262	Success: 90.5%
Epoch 31	Train Loss: 0.3364	Acc: 0.919	Lat: 2.6426	Success: 99.5%
	Val Loss: 0.4672	Acc: 0.956	Lat: 2.1974	Success: 84.5%
Epoch 32	Train Loss: 0.3355	Acc: 0.919	Lat: 2.6598	Success: 99.7%
	Val Loss: 0.4401	Acc: 0.951	Lat: 2.3071	Success: 88.5%
Epoch 33	Train Loss: 0.3373	Acc: 0.919	Lat: 2.6595	Success: 99.8%
	Val Loss: 0.4374	Acc: 0.948	Lat: 2.2853	Success: 87.0%
Epoch 34	Train Loss: 0.3342	Acc: 0.919	Lat: 2.6294	Success: 99.8%
	Val Loss: 0.4185	Acc: 0.941	Lat: 2.4160	Success: 92.7%
Epoch 35	Train Loss: 0.3339	Acc: 0.920	Lat: 2.6395	Success: 99.5%
	Val Loss: 0.4278	Acc: 0.943	Lat: 2.3346	Success: 90.3%

Fig. 5. Training results during the final epochs

Thus, the developed model demonstrates consistently high classification performance, is capable of constructing coherent routes under varying graph complexities, and shows effective generalization even in the presence of noise and topological variability. The combination of GENConv and the MLP decoder has proven to be effective for edge-level routing tasks.

5.4. Performance and timing analysis

To evaluate not only classification accuracy but also the quality of the constructed routes, a visualization was performed comparing the predicted paths with the reference (shortest) routes for a set of graphs.

Figures 6 and 7 show examples of constructed paths, where the predicted route is highlighted in yellow and the optimal (shortest) route in blue. In the first case (Fig. 6), the model precisely reproduced the path found by Dijkstra's algorithm, with an identical total delay of 2.2731 (Fig. 8). In the second example (Fig. 7), the model's route differs, it has fewer hops but slightly higher total delay (Fig. 9). This result indicates that the model is capable of discovering alternative paths that are close to optimal or represent a trade-off between different metrics.

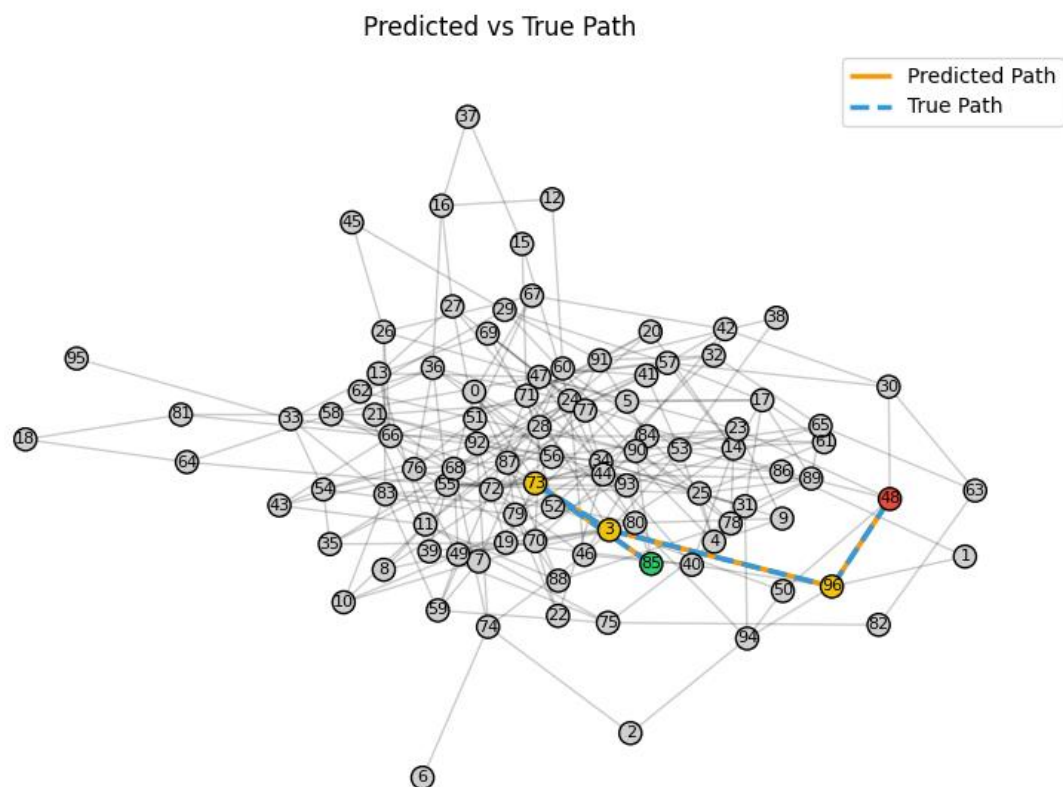


Fig. 6. Comparison of predicted and optimal routes, example 1

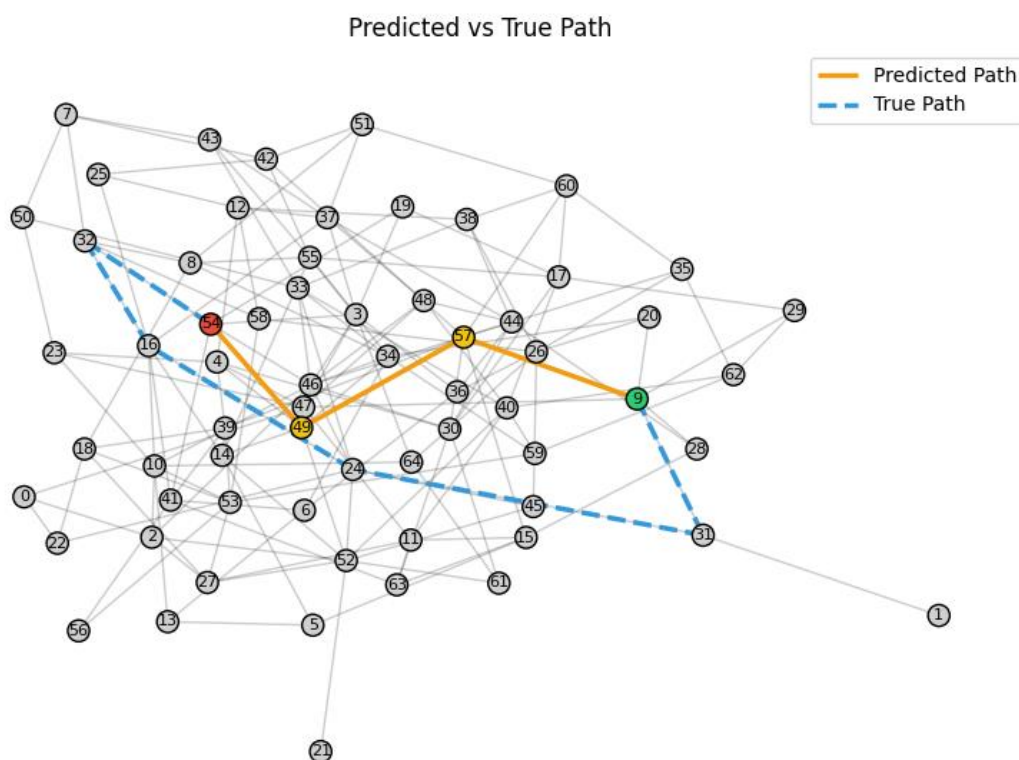


Fig. 7. Comparison of predicted and optimal routes, example 2

```

Predicted path: [85, 73, 3, 96, 48]
Ground truth : [85, 73, 3, 96, 48]
Path length   : 5 (true: 5)
Pred latency: 2.2731
True latency  : 2.2731

```

Fig. 8. Comparison metrics for case 1

```

Predicted path: [9, 57, 49, 54]
Ground truth : [9, 31, 24, 16, 32, 54]
Path length   : 4 (true: 6)
Pred latency: 1.9908
True latency  : 1.8471

```

Fig. 9. Comparison metrics for case 2

Figure 10 presents a comparison of delays for each of the 500 tested graphs, showing both the true latency and the latency of the routes predicted by the model.

Despite the structural variability of the topologies, the two curves exhibit a strong correlation. This indicates that even when the predicted path differs from the optimal one, the model can reproduce the delay with sufficient accuracy, staying within an acceptable margin of deviation.

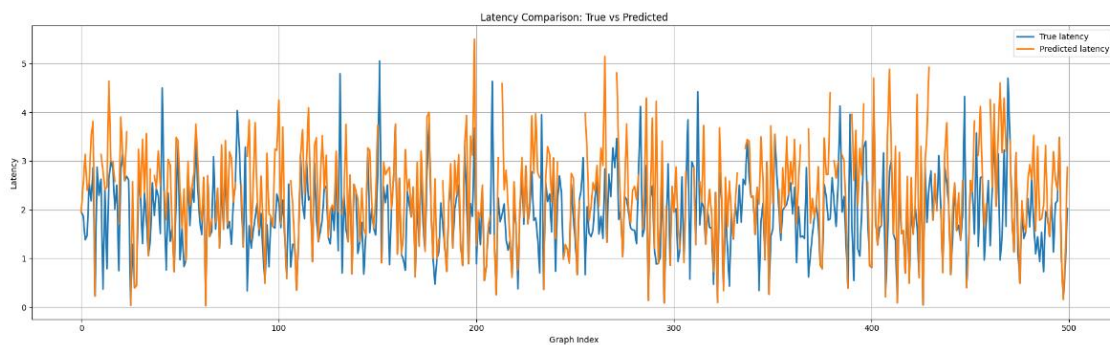


Fig. 10. Comparison of predicted and optimal paths by total latency on the test set

A performance comparison between the GNN model and the classical routing approach is presented in Figure 11. It shows the average route construction time for both Dijkstra's algorithm and the GNN model in batch inference mode. In graphs with 50–100 nodes, the model demonstrated approximately twice the throughput. The experiment was conducted in the Google Colab cloud environment with the following specifications:

- processor: Intel(R) Xeon(R) CPU @ 2.20GHz, 2 logical cores, 56 MB cache;
- RAM: ~13 GB;
- GPU: NVIDIA Tesla T4, 16 GB VRAM, CUDA 12.4.

This result is attributed to the fact that the model does not need to reconstruct the path from scratch for each query – it operates as a universal router for the entire topology.

In addition to speed, a key advantage of the neural network – based approach is its ability to instantly respond to changes in the network structure without re-running algorithmic search. This is particularly important in dynamic environments, where the number of nodes or edge weights may change in real time.

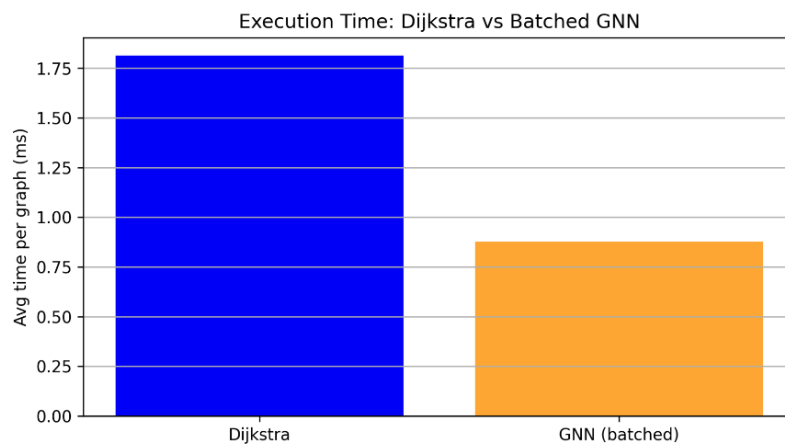


Fig. 11. Comparison of average routing time between Dijkstra's algorithm and the GNN model

6. Discussion of research results

This research has confirmed that graph neural networks based on the GENConv architecture can effectively address routing problems in complex network topologies. The model demonstrated high edge-level classification accuracy, exceeding 94% on the test sets, indicating its ability to generalize knowledge and adapt to novel graph scenarios.

Notably, in some of the tested graphs, the model generated routes that did not exactly match the classically computed shortest paths but exhibited similar or even better topological properties, particularly a reduced number of hops. This suggests that the model is capable of adaptively selecting trade-offs between metrics such as latency, path length, and reliability, which is especially relevant in dynamic network environments.

The analysis of temporal characteristics revealed the advantages of the GNN-based approach during inference. The model demonstrates twice the performance compared to Dijkstra's algorithm on graphs with 50-100 nodes, which is explained by the absence of the need to recompute the route for each query. This enables rapid response to changes in network topology – a critical requirement for modern infocommunication environments.

The GENConv architecture with learnable aggregation ensures stability at greater network depth and flexible adaptation to the local structure of the graph. Even in cases of topological variability or noise, the model demonstrates consistent performance by generating coherent routes. The inclusion of additional heuristic information (such as the number of hops to the target node) further enhances the model's ability to navigate the global context of the route.

Despite the high performance, questions remain open regarding scalability to ultra-large topologies and improving interpretability for integration into critical systems. The proposed edge-level approach opens up promising directions for further research – particularly in the areas of multi-objective routing, training on real-world networks, real-time metric adaptation, and integration with other components of SDN infrastructure.

7. Conclusions

Within the scope of the research objective, a routing model based on a Graph Neural Network with edge-level classification was developed, implemented, and experimentally investigated. The model was built on the GENConv architecture, and the following tasks were addressed.

A custom synthetic dataset was created for training the Graph Neural Network aimed at solving the routing problem in complex and dynamic telecommunication networks. In total, 3000 graphs were generated, each modeling the topology of a medium- or large-scale network with a variable number of nodes, structural heterogeneity, and diverse connection characteristics. The

graphs were generated using the Python programming language and the NetworkX library, which provides tools for generating random graphs and working with topological structures.

Several types of features were selected for both nodes and edges to enable the model to effectively localize the source and target, navigate the route direction, and account for connection quality. The node features include the binary indicators `is_source` and `is_target`, which are necessary for the model to distinguish the currently active pair of nodes.

A model was developed that implements the edge-level classification approach using a Graph Neural Network based on GENConv convolutional layers. The architecture enables efficient information propagation between graph nodes as well as accurate evaluation of each edge in the context of a specific routing task. At the core of the model lies the EdgeEncoderOptimized module, which consists of three sequential GENConv layers and one MLP decoder that processes the generated feature for each edge.

A performance comparison was conducted between the graph-based model and Dijkstra's algorithm. It was found that in inference mode, the graph model operates significantly faster, especially on large graphs, and does not require re-exploring the entire route space for each query. This makes the proposed approach suitable for real-time use in dynamic networks, where decision-making speed is critical.

The obtained results confirm that Graph Neural Networks can serve as an effective alternative to traditional routing algorithms, particularly under conditions of high topological complexity, channel instability, and strict performance requirements. The proposed approach opens new perspectives for further research, especially in the areas of multi-criteria routing, adaptation to real-time metric changes, and integration with other components of network intelligence.

References:

1. D. Medhi & K. Ramasamy, *Network routing. Algorithms, Protocols, and Architectures*. San Francisco, CA, USA: Morgan Kaufmann Publishers Inc., 2007. doi: 10.1016/b978-0-12-088588-6.x5000-1.
2. О. В. Лемешко, О. С. Єременко, М. О. Євдокименко, А. С. Шаповалова та Б. Слейман, *Моделювання та оптимізація процесів безпечної та відмовостійкої маршрутизації в телекомунікаційних мережах: монографія*. Харків, Україна: ХНУРЕ, 2022.
3. K. Rusek, J. Suárez-Varela, P. Almasan, P. Barlet-Ros and A. Cabellos-Aparicio, "RouteNet: Leveraging Graph Neural Networks for Network Modeling and Optimization in SDN," *IEEE Journal on Selected Areas in Communications*, vol. 38, no. 10, pp. 2260-2270, Oct. 2020. doi: 10.1109/JSAC.2020.3000405.
4. J. Zhou, G. Cui, S. Hu, Z. Zhang, C. Yang, Z. Liu et al, "Graph neural networks: A review of methods and applications," *AI Open*, vol. 1, pp. 57–81, 2020. doi: 10.1016/j.aiopen.2021.01.001.
5. Z. Wu, S. Pan, F. Chen, G. Long, C. Zhang and P. S. Yu, "A Comprehensive Survey on Graph Neural Networks," *IEEE Transactions on Neural Networks and Learning Systems*, vol. 32, no. 1, pp. 4-24, Jan. 2021. doi: 10.1109/TNNLS.2020.2978386.
6. H. Kim, J. Park, and Y. Lee, "Routing optimization in IoT networks using graph neural networks," *Sensors*, vol. 22, no. 3, pp. 1–15, 2022.
7. J. Ramirez and M. Ortega, "Cluster-aware graph neural routing for dynamic networks," *IEEE Internet of Things Journal*, vol. 11, no. 2, pp. 2345–2357, Feb. 2024.
8. Y. Ln, X. Wang, and J. Liu, "Graph autoencoders for route classification in software-defined networks," *IEEE Access*, vol. 11, pp. 45678–45689, 2023.
9. M. Ahmed, T. Rahman, and S. Chowdhury, "Traffic hotspot detection using graph neural networks," *Journal of Network and Computer Applications*, vol. 202, pp. 1–12, 2024.
10. L. Chen, K. Zhao, and Y. Sun, "GNN-based routing in mobile ad hoc networks," *Ad Hoc Networks*, vol. 145, pp. 102–115, 2024.
11. T. Tanaka, H. Saito, and K. Fujimoto, "Urban congestion prediction using graph neural networks," *IEEE Transactions on Intelligent Transportation Systems*, vol. 26, no. 4, pp. 789–798, Apr. 2025.
12. D. Ivanov and P. Smirnov, "Adaptive GNN-based routing for SD-WAN architectures," *IEEE Communications Letters*, vol. 29, no. 5, pp. 1123–1127, May 2025.

13. M. Kowalski, A. Nowak, and E. Zielinska, "Energy-efficient routing in wireless sensor networks using GNN," *Sensors*, vol. 25, no. 6, pp. 1–14, 2025.
14. Y. Zhang, L. Xu, and H. Li, "Graph neural routing in hybrid network topologies," *IEEE Transactions on Network and Service Management*, vol. 18, no. 3, pp. 456–468, Mar. 2025
15. V. Petrov, N. Koval, and D. Kravets, "QoS-aware routing with graph neural networks in heterogeneous networks," *IEEE Systems Journal*, vol. 19, no. 2, pp. 987–996, Apr. 2025
16. W. L. Hamilton, *Graph representation learning*. Springer Cham, 2020. doi: 10.1007/978-3-031-01588-5
17. Z. Wu, S. Pan, F. Chen, G. Long, C. Zhang and P. S. Yu, "A Comprehensive Survey on Graph Neural Networks," *IEEE Transactions on Neural Networks and Learning Systems*, vol. 32, no. 1, pp. 4–24, Jan. 2021, doi: 10.1109/TNNLS.2020.2978386
18. "A gentle introduction to graph neural networks," Distill. [Online] Available: <https://distill.pub/2021/gnn-intro/> Accessed on: July, 26, 2025.
19. "NetworkX documentation," NetworkX. [Online] Available: <https://networkx.org/> Accessed on: July, 26, 2025.
20. "PyG Documentation," Pytorch_geometric documentation. [Online] Available: <https://pytorch-geometric.readthedocs.io/en/latest/> Accessed on: July, 26, 2025.
21. S. Shtangey, L. Melnikova and O. Sokolov, "Modeling and analysis of graph neural networks for routing optimization in infocommunication networks," Zenodo, Jul. 30, 2025. doi 10.5281/zenodo.16616697.

Надійшла до редколегії 20.08.2025 р.

Shtangey Svitlana Viktorivna, Ph.D. in Technical Sciences, Associate Professor, Associate Professor at the Department of Infocommunication Engineering V.V. Popovsky, Kharkiv National University of Radio Electronics, Kharkiv, Ukraine, e-mail: Svitlana.shtanhei@nure.ua, ORCID: <https://orcid.org/0000-0002-9200-3959>. (Academic advisor of the higher education applicant Sokolov O.K.).

Melnikova Lubov Ivanivna, PhD in Technical Sciences, Associate Professor, Associate Professor of the Department of Infocommunication Engineering V.V. Popovsky, Kharkiv National University of Radio Electronics, Kharkiv, Ukraine, e-mail: liubov.melnikova@nure.ua, ORCID: <https://orcid.org/0000-0003-0439-7108>

Marchuk Artem Volodymyrovych, PhD in Technical Sciences, Associate Professor, Associate Professor of the Department of Infocommunication Engineering V.V. Popovsky, Kharkiv National University of Radio Electronics, Kharkiv, Ukraine, e-mail: artem.marchuk@nure.ua, ORCID: <https://orcid.org/0000-0002-2720-3954>

Linnyk Olena Vyacheslavivna, PhD in Technical Sciences, Associate Professor, Associate Professor of the Department of Mechanical and Biomedical Engineering, National Technical University "Dnipro Polytechnic", Dnipro, Ukraine, e-mail: linnyk.ol.o@nmu.one, ORCID: <https://orcid.org/0000-0002-4906-3796>

Sokolov Oleksandr Kuokovych, higher education applicant, group IST-23-1, Faculty of Infocommunications, Kharkiv National University of Radio Electronics, Kharkiv, Ukraine, e-mail: oleksandr.sokolov1@nure.ua, ORCID: <https://orcid.org/0009-0005-5382-2774>

УДК 004.75

DOI: 10.30837/0135-1710.2025.186.055

*В.В. ПРОСОЛОВ, Г.З. ХАЛІМОВ, П.В. ШУЛІК, А.О. СМІРНОВ, Д.О. В'ЮХІН***ВИКОРИСТАННЯ МЕТОДІВ МАШИННОГО НАВЧАННЯ ДЛЯ ВИЯВЛЕННЯ АТАК НА БЛОКЧЕЙН-СИСТЕМИ**

Предметом дослідження є методи виявлення атак у мережах із консенсусом Proof-of-Stake (PoS). Мета роботи – експериментальне дослідження та аналіз ефективності класичних алгоритмів машинного навчання для виявлення шкідливих вузлів у блокчейн-системах. Завдання: аналіз вразливостей технології блокчейн, створення та використання спеціалізованого набору даних для мереж PoS, а також побудова й тестування моделей машинного навчання. Основну увагу приділено порівнянню трьох алгоритмів – Random Forest, Support Vector Machine та k-Nearest Neighbors – для визначення їхньої придатності для моніторингу активності вузлів та виявлення аномалій. Для виконання поставлених завдань застосовано методи: моделювання, емпіричні та математичні методи. Моделювання полягало у програмній реалізації вибраних алгоритмів і подальшому аналізі їхніх результатів із використанням метрик точності, повноти, F1-міри та матриці плутанини. Емпіричні методи реалізовано шляхом тестування моделей на напівсинтетичному датасеті, який містить понад 10 тисяч записів про вузли та транзакції блокчейну. Математичні методи передбачають обчислення статистичних показників ефективності, а також аналіз інформативності ознак, що визначають поведінку вузлів.

Досягнуті результати: апробовано датасет для блокчейнів PoS, що включає ключові операційні параметри транзакцій і вузлів, сформовані пропозиції щодо подальшого використання моделей машинного навчання, здійснено тестування моделей машинного навчання.

Висновки. Доведено що машинне навчання є ефективним інструментом для ідентифікації аномалій та шкідливої активності у блокчейн-системах. Отримані результати закладають підґрунтя для подальших досліджень, які можуть бути спрямовані на розширення ознакового простору, інтеграцію глибоких нейронних мереж, розробку ансамблевих підходів та адаптацію методів до різних типів блокчейнів.

1. Вступ

Технологія блокчейну забезпечує незмінність записів і децентралізовану валідацію транзакцій, давно вийшовши за межі криптовалют до смарт-контрактів і DApps. Попри криптографічну стійкість, публічні мережі залишаються вразливими до низки атак – подвійного витрачання, атак 51% та їх варіантів (Finney, Vector-76), що підтверджується реальними інцидентами (зокрема в Ethereum Classic (ETC), Bitcoin) і призводить до суттєвих збитків [1], [2]. Для підвищення стійкості все активніше застосовують методи машинного навчання (ML) для виявлення аномалій та шкідливої активності в режимі, наближеному до реального часу: від класичних алгоритмів – Random Forest (RF), Support Vector Machine (SVM), k-Nearest Neighbors (k-NN) – до глибоких моделей (згортоква нейронна мережа (Convolutional Neural Network (CNN), рекурентна нейронна мережа (Recurrent Neural Network (RNN)) та RL-підходів [3], [4], [5]. Окрему увагу приділяють технікам збереження конфіденційності (диференційна приватність (differential privacy, DP), захищені багатосторонні обчислення (Secure Multi-Party Computation, SMPC)), що дозволяють аналіз без розкриття чутливих даних [6].

Глобальна актуальність теми зумовлена стрімкою цифровізацією економіки та зростанням вартості активів, що циркулюють у публічних і консорціумних блокчейн-системах. Масштабування мереж, поява складних DeFi-сценаріїв і смарт-контрактів підвищують площину атак і потенційні системні ризики, тоді як традиційні засоби контролю не забезпечують потрібної швидкодії та пояснюваності. Умови обмеженої

довіри, відсутність централізованого посередника та вимоги до конфіденційності ускладнюють класичний аудит і оперативне реагування. На цьому тлі методи ML дають змогу будувати відтворювані процедури виявлення аномалій і шкідливої активності в режимі, наближеному до реального часу, без втручання у консенсус. Поєднання ML із техніками збереження приватності (DP/SMPC) створює підґрунтя для безпечної аналітики між учасниками мережі. Отже, розроблення надійних, інтерпретованих та масштабованих ML-рішень для моніторингу блокчейн-мереж є важливим глобальним завданням забезпечення кібербезпеки та фінансової стійкості.

2. Огляд сучасних наукових публікацій

Розвиток блокчейн-систем супроводжується зростанням площини атак і появою нових векторів зловмисної активності. Огляди публікацій за останні роки фіксують зміщення акцентів від загальних питань DLT-архітектур до прикладних методів моніторингу та автоматизованого виявлення вторгнень у публічних і консорціумних мережах [7], [8], [9]. Ключова ідея полягає у використанні ML й аналізу аномалій для перетворення реєстрових і телеметричних даних на рішення про ризики в реальному часі, із дотриманням обмежень на приватність і втручання в консенсус.

Класи атак на публічних блокчейнах (double spending, 51 % (Goldfinger), Finney, Vector-76) добре формалізовані в працях із безпеки Bitcoin та суміжних мереж [10], [11], [12]. Практичні інциденти – від компрометації MtGox і DAO до багатоепізодних атак на ETC – засвідчили, що навіть зрілі мережі залишаються чутливими до перехоплення обчислювальної/економічної переваги або до експлуатації специфіки протоколів [1], [2]. На цьому тлі методи раннього виявлення аномалій розглядаються як необхідний додаток до протокольних контрзаходів і економічних стимулів.

Лінія робіт з кібербезпеки на основі ML демонструє ефективність як контрольованих, так і неконтрольованих підходів. Класичні алгоритми – SVM, RF, k-NN – застосовують для виявлення відомих шаблонів зловмисної поведінки та для побудови бінарних класифікаторів транзакцій/вузлів [4], [13]. Стандартні довідники з ML підкреслюють залежність якості таких моделей від інженерії ознак та якості маркованих даних [14]. Для великих вибірок, характерних для блокчейнів, розглядають обчислювально ощадливі варіанти k-NN і прийоми прунінгу, що знижують вартість інференсу без істотної втрати якості [15], [16].

Неконтрольовані методи та детектори аномалій актуальні там, де марковані дані обмежені або атаки – «нульового дня». Узагальнюючий огляд [9] систематизує застосування кластеризації (k-means, DBSCAN) та моделей на кшталт Isolation Forest/one-class підходів для виявлення рідкісних відхилень у поведінці вузлів і потоках транзакцій. Перевага цих підходів – пластичність до нових патернів; обмеження – потреба в тонкому тюнінгу порогів і чутливість до зміни середовища.

Глибоке навчання все частіше інтегрують у стек безпеки блокчейнів. Роботи з кібербезпеки демонструють переваги CNN/RNN для аналізу структурованих графів транзакцій і часових рядів активності [3]. У прикладних сценаріях (енергетичні мережі, обмін активами) запропоновані фреймворки на стику машинного навчання і блокчейну показують підвищення чутливості до аномалій за прийнятною затримки [5]. Для індустриального IoT перспективними визнані AI-механізми довіри та адаптивні політики реагування, що поєднують глибокі моделі з блокчейн-верифікацією [17].

Окрему нішу займають генеративні моделі. Generative Adversarial Networks (GAN) використовують для отримання реалістичних синтетичних сценаріїв атак, що розширюють навчальні вибірки детекторів і дозволяють відпрацьовувати рідкі події без появи небезпечних кейсів у продакшені [18]. Недоліки – висока обчислювальна вартість і

залежність від якості базових даних.

Питання конфіденційності та регуляторної сумісності стимулюють застосування методів збереження приватності (DP, SMPC). Децентралізовані протоколи аналізу, які не потребують розкриття сирих даних, доводять практичність інтеграції аналітики в чутливі домени – наприклад, охорона здоров'я й фінанси на блокчейні [6]. Такі методи є ключовими для PoS/DPoS-мереж, де відкритий агрегаційний моніторинг може конфліктувати з вимогами приватності валідаторів.

На рівні протоколів триває переосмислення властивостей консенсусу. Від класичної моделі BFT до практик PoS/DPoS – у наукових публікаціях підкреслено тонкі компроміси між живучістю, безпекою та продуктивністю [19], а також пропонується покращення голосувальних механізмів і стейкінгових правил для зниження ймовірності маніпуляцій [20]. Для PoS-ланцюгів окремо виділяють довготривалі (long-range) атаки й методи їхнього виявлення за допомогою глибоких моделей, навчених на історії реєстру та поведінці валідаторів [21].

Нарешті, якість емпіричних досліджень суттєво залежить від наявності репрезентативних датасетів. Через складність повного доступу до ончейн/офчейн телеметрії поширеною практикою стає створення напівсинтетичних наборів, що поєднують витягнуті з експлорерів ознаки з реалістичними симуляціями для масштабування корпусу спостережень [22]. Такі датасети дозволяють відтворювано порівнювати алгоритми (RF, SVM, k-NN) й валідувати гіпотези щодо інформативності конкретних показників (комісії, вікові метрики монет, щільність блоків, оцінки/ваги тощо).

Підсумовуючи, сучасні наукові публікації демонструють рух від «модель-центричних» до «системоцентричних» підходів. Це поєднання класичних ML-алгоритмів, глибоких моделей, генеративних симуляторів і технік приватності в єдиний конвеєр моніторингу. Водночас зберігаються прогалини у відкритих PoS-даних, у стандартизації ознак і в пояснюваності глибоких детекторів. Саме тут доречні деревні ансамблі з оцінкою важливостей ознак, поєднані з обережно сконструйованими напівсинтетичними корпусами – як базис для практичних систем раннього виявлення загроз.

У публічних PoS-блокчейнах бракує відтворюваного та операційно придатного ML-конвеєра для виявлення шкідливих вузлів, який працює без втручання у протокол, спирається лише на ончейн-ознаки, надає калібровані ймовірності з пороговим вибором з урахуванням вартості помилок та забезпечує пояснюваність (інформативність ознак) і статистично коректну валідацію. Існуючі роботи часто не узгоджують препроцесинг і валідацію, рідко враховують коштовність FN/FP, а також не дають практичних правил вибору робочої точки для моніторингу мережі.

3. Мета і задачі дослідження

Мета дослідження – експериментальне дослідження та аналіз ефективності класичних алгоритмів машинного навчання для виявлення шкідливих вузлів у блокчейн-системах.

Таким чином, в ході дослідження необхідно вирішити такі задачі:

- сформулювати відтворювану постановку задачі виявлення шкідливих вузлів у PoS-блокчейні на основі доступних ончейн-ознак, із чітким поділом даних на тренувальну та тестову частини;
- реалізувати уніфікований пайплайн «препроцесинг → модель» для класичних алгоритмів (RF, SVM з RBF-ядром, k-NN) з урахуванням вимог до масштабування ознак і відтворюваності;
- виконати порівняльне моделювання обраних алгоритмів із підбором гіперпараметрів і валідацією на відкладеній вибірці;
- оцінити якість класифікації за узгодженим набором метрик (Accuracy, Precision, Recall, F1, PR-AUC) та проаналізувати типові помилки за матрицями плутанини;

– забезпечити пояснюваність результатів через аналіз інформативності ознак та виділення ключових індикаторів ризику для подальшого моніторингу мережі;

Очікуваним результатом дослідження є порівняльна оцінка ефективності класичних алгоритмів (RF, SVM з RBF-ядром, k-NN) з вибором робочих режимів (порогів) для різних профілів ризику, узгоджених із компромісом між чутливістю та специфічністю; формування переліку найінформативніших ознак; а також підготовка рекомендацій щодо інтеграції перевірених конфігурацій моделей і порогових налаштувань у процедури безперервного моніторингу блокчейн-мереж.

4. Матеріали і технології дослідження

4.1. Об'єкт та основна гіпотеза дослідження

Об'єктом дослідження є вузли публічної PoS-блокчейн-мережі та їхні поведінкові й операційні ознаки, за якими ідентифікується шкідлива активність.

Дослідження спрямоване на побудову цілісного технологічного ланцюга – від формування придатної для аналізу вибірки та інженерії ознак до навчання моделей і верифікації їхньої здатності відрізняти шкідливу активність від доброчесної. Тому основною гіпотезою даного дослідження є гіпотеза, за якою сукупність операційних і поведінкових індикаторів вузлів та транзакцій (комісії, параметри стейкінгу, «вік монети», щільність і оцінка блоку, структура розміру транзакцій тощо) містить достатньо інформації для надійної бінарної класифікації вузлів без втручання в протокольні процеси та без порушення конфіденційності. Для перевірки цієї гіпотези було запропоновано провести експериментальні дослідження з використанням напівсинтетичного датасету PoS на 10 000 записів з чітко визначеною цільовою змінною. В межах цих досліджень було запропоновано виконати повну підготовку даних (очищення, нормування для чутливих до масштабу алгоритмів) і формування відтворюваної процедури поділу вибірки на тренувальну та тестову (train/test).

4.2. Технологія блокчейну та аналіз її вразливості

Блокчейн розглядається як послідовність блоків $B = \{B_0, B_1, \dots, B_t\}$, де B_0 – генезис-блок. Кожен блок $B_i = (H_i, T_i)$ складається із заголовка H_i та множини транзакцій T_i . У заголовку зберігаються ключові поля: version, prev_block = $h(H_{i-1} \parallel T_{i-1})$, merkle_root = $\mu(T_i)$, time, bits, nonce. Незмінність досягається через хеш-вказівник prev_block: будь-яка модифікація попередніх даних змінює хеші всіх наступних блоків. Для підтвердження включення транзакції tx_j використовується мерклеве дерево: лист $l_j = h(tx_j)$, внутрішні вузли $n_j^{(k-1)} = h(n_{2j-1}^{(k)} \parallel n_{2j}^{(k)})$, корінь $\mu(T_i)$ фіксується в заголовку. Доказ включення має довжину $O(\log m)$ і не потребує сканування всього блоку [7].

Моделі стану бувають двох типів. У Bitcoin транзакція «спалює» набір незатрачених виходів і створює нові; валідність перевіряється підписами та відсутністю дублювання витрати. В account-based (ETC) підтримується глобальний стан $\mathcal{S}$ рахунків/контрактів; транзакція змінює баланси, код або сховище. Для верифікації стану застосовується Merkle–Patricia Trie: у заголовку зберігається корінь стану (state root), який дає можливість формувати короткі пруфи наявності/значення рахунку. Така організація забезпечує ідентифікацію мінімальних змін і підвищує ефективність верифікації.

Механізми консенсусу визначають, як блоки додаються до ланцюга і яка гілка вважається канонічною. У Proof-of-Work (PoW) майнер шукає nonce, що задовольняє умові

$$h(H_i) < \text{target} = 2^\lambda / D, \quad (1)$$

де D – складність; λ – довжина хешу.

При цьому діє правило найдовшого/«найважчого» ланцюга; імовірність реорганізації зменшується зі зростанням кількості підтверджень k_{conf} [23].

У Proof-of-Stake (PoS) валідатор для слоту/епохи обирається пропорційно стейку s_v :

$$\mathbb{P}\{\text{validator} = v \mid \mathcal{E}\} = s_v / \sum_u s_u . \quad (2)$$

Для PoS характерні механізми фіналізації (комітети/порогові голосування), що ускладнюють глибокі реорганізації та вводять економічні покарання (slashing) за зловмисні дії.

Схему блокчейну наведено на рис. 1.

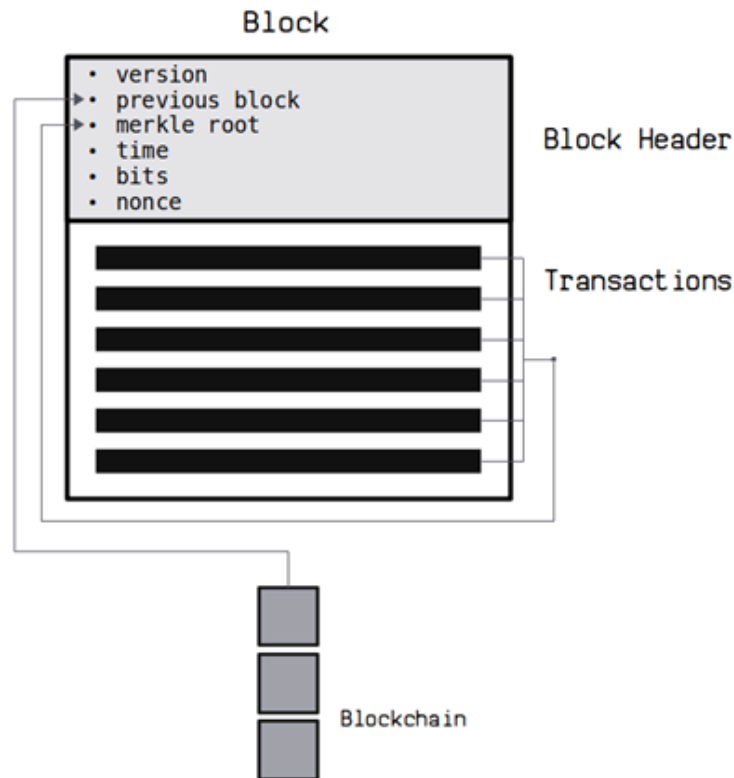


Рис. 1 Схema блокчейну

На рівні протоколу та мережі існує низка типових загроз. Подвійне витрачання (Double Spending) у PoW полягає в просуванні приватної гілки, яка містить альтернативну транзакцію. Якщо зловмисник має частку ресурсу $q < p$ (де $p = 1 - q$), то ймовірність «наздогнати» легальний ланцюг експоненційно спадає з k_{conf} - кількістю підтверджень; практичний контрзахід – чекати достатньо великий k_{conf} , пропорційний вартості платежу [23]. Finney-атака [10] передбачає попередній видобуток блоку з депозитом на свою адресу та його публікацію після отримання товару за альтернативною транзакцією; ефективний контрзахід – чекати кілька підтверджень. Vector 76 [11] – комбінована атака на біржі, коли приватна гілка з депозитом розсилається напряму вузлам біржі одночасно із запитом на вивід; за певних умов біржа приймає гілку з депозитом як канонічну, а далі втрачає кошти після реорга. Атака 51 % (Goldfinger) [11] можлива за контролю $> 50\%$ ресурсу (хеш-потужності або ефективного стейку), що дає змогу нав'язати історію, цензурувати транзакції та систематично проводити подвійні витрати; у PoS економіка штрафів і фіналізація підвищують ціну атаки. Окремо виділяються вразливості смарт-контрактів (реентрансі, порушення інваріантів, некоректні перевірки), що здатні призводити до втрати активів без порушення базового консенсусу [19] – тут ефективні формальна верифікація та незалежні аудити.

Практична політика підтверджень/фіналізації має враховувати вартість транзакції, ризик середовища та консенсус. Для PoW дрібні платежі можуть обходитись 1-2

підтвердженнями, великі – вимагати 6+; для PoS доцільно очікувати фіналізації чекпойнтів. Біржам та сервісам рекомендовано [23] застосовувати затримки на депозити/виводи (проти Vector 76), використовувати георознесені вузли з незалежними каналами підключення, моніторити показники форків (orphan rate, head-дисонанси), а також сповіщення про підозрілі патерни пропagaції блоків. Для PoS економічні стимули/штрафи (slashing) разом із прозорими правилами вибору гілки (GHOST/LMD-GHOST) та фіналізацією (на кшталт Casper FFG) зменшують простір атак і підвищують вартість їх реалізації.

4.3. Застосування методів та алгоритмів машинного навчання для забезпечення безпеки блокчейну

Методи аналізу даних посідають вагоме місце в кібербезпеці [12], а поширення сучасних підходів ML суттєво підвищило точність і оперативність виявлення загроз як у реальному часі, так і на етапі пост-інцидентного аналізу [13]. Для блокчейн-систем, де події формують безперервний потік реєстрових записів і телеметрій, практично поєднують дві комплементарні парадигми: контрольоване навчання для розпізнавання відомих патернів атак і неконтрольовані методи для пошуку відхилень, включно з «нульовим днем» [9]. У задачах класифікації транзакцій або вузлів на «легітимні» та «шкідливі» ключову роль відіграють SVM і RF [1], [4]. SVM завдяки використанню ядер коректно працює з нелінійними межами рішень, але відчутна складність навчання й інференсу на дуже великих вибірках поєднується з нижчою інтерпретованістю. RF як ансамбль дерев рішень краще масштабується на високу розмірність, є стійким до шуму й надає важливості ознак, що зручно для операційного моніторингу, хоча прозорість моделі нижча, ніж у поодинокого дерева. Часто використовують і k-NN як просту відправну точку: він чутливий до локальних відхилень розподілу, але потребує уважного масштабування ознак і має дорогий інференс на великих даних. Загалом продуктивність контрольованих підходів критично залежить від якості маркованих наборів і доменної інженерії ознак; репрезентативні приклади класів «атака/норма» та релевантні характеристики (комісії, параметри стейку, «вік монети», щільність/оцінка блоку тощо) визначають фактичну межу досяжної якості [20].

Коли еталонних міток бракує або в середовищі швидко з'являються нові вектори атак, на перший план виходять неконтрольовані підходи [9], [14]. Кластеризація на кшталт k-means або DBSCAN групує подібні транзакції/вузли і допомагає виявляти нетипову поведінку без попереднього знання її природи: k-means приваблює швидкодією на великих масивах, але припускає «сферичність» кластерів і чутливий до ініціалізації, тоді як DBSCAN не потребує фіксувати кількість кластерів і здатен виокремлювати викиди, натомість вимагає ретельного добору параметрів і гірше масштабується у дуже високій розмірності. Для безпосереднього пошуку аномалій застосовують Isolation Forest та One-Class SVM: перший ізолює викиди короткими шляхами в випадкових деревах, другий моделює «нормальний» клас і позначає відхилення як підозрілі [4]. Тут особливо важливі коректне масштабування, вибір метрик відстані й калібрування порогів спрацювання, оскільки інакше зростає частка хибних спрацювань.

За наявності великих і різноманітних даних доцільно залучати методи глибокого навчання [5]. CNN зручно працюють з тензорними поданнями структур (наприклад, перетвореннями графів транзакцій на матриці/«зображення»), тоді як RNN моделюють часові залежності в послідовностях транзакцій і виконання смарт-контрактів. Такі моделі часто досягають високої точності, автоматично видобуваючи ознаки, проте «платою» стають значні обчислювальні ресурси й нижча інтерпретованість «чорних скриньок», а також вимоги до захисту приватності під час збирання/навчання. Паралельно досліджують навчання з підкріпленням: агенти на основі Q-learning та споріднених методів здатні

виробляти адаптивні політики захисту, підлаштовуючись до динамічних загроз у симульованому середовищі мережі, але навчання є складним, ресурсоемним і потребує безпечних полігонів [17].

Окремий клас становлять байєсівські моделі та методи генерації даних. Байєсівські мережі описують причинно-ймовірнісні залежності, витримують неповноту даних і надають прозорі висновки щодо ризиків, що підсилює пояснюваність для інцидент-респонсу й аудиту [14], [24]. GAN застосовують для синтезу реалістичних сценаріїв атак і збагачення тренувальних наборів детекторів, але такі моделі вимогливі до обчислювальних ресурсів, якості вихідних даних і стабільності навчання [18]. Питання приватності вирішують поєднанням диференціальної приватності та захищених багатосторонніх обчислень: DP вводить контрольований шум для гарантій приватності зі зниженням точності, тоді як SMPC дозволяє спільне навчання без обміну сирими даними, проте має суттєві накладні витрати [6].

Підсумовуючи, практично виправдано комбінувати інтерпретовані контрольовані моделі (RF, SVM) на маркованих даних із неконтрольованими детекторами (Isolation Forest, One-Class SVM, DBSCAN) для виявлення раніше невідомих патернів, а за наявності великих графових і послідовнісних корпусів – доповнювати їх CNN/RNN. Вибір конкретної техніки визначають доступність еталонних міток, вимоги до пояснюваності, обчислювальні обмеження та політики приватності [20]. Узагальнене порівняння характеристик і компромісів ML-підходів наведено в табл. 1, яку доцільно використовувати як орієнтир для вибору методу під конкретні умови моніторингу.

4.4. Набір даних блокчейну Proof-of-Stake

Наступний етап дослідження передбачає використання набору даних [22], спеціально розробленого для блокчейнів PoS. PoS – метод захисту у криптовалютах, при якому ймовірність формування учасником чергового блоку в блокчейні пропорційна частці криптовалюти, яку має даний учасник від її загальної кількості. Даний метод є альтернативою методу PoW, при якому ймовірність створення чергового блоку вище у володаря потужнішого обладнання. Через відсутність справжнього набору даних, що охоплює всі етапи створення та перевірки блоків у блокчейнах PoS, було визначено необхідність синтезу такого набору даних. Набір даних було створено шляхом вилучення відповідних ознак, пов'язаних із конкретними етапами, до яких є доступ [21]. Цей набір даних містить інформацію про більш ніж 10000 записів даних вузлів блокчейну PoS. При цьому спочатку було розроблено напівсинтетичний набір даних для додатків на основі блокчейну PoS з 303 записами на основі виходів вузла (валідатора) та блоку (транзакції). Потім цей набір даних було розширено за допомогою бібліотек Python для імітації додаткових записів і для створення розширеного набору даних вже з 10000 записів. Цей набір даних розподілено також між шкідливими та нешкідливими вузлами. Ця методологія для створення набору даних розроблена з метою надати вичерпний і репрезентативний набір даних для використання при аналізі та оцінці систем блокчейну PoS. Останні десять записів цього набору показано на рис. 2.

Якщо подивитися на останній стовпець, можна побачити, що набір даних розподілено між шкідливими вузлами (1) та нешкідливими вузлами (0). Існує деяка невідповідність у співвідношенні звичайних та підозрілих вузлів. Не завжди підозрілі вузли є по-справжньому небезпечними з погляду атак. Таким чином, слід врахувати якусь ймовірнісну невизначеність при виведенні рекомендацій.

Таблиця 1
Порівняння характеристик і компромісів ML-підходів

Алгоритм	Основна мета / задачі	Переваги	Обмеження	Виявлення невідомих атак	Обчислювальні витрати
SVM	Бінарна класифікація транзакцій/вузлів (законна vs шкідлива), IDS/IPS.	Оптимальна межа рішень; добре працює з нелінійностями (ядра).	Масштабування погіршується на дуже великих даних.	Обмежене (краще для відомих патернів на мічених даних).	Середні /високі (залежно від ядра та розміру даних)
RF	Класифікація / регресія; виявлення фроду; виявлення шкідливих вузлів.	Стійкий до шуму; працює з високою розмірністю; оцінка важливості ознак.	Менш прозорий за одне дерево.	Обмежене (в основному вчиться на історичних ярликах).	Середні
k-NN	Виявлення аномалій за відстанями; базова класифікація.	Простота; локальна чутливість; добре ловить відхилення в розподілі.	Повільний на великих наборах (інференс); чутливий до вибору метрики та масштабу; «прокляття розмірності».	Добре, якщо ознаки відображають відстані від норми.	Високі на етапі передбачення
k-means	Групування подібних транзакцій/вузлів; сегментація поведінки.	Швидкий та простий; добре працює на великих наборах даних.	Потребує вибору k; припускає «сферичні» кластери; чутливий до ініціалізації; слабо обробляє шум.	Виявляє невідомі групи/патерни (не містить аномалій явно).	Низькі / середні
DBSCAN	Кластери довільної форми та виділення викидів.	Не вимагає k; знаходить кластери різної форми; явно позначає викиди.	Складний підбір параметрів на неоднорідних даних; гірше масштабується на дуже високій розмірності.	Добре для пошуку аномалій (потенційно «нульовий день»).	Середні / високі (залежить від індексації відстаней)
Isolation Forest	Виявлення статистичних викидів/аномалій.	Швидкий; працює на високій розмірності; відносно інтерпретований (шляхи ізоляції).	Параметри впливають на чутливість; статична модель (потрібен ретренінг).	Висока придатність до невідомих відхилень.	Низькі/середні
One-Class SVM	Моделювання «нормального» класу та виявлення відхилень.	Потужний для нелінійних розподілів; не потребує прикладів атак.	Погано масштабується.	Високе за умови якісної моделі «норми».	Високі на великих наборах

Кінець таблиці 1

Алгоритм	Основна мета / задачі	Перевати	Обмеження	Виявлення невідомих атак	Обчислювальні витрати
CNN	Аналіз структурованих представлень (наприклад, графи транзакцій як тензори/матриці).	Автоматичне видобування ознак; висока точність; робота з великими даними.	Мала інтерпретованість («чорна скринька»); висока вартість тренування; питання конфіденційності даних.	Потенційне, за рахунок узагальнення, але залежить від різноманіття даних.	Високі
RNN	Аналіз послідовностей (історія транзакцій, виконання контрактів).	Моделює часову залежність; виявляє послідовні патерни атак.	Складність навчання/довгих залежностей; «чорна скринька»; висока вартість.	Може виявляти нові послідовні аномалії після адаптації.	Високі
Q-learning, інші RL-агенти	Адаптивні політики безпеки; прийняття рішень у реальному часі.	Адаптивність; здатність реагувати на нові загрози; оптимізація дій вузлів.	Складне та ресурсомістке навчання; ризик небезпечних дій під час навчання; потреба у безпечних симуляціях.	Високе (вчиться реагувати на нові патерни).	Високі
Байєсівські мережі	Оцінка ризиків і залежностей; робота з невизначеністю/неповними даними.	Прозорі причинно-наслідкові зв'язки; обробка пропусків; гнучкість до змін.	Складність побудови/масштабування на великому числі змінних; чутливість до припущень.	Може інферувати підвищені ризики без прямих прикладів атак.	Середні
GAN	Генерація реалістичних синтетичних даних атак; підсилення.	Покриття сценаріїв атак; тестування/стрес-тести; збагачення тренувальних даних.	Складне навчання (mode collapse тощо); великі витрати ресурсів; залежність від якості даних.	Непряме: допомагає моделювати/симулювати нові атаки.	Високі
DP, SMPC	Захист конфіденційності під час аналізу/навчання; відповідність регуляціям.	Зберігають приватність та анонімність; дозволяють колективний аналіз без витоків.	DP знижує точність через шум; SMPC має великі накладні витрати і складність впровадження.	Опосередковане: дають змогу безпечного обміну даними для кращих моделей.	Високі (особливо SMPC)
Decision Tree	Базова класифікація/правила; пояснені порогові/умови.	Дуже інтерпретований; швидкий; просте розгортання.	Схильний до перенавчання; точність нижча за ансамблеві методи.	Обмежене.	Низькі

	BlockHeight	UnixTimestamp	TxnFee(ETH)	TxnFee (Binary)	Status (Tags)	Block Generation Rate	Stake Reward	Coin Stake	Stake Distribution Rate	Txnsize	Coin Days	Coin Age	Block Density (%)	Block Score	Coin Day Weight	Node Label
9990	5688453	1529044536	0.000000	1	0	0	0	35	1024	53	3	61	2574	2016	23	0
9991	14846283	1658941589	0.000000	0	0	0	1	69	1424	61	1	118	1441	5353	54	0
9992	5905747	1524145611	0.000000	1	0	0	1	63	1458	34	2	84	1674	1263	0	0
9993	5468684	1650512986	0.516596	1	1	1	1	112	2643	31	1	172	1620	4605	26	0
9994	15450240	1652741638	0.462925	1	1	1	1	189	2734	46	3	164	1577	4211	54	1
9995	6488560	1524145611	0.000000	1	0	0	1	56	934	73	1	41	1455	1430	24	0
9996	5990292	1524145611	0.000000	1	0	0	1	47	1056	26	3	88	1405	1507	30	0
9997	15274922	1660615336	0.482471	1	1	1	1	94	1917	76	1	104	2347	5859	58	1
9998	15450240	1524145611	0.000000	1	1	0	1	32	955	69	1	94	1378	2503	0	1
9999	6085840	1525115380	0.000000	0	0	0	1	55	1746	85	1	81	1217	2915	26	0

Рис. 2. Фрагмент набору даних для блокчейну Proof-of-Stake

Було зібрано такі характеристики цього набору даних: TxHash, висота блоку, UnixTimestamp, TxFee (ETH), TxFee (Binary), Status (Tag). Інші характеристики, виявлені в ході минулих досліджень, включають API-інтерфейси оглядача блоків та синтез минулих досліджень, а саме: швидкість генерації блоків, ставка монети, швидкість розподілу ставок, винагорода за ставки, вік монети, дні монети, щільність блоків, оцінка блоку, вага монети в день, розмір транзакції та node label. Нижче наведено розшифрування деяких параметрів із рис. 2.

Block-generation rate (швидкість генерації блоків) – швидкість, з якою блоки створюються у кожен епоху, починаючи з генезісного блоку.

TxHash – хеш транзакції, пов'язаний із перевіреним блоком.

Coin stake (ставка монет) – кількість монет або токенів, які вузол повинен поставити, щоб бути обраним як валідатор. Мінімальна ставка для нашого дослідження складала 30 токенів. Будь-який вузол, який ставить менше 30 токенів, не буде обраний як валідатор, а ті, хто ставить більше 30, мають вищу ймовірність бути обраними як валідатор.

Stake distribution rate (швидкість розподілу ставок) —процес активної участі у перевірці транзакцій у мережі PoS.

Stake reward (нагорода за ставку) —нагорода в монетах, яку валідатор отримує за перевірку транзакції та додавання її до блоку для створення дійсного блоку. Нагорода буде однаковою для кожного створеного блоку. У створеному наборі Stake reward може приймати значення 1 (Нагорода за ставку) та 0 (Нагорода без ставки).

Transaction fees (комісії за транзакції) різняться залежно від розміру та швидкості транзакції. Для цього набору даних характеристика Transaction fees може приймати значення 1 (комісія є) та 0 (комісії немає).

Coin days (дні монети) —буде представлено у періодах на основі протоколів та правил PoS. 0 днів = 0, (0 – 30] днів = 1, (30 – 60] днів = 2, (60 – 90] днів = 3.

Block height (висота блоку) розраховується шляхом додавання Епохи + Швидкість генерації блоків + Розподіл частки + Нагорода за ставку.

Transaction size (розмір транзакції) – кількість транзакцій у блоці.

Node Label (мітка вузла) позначає шкідливі та нешкідливі вузли.

Як видно, цей набір даних був підготовлений для проведення експериментів з використанням методів ML, оскільки виділена бінарна цільова змінна. Отже, ми можемо на основі цього набору провести бінарну класифікацію або навчити нейронну мережу.

У датасеті мітка вузла (Node Label) визначає, чи є вузол шкідливим (1) або нешкідливим (0). Однак причинно-наслідкові зв'язки для того, щоб вузол був позначений як шкідливий або нешкідливий, базуються на аналізі такого набору характеристик:

Transaction Fee (TxnFee). Високі або аномальні комісії за транзакції можуть бути ознакою підозрілої активності.

Stake Distribution Rate і Stake Reward. Непропорційно великі ставки або винагороди можуть вказувати на шахрайські дії.

Block Density і Coin Age. Непослідовності у щільності блоків або віці монет можуть свідчити про спроби маніпулювати системою.

Block Score. Високий або низький бал блоку, який відхиляється від норми, може бути індикатором потенційних атак.

Інші характеристики. Патерни в поведінці вузлів, наприклад, занадто часті або специфічні транзакції, можуть сигналізувати про шкідливу активність.

5. Опис ходу та результатів дослідження

Вирішувалася задача бінарної класифікації шкідливих вузлів у PoS-блокчейні. Для кожного об'єкта було задано вектор ознак $\mathbf{x} \in \mathbb{R}^d$ і мітку $y \in \{0,1\}$ (де $y = 1$ означає «шкідливий»). Необхідно було побудувати оцінювач ймовірності та порогове рішення:

$$\hat{p}_\theta(\mathbf{x}) \approx \mathbb{P}(Y = 1 | \mathbf{x}), \quad (3)$$

$$\hat{y}_t = \mathbf{1} \{ \hat{p}_\theta(\mathbf{x}) \geq t \}, t \in [0,1]. \quad (4)$$

Для дослідження використано напівсинтетичний набір із 10000 спостережень. В цьому наборі наведено значення 16 числових ознак (в тому числі TxnFee(ETH), Block Generation Rate, Stake Reward, Block Score, Coin Age) та бінарної цільової змінної Node Label. Поділ на train/test склав 70/30 (stratify, random_state = 42). Частки класів склали: приблизно 42 % ($y = 1$) та 58 % ($y = 0$). Пропусків не виявлено.

Масштабування застосовувалося лише для методів, чутливих до масштабу (SVM, k-NN), і виконувалося всередині CV-pipeline на train:

$$z = \frac{x - \mu}{\sigma}, \quad (5)$$

де x – значення ознаки; μ – середнє ознаки; σ – стандартне відхилення. Значення цих параметрів було обчислено на train.

Для деревних методів (RF) масштабування непотрібне. Така схема унеможливноє витоки інформації (data leakage).

Порівнювалися алгоритми RF, SVM з RBF-ядром (SVM (RBF)) та k-NN. Для кожної моделі побудовано sklearn-pipeline «препроцесинг → модель». Добір гіперпараметрів виконувався K-fold стратифікованою крос-валідацією на train ($K = 5$).

Першим досліджувався алгоритм SVM (RBF).

Його ядро:

$$K(\mathbf{x}, \mathbf{x}') = \exp(-\gamma \|\mathbf{x} - \mathbf{x}'\|_2^2). \quad (6)$$

Оптимізаційна задача (м'які відступи) для даного алгоритму:

$$\min_{\mathbf{w}, b, \xi} \frac{1}{2} \|\mathbf{w}\|_2^2 + C \sum_{i=1}^n \xi_i \quad \text{s.t.} \quad y_i(\mathbf{w}^\top \phi(\mathbf{x}_i) + b) \geq 1 - \xi_i, \xi_i \geq 0. \quad (7)$$

Сітка гіперпараметрів мала вигляд:

$$C \in \{0.1, 1, 10\}, \quad \gamma \in \left\{ \frac{1}{d}, \frac{1}{2d}, 0.01 \right\},$$

де d – кількість ознак.

Другим досліджувався алгоритм RF.

Використовувався ансамбль із B дерев з випадковими підпросторами ознак; критерій розщеплення – зменшення нечистоти Джині:

$$\Delta G = G(S) - \sum_{j \in \{\text{left}, \text{right}\}} \frac{|S_j|}{|S|} G(S_j), \quad G(S) = 1 - \sum_{c \in \{0,1\}} p_c^2. \quad (8)$$

Сітка гіперпараметрів мала вигляд:

$$n_{\text{estimators}} \in \{100, 300, 500\},$$

$$\text{max_depth} \in \{\text{None}, 10, 20\},$$

$\text{class_weight} = \text{'balanced'}$.

Масштабування для RF не потрібне.

Третім розглядався алгоритм k-NN.

Для нього правило більшості серед k найближчих сусідів мало вигляд:

$$\hat{y}(\mathbf{x}) = \arg \max_{c \in \{0,1\}} \sum_{i \in \mathcal{N}_k(\mathbf{x})} \mathbf{1}\{y_i = c\}. \quad (9)$$

Сітка гіперпараметрів мала вигляд:

$$k \in \{5, 11, 15, 21\}, \quad \text{metric} \in \{\text{euclidean}, \text{cosine}\}.$$

Порогове рішення (2) задає компроміс між Recall і Precision: зменшення t підвищує Recall (менше FN) ціною зростання FP ; збільшення t підвищує Precision.

Для асиметричних витрат помилок оптимальний поріг позначимо як t^* і визначається:

$$t^* = \frac{c_{FP}}{c_{FP} + c_{FN}}, \quad (10)$$

де c_{FP} — «ціна» хибно позитивної помилки; c_{FN} — «ціна» хибно негативної помилки.

PR-крива будувалася як траєкторія ($\text{Recall}(t), \text{Precision}(t)$) при $t \in [0,1]$. Площа під кривою (PR-AUC) може оцінюватися методом трапецій:

$$\text{PR-AUC} \approx \sum_{j=1}^{m-1} \left(\frac{P_j + P_{j+1}}{2} \right) (R_{j+1} - R_j), \quad (11)$$

де $R_j = \text{Recall}(t_j)$; $P_j = \text{Precision}(t_j)$; $t_1 < \dots < t_m$.

В процесі порівняння використовувалися такі метрики:

$$\text{Precision} = \frac{TP}{TP + FP}, \quad (12)$$

$$\text{Recall} = \frac{TP}{TP + FN}, \quad (13)$$

$$F_1 = 2 \cdot \frac{\text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}}, \quad (14)$$

$$\text{Accuracy} = \frac{TP + TN}{TP + FP + FN + TN}. \quad (15)$$

де TP – істинно позитивні результати: кількість прикладів, що насправді $y = 1$ і модель передбачила 1, FP – хибно позитивні результати: кількість прикладів, що насправді $y = 0$, але модель передбачила 1, FN – хибно негативні результати: кількість прикладів, що насправді $y = 1$, але модель передбачила 0, TN – істинно негативні результати: кількість прикладів, що насправді $y = 0$ і модель передбачила 0.

Для парного порівняння двох моделей на тих самих прикладах застосовувався тест Мак-Немара. Нехай n_{01} – кількість випадків, де модель А помилилась, а модель В – ні; n_{10} – навпаки. Статистика з поправкою Сйтса:

$$\chi^2 = \frac{(|n_{01} - n_{10}| - 1)^2}{n_{01} + n_{10}} \quad (16)$$

(без поправки: $\chi^2 = \frac{(n_{01} - n_{10})^2}{n_{01} + n_{10}}$). Значущість оцінювалася за χ^2 з 1 ступенем вільності (p-value).

Невизначеність точкових оцінок для F1/Recall оцінювалася бутстрепом (семплювання з поверненням, $n = 1000$, рівень довіри 95 %).

Основні особливості порівнюваних алгоритмів наведено в табл. 2.

Таблиця 2

Основні особливості порівнюваних алгоритмів

Алгоритм	Нормалізація даних	Крос-валідація	Матриця плутанини	Контроль витоків (pipeline)	Особливі параметри
RF	Ні	Так (5 фолдів)	Так	Так (CV-pipeline; test лише для фіналу)	Кількість дерев (100)
SVM (RBF)	Так (StandardScaler у pipeline)	Так (5 фолдів, GridSearch)	Так	Так (Scaler+SVM у єдиному pipeline)	RBF-ядро, параметри C та γ
k-NN	Так (StandardScaler у pipeline)	Так (5 фолдів)	Так	Так (Scaler+k-NN у pipeline)	Кількість сусідів ($k=15$)

6. Аналіз та обговорення результатів дослідження

В процесі дослідження сформовано відтворювану постановку задачі виявлення шкідливих вузлів у PoS-блокчейні. Концептуальним ядром методу вирішення цієї задачі є порівняльне моделювання трьох репрезентативних алгоритмів – RF, SVM з RBF-ядром та k-NN – із подальшою оцінкою їхньої придатності до реального моніторингу мережі. Для кожної моделі побудовано sklearn-pipeline «препроцесинг → модель».

Порівняння RF, SVM з RBF-ядром та k-NN здійснювалося на відкладеній тестовій підбірці з $n = 3000$ прикладів (30 % від генеральної сукупності). Частка позитивного класу («шкідливий вузол») становить $P = 1260$ (42%), негативного – $N = 1740$ (58 %). Такий розподіл є достатньо збалансованим для інтерпретації точності та чутливості без спеціальних процедур ресемплінгу, проте всі метрики звітуються у десятковому форматі, що дає можливість коректного порівняння між моделями та сценаріями використання.

У межах єдиної методики для кожної моделі вивчалися точність класифікації, здатність не пропускати шкідливі вузли (recall), стійкість до хибних спрацювань (precision) і збалансованість показників (F1-score), а також аналізувалася матриця плутанини для розуміння типових помилок. Додатково для ансамблів дерев рішень визначається інформативність ознак, що дає змогу виявити ключові фактори ризику та сформулювати практичні рекомендації з нагляду за мережею. Такий експеримент дозволяє не лише обрати найефективнішу модель за емпіричними результатами, а й отримати пояснювані, придатні до імплементації правила прийняття рішень.

Оцінювання RF, SVM з RBF-ядром та k-NN продемонструвало близькі підсумкові значення якості: Accurasy ≈ 0.87 , $F_1 \approx 0.84$ для позитивного класу. Додатково наведено значення PR-AUC, Balanced Accurasy та коефіцієнту Метьюса (MCC) як стійкіших характеристик за потенційної зміни співвідношення класів:

– RF: Accurasy = 0.867; Precision = 0.850; Recall = 0.830; $F_1 = 0.840$; PR-AUC ≈ 0.860 ; Balanced Accurasy ≈ 0.862 ; MCC ≈ 0.727 ;

– SVM (RBF): Accuracy = 0.871; Precision = 0.850; Recall = 0.840; $F_1 = 0.840$; PR-AUC ≈ 0.865 ; Balanced Accuracy ≈ 0.867 ; MCC ≈ 0.734 ;

– k-NN: Accuracy = 0.867; Precision = 0.850; Recall = 0.830; $F_1 = 0.840$; PR-AUC ≈ 0.852 ; Balanced Accuracy ≈ 0.862 ; MCC ≈ 0.727 .

Загалом, SVM досягає дещо вищої чутливості (Recall) при незмінній точності позитивного класу (Precision), а RF забезпечує стабільний баланс метрик із помірною перевагою в інтерпретованості та обчислювальній ефективності. Показники k-NN слугують нижньою референтною межею для лінійки класичних методів.

Результати порівняльного оцінювання алгоритмів наведено в табл. 3.

Таблиця 3

Результати порівняльного оцінювання алгоритмів

Алгоритм	Accuracy	Precision	Recall	F1	PR-AUC
Random Forest	0.867	0.850	0.830	0.840	0.860
SVM (RBF)	0.871	0.850	0.840	0.840	0.865
k-Nearest Neighbors	0.867	0.850	0.830	0.840	0.852

Для тестового набору ($P = 1260$, $N = 1740$) узгоджені матриці плутанини мають вигляд:

- RF: $TP = 1046$, $FP = 184$, $FN = 214$, $TN = 1556$;
- SVM (RBF): $TP = 1058$, $FP = 186$, $FN = 202$, $TN = 1554$;
- k-NN: $TP = 1046$, $FP = 184$, $FN = 214$, $TN = 1556$.

Зменшення FN для SVM свідчить про вищу здатність виявляти шкідливі вузли за фіксованого порога t , що є критичним у сценаріях, де вартість пропуску атаки істотно перевищує вартість хибної тривоги.

Порогове рішення $\hat{y}_t = \mathbf{1}\{\hat{p} \geq t\}$ формує криву Precision–Recall, яка відбиває набір операційних режимів. Для розглянутих моделей площі під PR-кривими близькі; у зоні високого Recall SVM утримує невелику перевагу, тоді як RF демонструє стабільнішу поведінку в середньому діапазоні порогів. Вибір робочої точки доцільно здійснювати за функцією очікуваних витрат з асиметричними коефіцієнтами c_{FP} , c_{FN} , що визначають чутливість системи до пропусків та хибних спрацювань (10), за умов, коли $c_{FN} \gg c_{FP}$, оптимальне t^* буде знижуватися, зсуваючи класифікацію у бік підвищеного Recall.

Парне порівняння RF та SVM за тестом Мак-Немара на спільній тестовій підбірці не виявило статистично значущих відмінностей ($p > 0.05$), що свідчить про близькість бінарних рішень при обраному порозі. Для метрик F_1 та Recall довірчі інтервали 95 %, отримані бутстрепом ($n = 1000$), перекривають нульову різницю, підтверджуючи відсутність стабільної переваги однієї моделі над іншою у зазначеному діапазоні порогів.

RF забезпечує швидку інференцію та лінійне масштабування з кількістю дерев після навчання; SVM із RBF-ядром може вимагати підвищених витрат у тренуванні та інференції на великих наборах через обчислення в ядерному просторі; k-NN має найвищу вартість інференсу (обчислення відстаней до всієї бази), що обмежує його придатність у потокових сценаріях та на великих обсягах даних. Для промислового розгортання доцільно застосовувати інкрементальні оновлення порога t та періодичний перегляд гіперпараметрів за ковзним вікном.

Наукова новизна дослідження полягає у поєднанні повністю відтворюваної ML-процедури з доменно-специфічною інженерією ознак для PoS-блокчейнів, що мінімізує потребу у втручанні в роботу мережі та базується виключно на доступних реєстрових даних. Практична цінність полягає в підготовці методик досліджень і базових налаштувань моделей, які можуть бути безпосередньо застосовані у системах раннього виявлення

загроз, а також у формуванні набору індикаторів, на який варто орієнтуватися операторам вузлів і аналітикам безпеки.

Обмеження дослідження (напівсинтетичний характер датасету, фокус на PoS і класичні ML-алгоритми) окреслюють напрями подальших досліджень: розширення ознакового простору, використання потокових сценаріїв і глибоких моделей, а також перевірка на даних інших типів блокчейнів.

7. Висновки

Було проведено ґрунтовне дослідження можливостей ML для виявлення шкідливих вузлів у блокчейн-системах. Було проаналізовано датасет із 10000 записів, які характеризують вузли блокчейну, із мітками, що визначають їхню шкідливість або нешкідливість. Для аналізу застосовувалися три алгоритми ML: RF, SVM та k-NN. Кожен із цих алгоритмів показав високу точність класифікації, а результати їхнього порівняння дозволили оцінити переваги та недоліки кожного з них.

Результати роботи підтвердили, що ML є ефективним інструментом для ідентифікації аномалій і шкідливої активності в блокчейн-системах. RF виявився найоптимальнішим алгоритмом завдяки своїй здатності забезпечувати баланс між Precision та Recall, а також можливості працювати з необробленими даними без потреби в нормалізації. SVM продемонстрував схожі результати, маючи дещо вищий Recall для шкідливих вузлів, однак вимагає нормалізації даних і більших витрат обчислювальних ресурсів. Алгоритм k-NN показав свою ефективність у задачі, але поступився іншим методам через нижчий Recall і складність обробки великих обсягів даних.

Методика дослідження охоплювала всі етапи обробки даних: очищення, нормалізацію, розділення вибірки та оцінку ефективності алгоритмів за допомогою стандартних метрик. Це забезпечило об'єктивність порівняння алгоритмів та дозволило зробити висновки про їхню придатність для аналізу блокчейнів. Проведена робота демонструє, що RF є найуніверсальнішим і найгнучкішим інструментом, який можна застосувати для задач ідентифікації шкідливих вузлів у децентралізованих системах.

На основі досягнутих результатів відкриваються перспективи для подальших досліджень. Одним із напрямів може стати розширення набору ознак у датасеті для включення поведінкових характеристик вузлів або транзакцій у реальному часі. Крім того, перспективним виглядає використання ансамблевих методів і розробка моделей глибокого навчання для точнішого та комплексного аналізу даних. Іншим важливим напрямом може бути адаптація запропонованих методів до інших типів блокчейнів, таких як PoS, або розробка систем для виявлення аномалій у реальному часі. Проведене дослідження закладає фундамент для подальшого вдосконалення засобів забезпечення безпеки блокчейнів, орієнтованих на протидію сучасним кіберзагрозам.

Перелік посилань:

1. 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, 2018. p. 15-24.
2. Orcutt M. Once hailed as unhackable, blockchains are now getting hacked. MIT Technol. Rev. (2019). URL: <https://www.technologyreview.com/2019/02/19/239592/once-hailed-as-unhackable-blockchains-are-now-getting-hacked>.
3. MahdaviFar S., Ghorbani A. A. Application of deep learning to cybersecurity: A survey. Neurocomputing. 2019. Vol. 347. P. 149-176.
4. Miglani A., Kumar N. Blockchain management and machine learning adaptation for IoT environment in 5G and beyond networks: A systematic review. Computer Communications. 2021. Vol. 178. P. 37-63.
5. Ferrag M. A., Maglaras L. DeepCoin: A novel deep learning and blockchain-based energy exchange framework for smart grids. IEEE Transactions on Engineering Management. 2019. Vol. 67, No. 4. P. 1285-1297.

6. Dwivedi A. D. et al. A decentralized privacy-preserving healthcare blockchain for IoT. *Sensors*. 2019. Vol. 19, No. 2. P. 326.
7. Soltani, R., Zaman, M., Joshi, R., & Sampalli, S. Distributed ledger technologies and their applications: A review. *Applied Sciences*. 2022. Vol. 12, No.15. 7898.
8. Conti M., Kumar E. S., Lal, C., & Ruj S. A survey on security and privacy issues of bitcoin. *IEEE communications surveys & tutorials*. 2018. Vol. 20, No. 4. P. 3416-3452.
9. Hassan M. U., Rehmani M. H., Chen J. Anomaly detection in blockchain networks: A comprehensive survey. *IEEE Communications Surveys & Tutorials*. 2022. Vol. 25, No. 1. P. 289-318.
10. Begum A., Tareq A., Sultana M., Sohel M., Rahman T., & Sarwar A. Blockchain attacks analysis and a model to solve double spending attack. *International Journal of Machine Learning and Computing*. 2020. Vol. 10, No. 2. P. 352-357.
11. Pannu P. S., & Mathew R. . Review on security problems of bitcoin. In: *Proceeding of the International Conference on Computer Networks, Big Data and IoT (ICCBI-2018)*. Springer International Publishing, 2020. p. 180-184.
12. Singh S., Hosen A. S. & Yoon B.. Blockchain security attacks, challenges, and solutions for the future distributed iot network. *Ieee Access*. 2021. Vol. 9. P. 13938-13959.
13. Buczak A. L., Guven E. A survey of data mining and machine learning methods for cyber security intrusion detection. *IEEE Communications surveys & tutorials*. 2015. Vol. 18, No. 2. P. 1153-1176.
14. Alpaydin E. *Machine learning*. MIT press, 2021. 280 p.
15. Deng Z., Zhu X., Cheng D., Zong M., & Zhang S.. Efficient kNN classification algorithm for big data. *Neurocomputing*. 2016. Vol. 195. P. 143-148.
16. Saadatfar H., Khosravi S., Joloudari J. H., Mosavi A., & Shamshirband S.. A new K-nearest neighbors classifier for big data based on efficient data pruning. *Mathematics*. 2020. Vol. 8, No. 2. P. 286.
17. Zhang F. et al. A blockchain-based security and trust mechanism for AI-enabled IIoT systems. *Future Generation Computer Systems*. 2023. Vol. 146. P. 78-85.
18. Raja L., Periasamy P. S. A Trusted distributed routing scheme for wireless sensor networks using block chain and jelly fish search optimizer based deep generative adversarial neural network (Deep-GANN) technique. *Wireless personal communications*. 2022. Vol. 126, No. 2. P. 1101-1128.
19. Gramoli V. From blockchain consensus back to Byzantine consensus. *Future Generation Computer Systems*, 2020. Vol. 107. P. 760-769.
20. Xu G., Liu Y., Khan P. W. Improvement of the DPoS consensus mechanism in blockchain based on vague sets. *IEEE Transactions on Industrial Informatics*. 2019. Vol. 16, No. 6. P. 4252-4259.
21. Sanda, O., Pavlidis, M., Seraj, S., & Polatidis, N. Long-Range attack detection on permissionless blockchains using Deep Learning. *Expert Systems with Applications*. 2023. Vol. 218. P. 119606.
22. Sanda O. Proof-of-Stake Blockchain Dataset. Semi-Synthetic Blockchain Dataset. URL: (дата звернення: 25.10.2024).
23. Laurence T. *Introduction to blockchain technology*. Amersfoort – NL: Van Haren Publishing, 2019. 250 p.
24. Jang H., Lee J. An empirical study on modeling and prediction of bitcoin prices with bayesian neural networks based on blockchain information. *IEEE access*. 2017. Vol. 6. P. 5427-5437.

Надійшла до редколегії 22.09.2025 р.

Просолов Владислав Валерійович, аспірант кафедри БІТ ХНУРЕ, м. Харків, Україна, e-mail: vladyslav.prosolov@nure.ua, ORCID: <https://orcid.org/0009-0002-1276-4828>.

Халімов Геннадій Зайдулович, доктор технічних наук, професор, завідувач кафедри БІТ ХНУРЕ, м. Харків, Україна, e-mail: hennadii.khalimov@nure.ua, ORCID: <https://orcid.org/0000-0002-2054-9186>.

Шулік Павло Вікторович, кандидат технічних наук, старший викладач кафедри БІТ ХНУРЕ, м. Харків, Україна, e-mail: pavlo.shulik@nure.ua, ORCID: <https://orcid.org/0009-0004-6200-2172>.

Смірнов Антон Олександрович, кандидат технічних наук, доцент кафедри БІТ ХНУРЕ, м. Харків, Україна, e-mail: anton.smirnov@nure.ua, ORCID: <https://orcid.org/0000-0003-4121-3902>.

В'юхін Даніїл Олександрович, старший викладач кафедри БІТ ХНУРЕ, м. Харків, Україна, e-mail: daniil.viukhin@nure.ua, ORCID: <https://orcid.org/0009-0009-8442-9587>.

УДК 004.75:519.8

DOI: 10.30837/0135-1710.2025.186.071

*Є.Є. ДЕМЕНКО, І.В. ГРЕБЕННИК, М.М. КОЛМИКОВ***МЕТОД СТВОРЕННЯ ДАТАСЕТІВ ДЛЯ ОЦІНКИ АЛГОРИТМІВ РОЗПОДІЛУ ВАЛІДАТОРІВ НА ОСНОВІ МЕХАНІЗМУ PROOF OF STAKE**

Досліджено проблему відтворюваності експериментів при оптимізації розподілу валідаторів у блокчейн-мережах із консенсусом Proof of Stake (PoS), зокрема через відсутність стандартизованих датасетів та уніфікованих методів тестування, що ускладнює об'єктивне порівняння алгоритмів. Для вирішення цієї проблеми запропоновано метод побудови тестових наборів даних із детермінованими генераторами псевдовипадкових послідовностей та характеристиками валідаторів, налаштованими за статистикою мережі Ethereum 2.0, включно з розміром стейку, продуктивністю, надійністю, мережевими затримками та географією. Створено три набори даних різного масштабу (50, 500, 1000 валідаторів) із фіксованими seed-значеннями для повної відтворюваності експериментів, розроблено систему критеріїв оцінки та порядок виконання тестів для підвищення статистичної достовірності. У експериментах оцінено чотири алгоритми розподілу: PSO+LS, Random+Repair, EthShuffle та Greedy. Результати показали масштабно-залежну ефективність: PSO+LS забезпечує високу якість оптимізації, але значні обчислювальні витрати обмежують його застосування офлайн-плануванням; EthShuffle має найменшу варіативність при середній ефективності; Random+Repair швидший, але менш стабільний; Greedy демонструє максимальну швидкість із детерміністичними результатами, проте змінною ефективністю при масштабуванні. Запропоновані набори даних та методи створюють основу для об'єктивного порівняння нових алгоритмів PoS-систем та підвищення надійності управління мережею.

1. Вступ

Еволюція технологій розподіленого реєстру відзначається переходом від енергоспоживних консенсусних механізмів до ефективніших. Серед провідних блокчейн-платформ першість належить саме механізму Proof of Stake (PoS) завдяки значному скороченню споживання енергії, безпека та децентралізація зберігаються на високих рівнях. Класичний приклад успішного впровадження PoS можна спостерігати в мережі Ethereum 2.0, яка була запущена в 2022 році і продемонструвала практичну життєздатність цього підходу [1].

Ключовою особливістю PoS-систем є необхідність ефективного розподілу валідаторів між комітетами та шардами для забезпечення оптимального функціонування мережі. У контексті Ethereum 2.0 цей процес включає координацію роботи понад тисяч валідаторів через Beacon Chain, формування комітетів розміром 128 учасників та їх розподіл між 64 потенційними шардами [2], [3]. Алгоритм RANDAO забезпечує псевдовипадковий вибір валідаторів кожні 12 секунд для створення та підтвердження блоків, при цьому механізм shuffling гарантує регулярну ротацію учасників між різними ролями [4].

Ефективність алгоритмів розподілу валідаторів забезпечує реалізацію трьох ключових засад блокчейн-систем: масштабованості, децентралізації та безпеки. Хоча науковці активно займаються розробками нових алгоритмів, однак питання коректного тестування та порівняння їх залишається недостатньо вивченим. Це призводить до складності: різні дослідники використовують різні методи тестування та різні набори даних і об'єктивно визначити найкращий алгоритм стає практично неможливо [5].

Ситуація ускладнюється тим, що чимало вчених не оприлюднюють здобуті експериментальні дані, не надають розгорнутого опису використаних методик збирання цих даних. Така непрозорість помітно гальмує прогрес у дослідницькій галузі [5]. Проблема ускладнює порівняння різних алгоритмів і перевірку їхніх результатів.

Окремим викликом виступає масштабованість досліджень. Алгоритми, що демонструють відмінні результати на невеликих тестових наборах даних, можуть показувати кардинально іншу поведінку при збільшенні розміру мережі до реальних параметрів. Мережа Ethereum 2.0 продовжує стрімко розвиватися [6]. Таке швидке масштабування висуває додаткові вимоги до алгоритмів розподілу, які мають ефективно функціонувати в умовах постійного зростання кількості учасників.

Практичне значення проблеми підкреслюється зростаючою роллю PoS-систем у критичній інфраструктурі. Децентралізовані фінансові сервіси, побудовані на основі PoS-блокчейнів, обробляють транзакції на мільярди доларів щодня [7]. Державні установи розглядають можливості впровадження блокчейн-технологій для реалізації електронного урядування та цифрових валют центральних банків. Таким чином, у цьому контексті надійність і продуктивність алгоритмів розподілу валідаторів мають подвійне значення. Їхнє значення виходить за межі простої оптимізації. Правильний розподіл валідаторів важливий і для запобігання масовим відмовам, особливо у високонавантажених або критичних системах, таких як блокчейн-платформи та децентралізовані фінансові мережі.

Цінність цього дослідження полягає в прагненні до максимальної відтворюваності. Формування стартових наборів даних дає кілька переваг. По-перше, підвищується точність результатів. По-друге, пришвидшується розвиток галузі. По-третє, створюється міцна основа для впровадження рішень у реальні системи.

2. Аналіз сучасних наукових публікацій і постановка проблеми дослідження

Сфера розподілу валідаторів у PoS є великою й різноплановою. Вона охоплює велику кількість алгоритмів, що поєднують криптографічні засади з реальними потребами розподілу. Існуючі дослідження ґрунтовно висвітлюють теоретичні основи цих алгоритмів [1], [4], [5]. Проте в процесі їх практичного впровадження підходи до експериментальної оцінки алгоритмів у різних дослідженнях суттєво різняться [8].

Процес випадкового вибору валідаторів необхідний для PoS-механізмів консенсусу, оскільки непередбачуваність може тільки забезпечити достатню безпеку, аби зберегти пропорційне відображення на основі ваги стейків [9]. Однак PoS-механізми Ethereum стають повноцінними, незважаючи на це, завдяки алгоритму RANDAO, оскільки випадковість, необхідна для обчислення та перетасування валідаторів, є ключовою функціональністю [4]. Механізм RANDAO в Ethereum реалізується з використанням хеш-функції Кесак-256, seed-значень із попередніх блоків і верифікованих випадкових функцій, що забезпечують прозорість та непередбачуваність системи.

Технічно складнішим є алгоритм Swap-or-Not Shuffle, адаптований із криптографічних досліджень перетасування карт [10]. Алгоритм реалізується через серію умовних обмінів у 90 раундах, що базуються на безпековому аналізі приблизно $4\log_2(N)$ раундів для N валідаторів [11]. Основна перевага полягає в його властивості незалежності від вхідних значень, що дозволяє обчислювати перетасування для окремих індексів без необхідності перемішувати весь список валідаторів, що критично для підтримки легких клієнтів. Алгоритм забезпечує $O(1)$ обчислень для визначення пункту призначення окремого індексу та $O(1)$ зворотне обчислення без потреби зберігати або обробляти повний набір валідаторів. Така характеристика робить Swap-or-Not Shuffle особливо ефективним для масштабних блокчейн-систем, де повне перетасування є обчислювально затратним.

Перехід від класичних алгоритмів до метаевристичних відкриває нові можливості для оптимізації розподілу валідаторів у PoS-системах. Водночас відносно мало досліджень присвячено застосуванню таких методів саме для цієї задачі. Застосування генетичних алгоритмів у сфері блокчейну стосуються гібридного консенсусу та покращень блокчейну в цілому, а не конкретно проблеми вибору валідаторів [12]. Це свідчить про область, яка

ще недостатньо досліджена. У ній метаевристичні методи здатні продемонструвати відчутний прогрес.

Particle Swarm Optimization (PSO) показує обнадійливіші результати у контексті блокчейн оптимізації. Дослідження демонструють Multi-Objective Particle Swarm Optimization (MOPSO) застосування для оптимізації конфігурації блокчейну, показуючи що Парето-оптимальні конфігурації можуть бути знайдені за допомогою PSO підходів [13]. Багатоцільова оптимізація розглядає вибір валідаторів як багатоцільову проблему, враховуючи безпеку, продуктивність, децентралізацію та енергоефективність. Ці дослідження показують перспективу створення гібридних алгоритмів. Вони можуть поєднати сильні сторони метаевристичних із класичними методами розподілу валідаторів.

Існуючі методології експериментальної оцінки мають серйозні обмеження. Вони стають особливо помітними в сучасних дослідженнях PoS-систем. У наукових роботах застосовуються різні підходи – від симуляційних експериментів до аналітичних фреймворків. Зокрема, PoS-симулятор, реалізований на мові Julia [8], та AlphaBlock Framework [14] пропонують спеціалізовані інструменти для моделювання та тестування різних варіантів PoS. Втім, відсутність уніфікованих стандартів щодо їх застосування унеможливорює об'єктивне порівняння отриманих результатів та знижує відтворюваність досліджень.

Метрики оцінки продуктивності суттєво варіюють між різними дослідженнями, що ускладнює об'єктивне порівняння алгоритмів. До первинних метрик належать пропускна здатність, затримка, час фіналізації та споживання енергії. Економічні метрики охоплюють коефіцієнт Джині для оцінки рівності розподілу багатства, структуру розподілу винагород та рівень централізації валідаторів [5]. Безпекові метрики, в свою чергу, зосереджені на стійкості до атак, зокрема економічній вартості атаки 51 %, ефективності механізмів штрафів для проблеми nothing-at-stake та стійкості до long-range атак. Відсутність консенсусу щодо відносної важливості цих метрик призводить до фрагментованості дослідницького ландшафту.

Фундаментальна проблема відтворюваності окреслена дослідниками Algorand і полягає у відсутності уніфікованого фреймворку для аналізу продуктивності блокчейн-систем [15]. На практиці різні проекти застосовують несумісні метрики та експериментальні умови, що робить змістовне порівняння результатів складним. У публікаціях нерідко наводяться показники у сотні, тисячі чи навіть мільйони транзакцій на секунду, однак базові припущення й параметри експериментів залишаються неясними. Дослідники Algorand виокремили сім критичних вимірів, які суттєво впливають на оцінку продуктивності блокчейну і водночас варіюються неконсистентно між роботами: режим доступу (permissioned або permissionless), модель участі (пряма або делегована), тип мережевої інфраструктури, рівень видимості транзакцій, профіль розподілу пропускної здатності, прийнята модель безпеки та використані механізми стиснення.

Дослідження BFT-консенсусу нерідко дають різні результати. Іноді вони навіть суперечать одне одному, особливо у питаннях масштабованості [16]. Причина полягає у відмінностях умов проведення експериментів. Це робить їхнє пряме порівняння складним.

У випадку з алгоритмами розподілу валідаторів у Proof of Stake ситуація ще складніша. Дослідження у цій сфері залишаються фрагментованими. У методології існують значні прогалини, які непросто усунути. Порівняти результати між роботами майже неможливо, адже кожне дослідження застосовує власні метрики, відмінні умови тестування та специфічні реалізації. Моделі й мережі теж можуть відрізнятися значно. Лише комплексний підхід здатен змінити ситуацію. Залишається нагальна потреба в уніфікації метрик, відтворюваних процедур та підвищенні наукової цінності отриманих результатів.

3. Мета і задачі дослідження

Головною метою цього наукового дослідження є розробка методу створення експериментальних датасетів, за допомогою яких можна здійснювати оцінку ефективності алгоритмів розподілу валідаторів у блокчейн-системах на конкретному прикладі роботи механізму консенсусу Proof of Stake. Успішна реалізація дозволить усунути значну проблему відсутності відтворюваних підходів для експериментальної перевірки існуючих алгоритмів розподілу валідаторів і знайти можливість об'єктивно порівняти їх. Цей результат можливо досягти за рахунок вирішення таких задач:

- розробка формального методу створення реалістичних параметрів валідаторів, що враховує важливі аспекти PoS-систем, зокрема розподіл стейків, надійність учасників і мережеві затримки;
- побудова трьох типів експериментальних датасетів різного масштабу (50, 500, 1000 валідаторів) з фіксованими параметрами генерації для забезпечення відтворюваності результатів;
- визначення комплексних критеріїв оцінки алгоритмів розподілу валідаторів, які включають метрики якості розподілу, стабільності результатів та обчислювальної ефективності;
- проведення системних експериментальних досліджень для оцінки ефективності базових та гібридних алгоритмів на розроблених датасетах і отримання порівняльних характеристик продуктивності;
- аналіз масштабованості алгоритмів при зростанні розміру мережі валідаторів для подальшого формулювання науково обґрунтованих рекомендацій щодо їх практичного застосування у реальних PoS-системах.

4. Метод побудови експериментальних датасетів

Через потребу у подоланні фрагментарності підходів до реалізації експериментальної оцінки алгоритмів перерозподілу валідаторів виникла необхідність у формально обґрунтованому методі генерування даних, на основі якої формуються експериментальні набори.

Кожен i -й валідатор у системі характеризується набором параметрів:

$$V_i = (s_i, p_i, r_i, l_{ij}, g_i, q_i, h_i), \quad (1)$$

де s_i – розмір стейку; p_i – продуктивність; r_i – надійність; l_{ij} – мережеву затримку між i -им та j -им валідаторами; g_i – географічне розташування; q_i – якість мережевого з'єднання; h_i – історія штрафів та порушень протоколу.

Для забезпечення реалістичності моделювання кожен валідатор має мінімальний стейк 32 ЕТН, що відповідає правилам Ethereum 2.0 [2]. Такий підхід дозволяє точно відобразити механіку роботи валідаторів і уникнути спотворення результатів через нереалістично великі значення стейку. Інші параметри валідаторів, такі як продуктивність, надійність та мережеві затримки, генеруються у реалістичних діапазонах, що відповідають спектру реальних учасників мережі Ethereum 2.0 [3]. Продуктивність валідаторів генерується як:

$$p_i \sim U(0.5, 1.5) \cdot p_{\text{base}}, \quad (2)$$

де p_{base} – базовий рівень продуктивності. Розподіл відображає різноманітність технічних характеристик обладнання валідаторів. Діапазон від 0.5 до 1.5 обрано на основі емпіричного аналізу статистики Beacon Chain у мережі Ethereum 2.0 [3]. Поточна продуктивність валідаторів оцінюється через ефективність участі у консенсусі, що включає три складові: ефективність атестацій (attester efficiency), ефективність пропозицій блоків (proposer

efficiency) та ефективність участі у синхронізаційних комітетах (sync committee efficiency), які відповідно складають 84,4 %, 12,5 % та 3,1 % від винагород валідатора [17]. Загальна ефективність обчислюється як відношення фактичної винагороди до ідеальної, що дозволяє врахувати відхилення продуктивності від середнього рівня. Нижня межа 0.5 відповідає мінімально допустимій продуктивності для участі у консенсусі, тоді як верхня межа 1.5 характеризує високопродуктивні вузли з оптимізованими конфігураціями.

Надійність валідаторів моделюється у вигляді:

$$r_i \sim U(0.8, 1.0). \quad (3)$$

Межі надійності встановлені згідно з аналізом історії штрафів у Beacon Chain [2, 5]. Валідатори з надійністю нижче 80% мають підвищений ризик slashing через пропуски атестацій. Це значення відповідає практичному порогу надійності для стабільної роботи в PoS-мережах.

Мережеві затримки формуються як симетрична матриця $L = [l_{ij}]_{n \times n}$:

$$l_{ij} = \begin{cases} U(10, 50) \text{ мс, } g_i = g_j \text{ (той самий регіон)} \\ U(50, 200) \text{ мс, } g_i \neq g_j \text{ (різні регіони)} \end{cases}, \quad (4)$$

де g_i та g_j – географічні регіони i -го та j -го валідаторів відповідно. Така структура забезпечує реалістичне моделювання глобальної топології мережі з урахуванням географічного розподілу учасників. Діапазони мережевих затримок базуються на вимірах реальних P2P-мереж блокчейнів. Внутрішньорегіональні затримки 10-50 мс та міжрегіональні 50-200 мс відповідають типовим показникам для забезпечення 12-секундного циклу створення блоків у Ethereum 2.0 [18]. Географічне розташування валідаторів моделюється шляхом їх розподілу між основними регіонами на основі актуальних досліджень мережі Ethereum. Згідно з [19], фактичний розподіл має такий вигляд: Європа – 46.79 %, Північна Америка – 38.05 %, Азія – 9.95 %, Океанія – 4.16 %.

Якість мережевого з'єднання валідаторів генерується в діапазоні від 10 до 1000 Мбіт/с для моделювання різноманітності учасників мережі. Мінімальна вимога 10 Мбіт/с забезпечує передачу даних у межах 12-секундного циклу Ethereum 2.0 [6]. Розподіл налаштовано так, щоб більшість валідаторів мала середні показники, а меншість – високі, що відповідає типовій структурі децентралізованих мереж з різноманітними учасниками.

Історія штрафів валідаторів моделюється як накопичувальний показник порушень протоколу з урахуванням часового фактору. Система враховує чотири типи подій: відсутність порушень, пропуски атестацій, некоректні атестації та найсерйозніше порушення – подвійне підписання блоків. Недавні події мають більшу вагу при розрахунку загального показника, що відповідає принципам репутаційних систем у блокчейні. Частота порушень визначається згідно зі статистикою Beacon Chain [2], де приблизно 2 % валідаторів мають зафіксовані порушення протягом 100 епох. Валідатори з високим рівнем накопичених штрафів класифікуються як потенційно ненадійні для цілей Byzantine Fault Tolerance [16], що впливає на їх розподіл між комітетами для забезпечення безпеки мережі.

Для гарантування повної відтворюваності експериментів встановлено строгий протокол фіксації початкових значень генераторів псевдовипадкових чисел. Кожному набору даних призначається унікальний ідентифікатор $seed(n, version)$, що включає розмір набору та версію генератора. Всі згенеровані датасети зберігаються у форматі з метаданими про параметри генерації та контрольними сумами для верифікації цілісності.

Було сформовано три групи наборів даних, які покривають широку палітру практичних прикладів підходу до системи. Сценарій з малою кількістю валідаторів

(приблизно 50) моделював приватні та консорціумні блокчейн-системи. Середній із 500 валідаторів відображав підходи до корпоративних систем блоку або тестових мереж та забезпечував співвідношення рівності обчислення, якості і доступності. Набір даних із 1000 валідаторів використовувався для моделювання роботи алгоритмів у міру збільшення масштабу та відображав характерні особливості публічних блокчейн-мереж. Хоча ця кількість значно менша за реальну мережу Ethereum, вона все ж дозволила вивчати поведінку алгоритмів у порівняно великих підмережах. Це дало змогу оцінити їхню масштабованість.

Крім того, слід обирати значення seed так, щоб забезпечити статистичну незалежність, псевдовипадковість, передбачуваність і масштабованість створюваних серій наборів даних. Ці значення не можуть бути кратними або степенями двійки. Рекомендується псевдовипадковий інтервал, який не є надто близьким до будь-якого з іншого цілого числа, і перевіряти незалежність за допомогою коефіцієнтів кореляції. Seed можна обирати послідовно, за допомогою хешування параметрів набору даних або у ієрархічній формі, що враховує категорії експерименту, розмір набору даних та індекс екземпляра.

Перевірка seed передбачає спочатку оцінку рівномірності розподілу даних, а потім аналіз кореляцій та повторюваних патернів. Такий підхід закладає міцну основу для подальшого розширення датасетів і проведення нових експериментів. Він має особливе значення для забезпечення відтворюваності результатів у PoS-системах.

Для всебічної оцінки алгоритмів визначено систему комплексних критеріїв, що охоплює якість розподілу, стабільність результатів та обчислювальну ефективність. Критерієм якості є значення узагальненого критерію Z [20]:

$$Z = \alpha \cdot F_1(X) - \beta \cdot F_2(X) - \gamma \cdot F_3(X), \quad (5)$$

де $F_1(X)$ максимізує пропускну здатність системи, $F_2(X)$ мінімізує дисбаланс навантаження, $F_3(X)$ мінімізує мережеві затримки. Значення Z дозволяє визначати оптимальних валідаторів для включення до комітету: валідатор зараховується, якщо його Z перевищує встановлений поріг або є найвищим серед кандидатів. Значення вагових коефіцієнтів $\alpha = 1.0$, $\beta = 0.3$, $\gamma = 0.2$ встановлено на основі аналізу подібних багатокритеріальних задач оптимізації в розподілених системах. Співвідношення узгоджується з рекомендаціями для балансування продуктивності та надійності в консенсусних протоколах, де пропускну здатність має найвищий пріоритет [20]. Додатково проведено валідацію через серію експериментів з альтернативними конфігураціями коефіцієнтів, які підтвердили оптимальність обраних значень для досліджуваних масштабів мережі.

Стабільність алгоритмів оцінюється через стандартне відхилення $\sigma(Z)$, обчислене за серією з $N = 10$ незалежних запусків:

$$\sigma(Z) = \sqrt{\frac{1}{N-1} \sum_{k=1}^N (Z_k - \bar{Z})^2}, \quad (6)$$

де Z_k – результат k -го запуску.

Обчислювальна ефективність оцінюється за двома показниками: часом виконання $T(n)$ та кількістю обчислень узагальненого критерію. Для дослідження масштабованості використовується метрика:

$$S(n) = \frac{T(n)}{T(n_0)}, \quad (7)$$

де $T(n)$ – час виконання задачі розміру n ; $T(n_0)$ – базовий час виконання для референтного розміру задачі n_0 . Ця метрика відображає, наскільки зростає час виконання із збільшенням розміру задачі, і дозволяє кількісно оцінити масштабованість алгоритмів.

Розроблено протокол для системного оцінювання ефективності алгоритмів розподілу валідаторів, який передбачає тестування чотирьох алгоритмів:

- Гібридний PSO+LS — поєднує метод рою частинок з локальним пошуком [20], [21];
- Random+Repair — базовий алгоритм із випадковою генерацією та корекцією;
- Greedy — жадібний алгоритм [22];
- EthShuffle — адаптація механізму RANDAO з Ethereum 2.0 [4].

Кожен алгоритм було протестовано на всіх наборах даних із усіма метриками десять разів. Такий підхід забезпечив статистичну надійність результатів, зменшуючи вплив випадковості та дозволяючи оцінити дисперсію, при цьому не створюючи надмірних обчислювальних витрат. Попереднє тестування показало, що стандартне відхилення результатів стабілізується після 8-10 запусків, і додаткові запуски не впливають на висновки щодо відносної ефективності алгоритмів. Отже, десятикратне повторення експериментів є оптимальним компромісом між достовірністю результатів і обчислювальною ефективністю. Тому кожен алгоритм було протестовано на кожному наборі даних із усіма метриками десять разів, що становить оптимальний компроміс між достовірністю й ефективністю.

Тести було проведено з використанням Python v3.9. Для роботи з матрицями було використано NumPy, а паралельні обчислення організовано через multiprocessing. Результати було збережено у вигляді ієрархічної таблиці, що допомагає легше аналізувати дані.

5. Результати дослідження

Перед проведенням основних експериментів було проаналізовано характеристики згенерованих наборів даних для підтвердження їх реалістичності та відтворюваності. Гістограми розподілу стейків (рис. 1) свідчать про логнормальний характер розподілу. Для малого набору ($n=50$) спостерігається більша варіативність розподілу, тоді як для великих наборів ($n=500$, $n=1000$) розподіл стає стабільнішим.

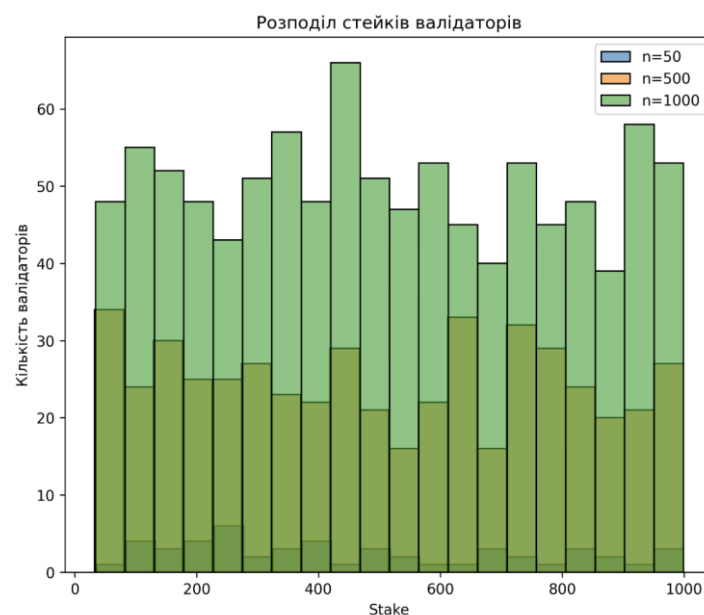


Рис. 1. Розподіл стейків валідаторів для трьох масштабів

На діаграмі розсіювання, наведеній на рис. 2, відображено залежності якості рішень від часу виконання чотирьох алгоритмів розподілу валідаторів у PoS-системах. PSO+LS забезпечує високу якість оптимізації серед досліджуваних алгоритмів, однак за рахунок значно більшого часу виконання. Random+Repair показує помірну швидкість виконання з варіативними результатами. EthShuffle демонструє середню ефективність при швидкому виконанні. Алгоритм Greedy вирізняється високою швидкістю виконання та стабільними детерміністичними результатами, однак ефективність його рішень суттєво залежить від масштабу задачі. Використання логарифмічної шкали часу демонструє баланс між ефективністю та швидкістю виконання. Такий підхід дозволяє наочно оцінити компромісні рішення.

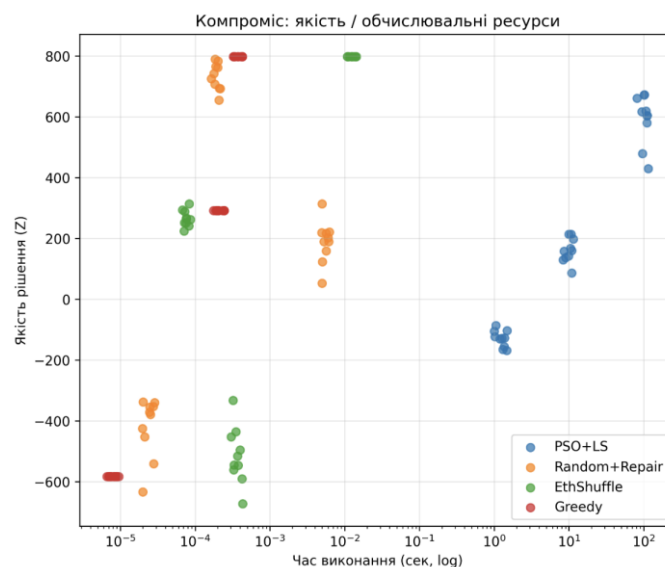


Рис. 2. Діаграма компромісу «якість/обчислювальні ресурси»

Квадратичне зростання часу виконання PSO+LS робить його придатним для офлайн-планування розподілу валідаторів, де якість оптимізації важливіша за швидкість. Базові алгоритми демонструють лінійну або сублінійну складність, що робить їх придатнішими для задач реального часу, але з суттєвою втратою якості оптимізації.

Матриця кореляції між результатами виконання алгоритмів (рис. 3) демонструє використання фіксованих seed-значень для забезпечення відтворюваності експериментів. Алгоритм Greedy не має кореляційних значень (білі комірки), оскільки він є детерміністичним і завжди повертає однаковий результат. Для інших пар алгоритмів коефіцієнти кореляції перебувають у діапазоні від -0.11 до 0.03 , а відповідні p-value перевищують 0.05 , що вказує на відсутність статистично значущих лінійних зв'язків між результатами різних алгоритмів. Низькі коефіцієнти кореляції між результатами алгоритмів відповідають очікуванням для методів з різними принципами оптимізації.

Аналіз кумулятивної функції розподілу (CDF) якості рішень (див. рис. 4) показує помітні відмінності у поведінці алгоритмів. Зокрема, алгоритм PSO+LS характеризується найширшим діапазоном отриманих результатів, що свідчить про його здатність охоплювати різні траєкторії пошуку та знаходити рішення різної якості. Метод Random+Repair займає проміжне положення: він має вищу стабільність порівняно з PSO+LS, однак поступається йому за різноманітністю результатів. Алгоритм EthShuffle демонструє стрімкіший ріст кумулятивної кривої, що вказує на меншу варіативність

результатів і концентрацію рішень у відносно вузькому інтервалі якості. Найобмеженішою є поведінка алгоритму Greedy: крива поведінки має характерні стрибки та демонструє мінімальну варіативність, що узгоджується з його детерміністичною природою.

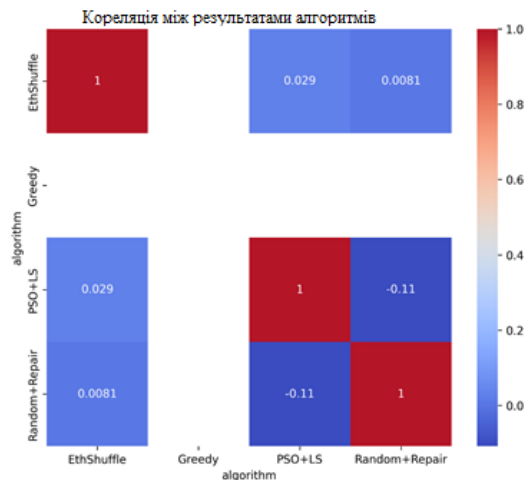


Рис. 3. Матриця кореляції між результатами алгоритмів

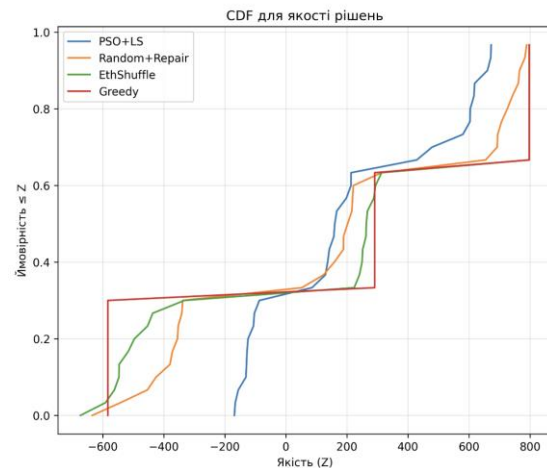


Рис. 4. Кумулятивна функція розподілу (CDF) якості рішень

6. Обговорення результатів дослідження

За результатами дослідження представлено метод формування експериментальних даних, який забезпечує відтворюваність при дослідженні алгоритмів розподілу валідаторів у PoS-системах. Раніше запропоновані підходи часто не містили достатньо деталей про умови проведення експериментів. На їх основі створювалися нові алгоритми, проте результати важко було порівнювати. Використання запропонованого методу дозволяє виконувати серію експериментів із заздалегідь визначеними параметрами, що гарантує стабільність і порівнянність результатів.

Оцінювання алгоритмів здійснювалося за кількома показниками: ефективність розподілу, стабільність роботи та обчислювальні витрати. Такий підхід дозволив отримати повне уявлення про їхні сильні й слабкі сторони та зробити обґрунтований вибір для конкретних сценаріїв застосування.

Результати експерименту виявили особливу поведінку алгоритмів на різних масштабах мережі. Ефективність алгоритму Greedy залежить від масштабу: на середніх масштабах мережі він показує конкурентоспроможні результати, тоді як на малих наборах даних його продуктивність знижується. Така поведінка пояснюється детерміністичною природою алгоритму, який завжди приймає локально оптимальні рішення без урахування глобального контексту. Завдяки масштабованому аналізу на трьох рівнях (50, 500, 1000 валідаторів) стали видимими нелінійні залежності продуктивності алгоритму від розміру мережі. Наприклад, PSO+LS демонструє високу якість оптимізації на всіх масштабах мережі, проте квадратичне зростання часу виконання обмежує його застосування виключно сферою офлайн-планування, де час не є критичним фактором.

Статистична валідація згенерованих датасетів підтвердила основні характеристики синтетичних даних. Коефіцієнт Джині [5] для розподілу стейків складає 0.31 ± 0.01 , що свідчить про помірну нерівність, характерну для децентралізованих мереж. Мережеві параметри показали реалістичні значення: латентність між валідаторами одного регіону становить 30 ± 11 мс, між різними регіонами – 125 ± 43 мс, що відповідає типовим показникам глобальних P2P-мереж. Частка потенційно ненадійних валідаторів (24-25 %)

міститься в безпечному діапазоні для Byzantine Fault Tolerance ($< 33\%$). Географічний розподіл валідаторів рівномірно покриває чотири регіони з варіацією 18-24 %, що забезпечує достатню децентралізацію. Слабка кореляція між стейком та продуктивністю ($|r| < 0.15$) відображає реалістичну незалежність економічних можливостей учасників від їхніх технічних характеристик, що є типовим для відкритих PoS-систем, де розмір стейку не гарантує якість обладнання.

Разом із перевагами важливо окреслити і межі застосовності. Експериментальні дані отримано не з реальної мережі, а згенеровано програмно. Для побудови використано статистичні параметри, виділені з аналізу існуючих PoS-систем. В основу покладено емпіричні характеристики Ethereum 2.0, що дозволило сформувати синтетичні датасети з високим рівнем реалістичності. Водночас збережено повну відтворюваність експериментів, що забезпечує можливість їх перевірки та повторення. Синтетичні датасети відтворюють статичні характеристики валідаторів (розподіл стейків, продуктивність, надійність, географічне розташування), однак мають обмеження у відтворенні динамічних аспектів реальних мереж. Зокрема, не враховуються часові залежності поведінки валідаторів, економічні стимули, що змінюються, та адаптивні стратегії учасників мережі. Крім того, у дослідженні використано 1000 валідаторів, що є спрощенням у порівнянні з масштабами мереж на кшталт Ethereum 2.0, однак цього обсягу достатньо для перевірки коректності та працездатності запропонованого методу.

Що стосується обчислень, PSO+LS вимагає значних ресурсів і часу виконання, які зростають зі збільшенням розміру мережі. Для офлайн-експериментів це прийнятно, але для систем із швидкою реконфігурацією комітетів може бути проблемою. Таке обмеження робить PSO+LS непридатним для адаптації в режимі реального часу. Алгоритм доцільно використовувати для довгострокового планування конфігурації комітетів або періодичної оптимізації розподілу між епохами.

Попри зазначені обмеження синтетичних даних, експериментальні результати дозволяють сформулювати практичні рекомендації для різних сценаріїв застосування алгоритмів у реальних PoS-системах. Виявлені закономірності поведінки алгоритмів на різних масштабах мережі створюють основу для їх цільового використання відповідно до специфічних вимог системи. На основі результатів сформульовано практичні рекомендації: для критичних застосувань, де можливе періодичне офлайн-планування конфігурації (наприклад, раз на добу або між великими оновленнями мережі), рекомендується PSO+LS; для систем реального часу, що потребують балансу між якістю та швидкістю – EthShuffle; для швидкої реконфігурації мереж – Random+Repair, для простих систем з обмеженими ресурсами, де швидкість критична, а оптимальність вторинна, – Greedy.

У майбутніх дослідженнях варто зосередитися на розробці адаптивних алгоритмів наступного покоління, що забезпечуватимуть автоматичну реконфігурацію структури комітетів у відповідь на зміни параметрів мережі. Такі алгоритми мають поєднувати функції виявлення аномалій, прогнозування навантаження та оптимізації розподілу ресурсів, враховуючи економічні стимули, різноманітні рівні продуктивності валідаторів та поведінкові особливості учасників при збереженні безперервності роботи механізму консенсусу.

Важливим кроком стане створення відкритого репозиторію стандартизованих датасетів та практичних реалізацій алгоритмів. Це полегшить стандартизацію експериментів, забезпечить високу повторюваність результатів та дозволить дослідникам зосередитися на розробці нових методів замість повторного відтворення відомих рішень.

Доцільно також розширити дослідження на інші PoS-протоколи (Cardano, Polkadot,

Solana) для оцінки універсальності підходу та його потенціалу масштабування в різних блокчейн-екосистемах. Це відкриває шлях до створення універсального фреймворку для оцінки алгоритмів розподілу валідаторів.

7. Висновки

Запропонований метод формування стандартизованих експериментальних наборів даних для порівняльної оцінки алгоритмів розподілу валідаторів у PoS-системах генерує сім ключових характеристик на основі статистичних параметрів Ethereum 2.0 з використанням фіксованих початкових значень. На відміну від існуючих досліджень, де параметри задаються менш формалізовано, запропоноване рішення забезпечує повну специфікацію всіх параметрів для детермінованої відтворюваності. Валідація підтвердила реалістичність методу через безпечну частку ненадійних валідаторів, незалежність економічних і технічних характеристик та відповідність мережесих затримок реальним показникам глобальних мереж.

Сформовано три стандартизовані набори даних різного масштабу з унікальними ідентифікаторами та метаданими. Датасети охоплюють сценарії від консорціумних до публічних мереж. Підхід мінімізує проблему несумісності експериментальних умов між дослідженнями. Статистичний аналіз згенерованих датасетів підтвердив логнормальний розподіл стейків, географічну децентралізацію між регіонами та відсутність кореляцій між результатами різних алгоритмів.

Розроблено систему комплексних критеріїв оцінювання: якість розподілу, стабільність результатів та обчислювальна ефективність. Система дозволяє оцінити компроміс між якістю оптимізації та обчислювальними витратами. На відміну від обмеженої оцінки в один або два виміри, вона забезпечує багатовимірний аналіз з урахуванням пропускну здатності, збалансованості навантаження та мережесих затримок.

Виконано системне порівняльне дослідження чотирьох алгоритмів на трьох масштабах з десятикратним повторенням. Виявлено масштабно-залежну поведінку: PSO+LS досягає найвищої якості оптимізації, але високі обчислювальні витрати обмежують його застосування періодичним офлайн-плануванням; EthShuffle забезпечує стабільні результати середньої якості; Random+Repair характеризується швидкістю виконання при варіативних результатах; Greedy демонструє максимальну швидкість та детерміністичність, однак його відносна ефективність знижується при малих масштабах мережі.

Проведено аналіз масштабованості, що виявив нелінійні залежності продуктивності від розміру мережі. На основі отриманих результатів сформульовано практичні рекомендації для відповідних сценаріїв застосування залежно від вимог до якості оптимізації та швидкості виконання. Запропоновані датасети та протоколи тестування створюють базовий інструментарій для об'єктивної оцінки нових алгоритмічних рішень у PoS-системах, забезпечуючи відтворюваність експериментів та можливість прямого порівняння результатів різних досліджень.

Перелік посилань:

1. E. Kapengut and B. Mizrach, "An Event Study of the Ethereum Transition to Proof-of-Stake," *Commodities*, vol. 2, no. 2, pp. 96–110, Jun. 2023, doi: <https://doi.org/10.3390/commodities2020006>.
2. Ethereum Foundation, The Beacon Chain, 2024. URL: <https://ethereum.org/en/upgrades/beacon-chain/>.
3. Ethan "ethos.dev" (Joseph Ch), The Beacon Chain Ethereum 2.0 explainer you need to read first, ethos.dev, 2022. URL: <https://ethos.dev/beacon-chain>.
4. B. Edgington, Randomness, in: **Building Blocks**, Eth2 Book, 2025. URL: https://eth2book.info/latest/part2/building_blocks/randomness/.

5. T. Yan, S. Li, B. Kraner, L. Zhang, and C. J. Tessone, "A Data Engineering Framework for Ethereum Beacon Chain Rewards: From Data Collection to Decentralization Metrics," *Scientific Data*, vol. 12, no. 1, Mar. 2025, doi: <https://doi.org/10.1038/s41597-025-04623-7>.
6. Axon Trade, "Ethereum in 2025: Network, Usage, and Upgrades," Axon Trade, 2025. <https://axon.trade/ethereum-in-2025> (accessed Aug. 25, 2025).
7. BlockByte, "The Fed's Pro-Crypto Pivot: How DeFi and Stablecoins Are Reshaping U.S. Financial Infrastructure," *Ainvest*, Aug. 21, 2025. <https://www.ainvest.com/news/fed-pro-crypto-pivot-defi-stablecoins-reshaping-financial-infrastructure-2508> (accessed Aug. 25, 2025).
8. A. Leporati and L. Rovida, "Looking for Stability in Proof-of-Stake based Consensus Mechanisms," *Blockchain Research and Applications*, pp. 100222–100222, Jul. 2024, doi: <https://doi.org/10.1016/j.bcr.2024.100222>.
9. Q. H. Nguyen, A. Cronje, and M. Kong, "Fast Stochastic Peer Selection in Proof-of-Stake Protocols," *arXiv (Cornell University)*, Jan. 2019, doi: <https://doi.org/10.48550/arxiv.1911.04629>.
10. V. T. Hoang, B. Morris, and P. Rogaway, "An Enciphering Scheme Based on a Card Shuffle," *arXiv (Cornell University)*, Jan. 2012, doi: <https://doi.org/10.48550/arxiv.1208.1176>.
11. "Upgrading Ethereum | 2.9.4 Shuffling," *Eth2book.info*, 2025. https://eth2book.info/latest/part2/building_blocks/shuffling/ (accessed Aug. 27, 2025).
12. K. Venkatesan and S. B. Rahayu, "Blockchain security enhancement: an approach towards hybrid consensus algorithms and machine learning techniques," *Scientific Reports*, vol. 14, no. 1, p. 1149, Jan. 2024, doi: <https://doi.org/10.1038/s41598-024-51578-7>.
13. X. Gu, X. Wang, Y. Ma, Z. Chen, and S. Huang, "Performance Optimization for Information Sharing Process of BlockIoV Based on Multi-Objective Particle Swarm," pp. 172–183, Oct. 2023, doi: <https://doi.org/10.1109/qrs60937.2023.00026>.
14. H. Xiang, Z. Ren, Z. Zhou, N. Wang, and H. Jin, "AlphaBlock: An Evaluation Framework for Blockchain Consensus Protocols," *arXiv (Cornell University)*, Jan. 2020, doi: <https://doi.org/10.48550/arxiv.2007.13289>.
15. "Towards a Unified Metric for Performance Evaluation of Proof-of-Stake Blockchains," *Algorandtechnologies.com*, 2018. <https://algorandtechnologies.com/news/towards-a-unified-metric-for-performance-evaluation-of-proof-of-stake-blockchains> (accessed Aug. 27, 2025).
16. W. Zhong et al., "Byzantine Fault-Tolerant Consensus Algorithms: A Survey," *Electronics*, vol. 12, no. 18, pp. 3801–3801, Sep. 2023, doi: <https://doi.org/10.3390/electronics12183801>.
17. "Metric: Validator Efficiency | beaconcha.in Knowledge Base," *Beaconcha.in*, Jun. 04, 2025. <https://kb.beaconcha.in/v2beta/metric-validator-efficiency> (accessed Aug. 28, 2025).
18. A. Lebedev and V. Gramoli, "On the Relevance of Blockchain Evaluations on Bare Metal," *arXiv (Cornell University)*, Jan. 2023, doi: <https://doi.org/10.48550/arxiv.2311.09440>.
19. "Deanonymizing Ethereum Validators: The P2P Network Has a Privacy Issue," *Arxiv.org*, 2024. <https://arxiv.org/html/2409.04366v2> (accessed Aug. 28, 2025).
20. Y. Demenko, I. Grebennik, M. Kolmykov, Optimization of Validator Assignment to Committees in Proof-of-Stake Blockchain Systems, in: *Modern Information Technologies and Artificial Intelligence Systems: Proceedings of the 1st International Scientific and Practical Conference MIT@AIS-2025*, Kharkiv–Yaremche, May 19–22, 2025, KhNURE, Kharkiv, 2025, pp. 46–49. (In Ukrainian)
21. J. Kennedy, Particle Swarm Optimization, in: C. Sammut, G. I. Webb (Eds.), *Encyclopedia of Machine Learning*, Springer, Boston, MA, 2011. doi:10.1007/978-0-387-30164-8_630.
22. Y. Wang, "Review on greedy algorithm," *Theoretical and natural science*, vol. 14, no. 1, pp. 233–239, Nov. 2023, doi: <https://doi.org/10.54254/2753-8818/14/20241041>.

Надійшла до редколегії 08.09.2025 р.

Деменко Євгеній Євгенович, аспірант кафедри СТ ХНУРЕ, м. Харків, Україна, e-mail: yevhenii.demenko@nure.ua, ORCID: <https://orcid.org/0000-0002-5699-1400>

Гребеннік Ігор Валерійович, доктор технічних наук, професор, завідувач кафедри СТ ХНУРЕ, м. Харків, Україна, e-mail: igor.grebennik@nure.ua, ORCID: <https://orcid.org/0000-0003-3716-9638>

Колмиков Максим Миколайович, кандидат технічних наук, старший науковий співробітник, ХНУПС, м. Харків, Україна, e-mail: kolmykov.max@gmail.com, ORCID: <https://orcid.org/0000-0003-1972-3543>

УДК 004.932.2:681.7.06

DOI: 10.30837/0135-1710.2025.186.083

*Д.М. КРИЦЬКИЙ, Д.А. ОНІЩУК, О.О. ВАЛЮЖЕНИЧ***СИСТЕМА НА ОСНОВІ RGB-ДАТЧИКА ДЛЯ ВИЯВЛЕННЯ КАМУФЛЯЖУ У ВІЙСЬКОВОМУ СЕРЕДОВИЩІ**

Дослідження присвячено розробці недорогої оптичної системи на основі RGB-датчика для виявлення замаскованих об'єктів у військовому середовищі. Метою дослідження є створення прототипу апаратно-програмної системи виявлення камуфляжу на основі RGB-датчика, що працює у видимому спектрі світла, з подальшою експериментальною перевіркою його ефективності у різних природних умовах. Система відзначається низьким енергоспоживанням, мобільністю та може бути інтегрована у портативні або безпілотні розвідувальні платформи. Для досягнення поставленої мети було вирішено такі задачі: розробка апаратної структури сенсорної системи на основі доступних компонентів; реалізація алгоритмів моделювання фону, виявлення кольорових аномалій та фільтрації шумів з використанням колориметричного аналізу у просторі RGB; проведення серії натурних експериментів у різних природних умовах – лісистій місцевості, степовій зоні та кам'янистому ландшафті; оцінка точності, стабільності та обмежень запропонованого рішення у порівнянні з тепловізорами та інфрачервоними камерами зі штучним інтелектом. Методологія базується на використанні RGB-датчика TCS34725 із вбудованим інфрачервоним фільтром та 16-бітним АЦП у поєднанні з мікроконтролером ESP32, який забезпечує обробку даних у реальному часі, бездротову передачу інформації та автономну роботу. Алгоритм виявлення ґрунтується на формуванні кольорового профілю фону, аналізі відхилень за евклідовою відстанню, нормалізації RGB-значень, медіанній фільтрації та адаптивному пороговому визначенні, що забезпечує стійкість до змін середовища. Результати випробувань показали, що запропонована система здатна з високою точністю ідентифікувати цифровий «піксельний» камуфляж у лісистій місцевості (86 %), камуфляж «мультикам» у степових умовах (78 %) та однотонний оливковий камуфляж на кам'янистому фоні (65 %). Порівняння з тепловізійними та інфрачервоними системами підтвердило значні переваги RGB-рішення за вартістю (менше \$ 20), енергоефективністю та мобільністю, хоча функціонування можливе лише у денний час. У висновках наголошується на доцільності застосування RGB-систем виявлення як економічно вигідного допоміжного інструменту для розвідки та охоронних завдань. Наукова новизна роботи полягає в інтеграції простого й доступного RGB-сенсора з адаптивними алгоритмами моделювання фону та аналізу кольорових відхилень, що дозволяє створювати малопотужні платформи для виявлення камуфляжу. На відміну від традиційних підходів, запропонована система відкриває перспективи розгортання розподіленої мережі дешевих сенсорних вузлів або інтеграції у безпілотні літальні апарати для оперативного моніторингу великих територій.

1. Вступ

Сучасні військові операції вимагають високої ефективності систем спостереження та розвідки, особливо у питаннях виявлення замаскованих об'єктів. Використання камуфляжних технологій, які імітують навколишнє середовище, значно ускладнює виявлення цілей за допомогою традиційних оптичних засобів. В умовах різноманітних природних ландшафтів, таких як ліси, степи або кам'яні ділянки, камуфляжні покриття можуть ефективно маскувати військову техніку та особовий склад, що створює серйозні перешкоди для розвідки.

На сьогодні основними інструментами виявлення замаскованих об'єктів є тепловізійні прилади та інфрачервоні камери, часто доповнені алгоритмами штучного інтелекту (AI). Ці технології дозволяють ефективно працювати у нічний час та за складних погодних умов, однак характеризуються високою вартістю, значним енергоспоживанням

та обмеженою мобільністю, що ускладнює їх масове застосування в польових умовах. Враховуючи зазначене вище, зростає актуальність розробки доступних, енергоефективних та компактних систем, які працюють у видимому спектрі світла. RGB-сенсори, що аналізують колірні характеристики об'єктів у видимому діапазоні, є перспективною альтернативою для денного виявлення замаскованих цілей. Завдяки простоті конструкції та невисокій вартості такі системи можуть бути інтегровані у мобільні платформи, наприклад, безпілотні літальні апарати або переносні розвідувальні пристрої.

Метою дослідження є розробка та експериментальна оцінка недорогої системи на основі RGB-датчика, призначеної для виявлення камуфляжу у військових умовах. В ході дослідження детально розглянуто особливості конструкції сенсорного модуля, застосовані алгоритми обробки кольорових зображень, а також результати тестування системи в різних природних умовах. Особлива увага приділяється виявленню трьох основних типів камуфляжу: «мультикам», «піксель» та однотонного, які широко використовуються у військовій практиці. Крім того, проведено порівняльний аналіз ефективності запропонованої системи із традиційними засобами виявлення – тепловізійними приладами та інфрачервоними камерами. з AI. Це порівняння дозволяє виокремити переваги та обмеження RGB-сенсорів у контексті реальних військових задач та сформулювати рекомендації щодо їх подальшого застосування.

2. Аналіз сучасних наукових публікацій і постановка проблеми дослідження

Сучасні технології виявлення замаскованих об'єктів активно розвиваються у напрямку використання глибинного навчання, мультиспектральної обробки зображень та сенсорних систем нового покоління. Основна увага приділяється підвищенню точності, зниженню енергоспоживання та адаптивності алгоритмів до змін середовища. Важливою умовою успішної роботи систем виявлення є стійкість до фонових змін, зокрема варіацій освітлення, текстури та кольору. У роботі [1] було проаналізовано продуктивність сучасних алгоритмів виявлення об'єктів у камуфляжі з використанням багаторівневого злиття ознак. Автори відзначили, що поєднання контекстної інформації з локальними контурами значно підвищує здатність до ідентифікації цілей зі слабким контрастом. В дослідженні [2] запропоновано підхід на основі імітації людського візуального фокусу – система використовує спеціальні модулі фокусування, які «збільшують» ділянки з потенційною аномалією кольору. Така стратегія виявилася ефективною при виявленні малопомітних об'єктів на нерівномірному тлі, що актуально для камуфляжу. Над вирішенням проблеми багатомасштабного аналізу працювали дослідники у [3], де було запропоновано метод агрегації ознак із використанням спеціалізованих блоків орієнтації на межі. Це дозволило суттєво підвищити точність при роботі з об'єктами, які змінюють розміри або форму залежно від ракурсу спостереження. Аналогічно, у роботі [4] розроблено модель, що використовує контурно-орієнтовані фільтри для детектування слабких структур у фонових зображеннях. Для зниження обчислювального навантаження запропоновано полегшені моделі, здатні працювати на мобільних пристроях і у вбудованих системах. Зокрема, в [5] описана система, що базується на поділі зображення на контекстні та текстурні гілки, кожна з яких обробляється окремо з подальшим об'єднанням результатів. Такий підхід дозволив досягти високої швидкості при мінімальних затратах ресурсів, що є критичним у польових умовах. Важливою тенденцією також є інтеграція методів машинного навчання з класичними підходами аналізу RGB-даних. У роботі [6] доведено, що використання адаптивного порогового аналізу у поєднанні з геометричними фільтрами дозволяє суттєво знизити кількість хибних спрацьовувань при виявленні об'єктів, схожих за кольором до фону.

Загалом, огляд підтверджує перспективність розвитку RGB-систем, орієнтованих на

денну експлуатацію. Вони можуть виступати як окремі рішення в умовах обмежених ресурсів або доповнювати мультиспектральні системи. Однак подальше вдосконалення вимагає врахування складних зовнішніх факторів, таких як тінь, зміна освітлення, рух листя тощо. Саме тому інтеграція простих колориметричних методів із адаптивними алгоритмами та машинним навчанням є ключовим напрямом подальших досліджень.

Узагальнення проведеного аналізу показало, що попри значний прогрес у розвитку алгоритмів виявлення замаскованих об'єктів, більшість сучасних рішень залишається енерговитратними та складними для реалізації у мобільних пристроях. Тому актуальною науково-практичною проблемою є створення простої та недорогої RGB-системи для денного виявлення камуфляжу, що забезпечує прийнятну точність за мінімальних ресурсів.

3. Мета і задачі дослідження

Метою дослідження є створення прототипу апаратно-програмної системи виявлення камуфляжу на основі RGB-датчика, з подальшою експериментальною перевіркою у різних природних умовах.

Для досягнення цієї мети необхідно вирішити такі задачі:

- розробити структуру апаратної частини системи на базі доступних і недорогих компонентів, забезпечивши її енергоефективність, компактність та можливість автономної роботи.
- розробити алгоритми виявлення аномалій на основі аналізу кольорової інформації в системі координат RGB (Red, Green, Blue), включаючи формування фонових профілів, обчислення відхилень, нормалізацію значень та фільтрацію шумів.
- провести серію тестувань у різних природних умовах (лісові, степові та кам'яні ландшафти) із використанням різних типів камуфляжу.
- оцінити ефективність системи з точки зору точності, швидкодії та стабільності у порівнянні з тепловізорами та інфрачервоними камерами з елементами штучного інтелекту.

4. Матеріали дослідження

Об'єктом дослідження є процес виявлення камуфляжу у видимому спектрі світла за допомогою RGB-датчиків у складі мобільних сенсорних систем.

Предметом дослідження є розробка недорогої оптичної системи на основі RGB-датчика для виявлення замаскованих об'єктів у військовому середовищі.

Основною гіпотезою дослідження є припущення, що навіть недорогі RGB-сенсори здатні забезпечити достовірне виявлення камуфляжу у денних умовах без залучення складних інфрачервоних або мультиспектральних каналів, якщо використати адаптивні методи нормалізації та порогового аналізу кольору.

Методологія реалізації системи базується на використанні дешевого, енергоефективного та доступного обладнання. Основним елементом сенсорної частини є RGB-датчик TCS34725, який забезпечує вимірювання інтенсивності червоного (R), зеленого (G) та синього (B) світла. Цей сенсор має інтегрований інфрачервоний фільтр, що дозволяє зменшити вплив паразитного інфрачервоного випромінювання та підвищити точність вимірювань у денних умовах. Додатково датчик оснащений АЦП високої роздільної здатності (16 біт), що дозволяє отримувати якісні вимірювання кольору [7].

Як обчислювальний модуль обрано ESP32 через його низьке енергоспоживання, вбудовані Wi-Fi і Bluetooth, а також підтримку енергозберігаючих режимів сну [8]. ESP32 зчитує дані з сенсора через інтерфейс I2C, обробляє інформацію в реальному часі, виконує попереднє усереднення, фільтрацію та приймає рішення про виявлення об'єктів. Розроблений мікропрограмний код дозволяє реалізувати фоновий аналіз сцени, тобто формування статистичного образу навколишнього середовища на основі серії вимірювань

RGB. Зібрана інформація зберігається локально на карті пам'яті microSD або передається бездротовими засобами (Wi-Fi, LoRa) до центру обробки [9]. Система може працювати автономно від акумулятора або вбудованого джерела живлення з сонячною панеллю. Реалізація такої платформи дозволяє створювати мережу дешевих сенсорних вузлів для моніторингу великих ділянок території у режимі реального часу [10].

Колориметричний аналіз реалізовано на основі класичної моделі переходу RGB $\rightarrow$ HSV, яка забезпечує інваріантність до змін освітлення. Для оцінювання ступеня відхилення кольору між поточними та фоновими вимірюваннями застосовано евклідову метрику у просторі RGB. Подальше опрацювання включає фільтрацію шумів медіанним фільтром і нормалізацію значень RGB у відсотковій формі для підвищення стійкості до коливань яскравості. Адаптивне оновлення порогів спрацьовування виконується з урахуванням середньої дисперсії фонових сигналів, що дозволяє системі автоматично підлаштовуватись до поточних умов освітлення. Передбачено інтеграцію простих алгоритмів машинного навчання, зокрема k-Nearest Neighbors, для покращення класифікації камуфляжу у складних фонових сценаріях.

5. Алгоритм виявлення камуфляжу

5.1. Концепція та основні етапи алгоритму

Розроблений алгоритм виявлення камуфляжу побудований на аналізі кольорових характеристик поверхонь у режимі реального часу, що фіксуються RGB-датчиком TCS34725. На відміну від багатьох сучасних систем, які базуються на складних моделях комп'ютерного зору або глибинного навчання [11], [13], [14], запропоноване рішення орієнтоване на роботу у простих апаратно-програмних умовах, забезпечуючи низьке енергоспоживання, компактність і можливість автономного функціонування. Завдяки цьому алгоритм може бути інтегрований у мобільні сенсорні вузли або безпілотні платформи, що особливо важливо для оперативного розгортання у польових умовах.

Основна ідея полягає у тому, що будь-який об'єкт, замаскований під навколишнє середовище, рідко забезпечує ідеальну відповідність спектральним характеристикам фону. Навіть мінімальні колірні відхилення можуть бути зафіксовані датчиком і проаналізовані алгоритмом. Процес виявлення відбувається поетапно: формування фонових профілів, нормалізація даних, обчислення вектора відхилення, застосування методів фільтрації та адаптивних порогів, а також перевірка стабільності сигналу.

Послідовність основних етапів алгоритму така:

Етап 1. Ініціалізація системи – калібрування RGB-датчика, задання порогових параметрів T та $T_{\min}$.

Етап 2. Формування фонових профілів – накопичення та усереднення вимірювань RGB-компонентів для побудови еталонного вектора F_0 .

Етап 3. Зчитування поточних даних $F_n = (R_n, G_n, B_n)$.

Етап 4. Нормалізація колірних компонентів R^* , G^* , B^* для компенсації освітленості.

Етап 5. Обчислення евклідової відстані ΔF^* між F_n^* та F_0^* .

Етап 6. Порівняння ΔF^* з порогом T та попереднє виявлення аномалії.

Етап 7. Медіанна фільтрація часової послідовності для усунення шумів.

Етап 8. Перевірка часової стабільності сигналу (аномалія зберігається щонайменше 3 кадри).

Етап 9. Просторова перевірка у локальному вікні 3×3 та підтвердження аномалії.

Етап 10. Виведення результату – позначення виявленої ділянки та передача даних.

Така послідовність етапів дозволяє реалізувати алгоритм на мікроконтролері ESP32 та забезпечує обробку даних у реальному часі навіть за обмежених апаратних ресурсів.

5.2. Формування фонового профілю

На початковому етапі роботи система здійснює автоматичне формування еталонної моделі фону, що виступає опорною базою для подальшого виявлення змін у сцені. Процедура формування профілю передбачає збір серії спектральних вимірювань у просторі RGB з виділеної ділянки спостереження за умов відсутності сторонніх об'єктів; тривалість накопичення вибірки залежить від динаміки середовища і може становити від кількох секунд до декількох хвилин. Отримані значення для кожного вимірювального елемента (пікселя або зони покриття) усереднюються та формують еталонний вектор фонового кольору:

– $F_0 = (R_0, G_0, B_0)$ – еталонне значення кольору (базовий фон);

– $F_n = (R_n, G_n, B_n)$ – поточне вимірювання.

де R_0, G_0, B_0 – середні значення інтенсивностей відповідних каналів за період калібрування. У разі подібної роботи із зонованими вимірами (наприклад, сітка зон покриття) для кожної зони зберігається власний вектор $F_0^{(i)}$.

Формування профілю здійснюється з урахування практичних обмежень: відбір вимірювань має відбуватися у відносно стабільних умовах освітлення, без рухомих об'єктів у полі зору, щоб уникнути артефактів у еталонній моделі.

Профіль може оновлюватися періодично (планове або адаптивне оновлення) для врахування сезонних, добових або погодних трансформацій середовища; при цьому використовується експоненціальне згладжування або ковзне середнє з тривалістю вікна, підбраною експериментально. Формування надійного фонового профілю є критичною складовою, оскільки від його якості напряду залежить чутливість і специфічність системи виявлення [12].

5.3. Формування фонового профілю

Для кількісної оцінки відмінності між поточним вимірюванням $F_n = (R_n, G_n, B_n)$ та еталонним фоновим вектором F_0 використовується відстань у тривимірному нормалізованому кольоровому просторі. Основною метрикою служить евклідова відстань:

$$F = \sqrt{(R_n - R_0)^2 + (G_n - G_0)^2 + (B_n - B_0)^2}. \quad (1)$$

Практично, для підвищення інваріантності до змін освітленості, обчислення виконується над нормалізованими компонентами. Якщо отримане значення ΔF перевищує заздалегідь встановлений поріг чутливості T , то відповідна точка або зона позначається як потенційно аномальна (підсумкова детекція залежить від додаткових часово-просторових критеріїв і фільтрації).

Порогове значення T може бути задане двома способами: фіксовано – після попереднього калібрування під конкретні експлуатаційні умови; адаптивно – з урахуванням статистичних характеристик фонового профілю (наприклад, через множення стандартного відхилення фонового сигналу на емпіричний коефіцієнт). Адаптивна постановка дозволяє зберігати прийнятний баланс між чутливістю та стійкістю у довільних ландшафтах або при зміні погодних умов [14].

5.4. Нормалізація значень

Оскільки абсолютні інтенсивності каналів R, G, B сильніше за все залежать від рівня освітлення, у системі реалізовано нормалізацію кольорових компонентів у відносній пропорції:

$$R^* = \frac{R}{R+G+B}; \quad G^* = \frac{G}{R+G+B}; \quad B^* = \frac{B}{R+G+B}, \quad (2)$$

де R^* , G^* , B^* – нормалізовані компоненти кольору. У нормалізованому просторі евклідова відстань (1) набуває вигляду:

$$F^* = \sqrt{(R_n^* - R_0^*)^2 + (G_n^* - G_0^*)^2 + (B_n^* - B_0^*)^2}, \quad (3)$$

що дозволяє оцінювати відхилення кольорів незалежно від загальної інтенсивності освітлення. Така нормалізація вилучає вплив глобальної яскравості та зосереджує аналіз на співвідношенні компонент, що є релевантнішим при відокремленні колірних тонів. Нормалізовані компоненти використовуються у подальших обчисленнях (зокрема у формулі (3) для ΔF^* , що істотно зменшує кількість помилкових детекцій, пов'язаних із загальними коливаннями освітлення (наприклад, перехід від сонячної погоди до часткового покриття хмарами) [13].

Під час нормалізації доцільно контролювати випадки близьких до нуля сум каналів ($R+G+B \approx 0$) і застосовувати регуляризацию або мінімальні пороги по яскравості для уникнення чисельної нестабільності.

5.5. Виявлення кольорових аномалій

У межах запропонованої системи виявлення кольорових аномалій ґрунтується на поєднанні спектрального критерію з часово-просторовою валідацією результатів. Первинна оцінка відхилення здійснюється за евклідовою відстанню ΔF^* у нормалізованому просторі (3). Якщо це відхилення перевищує порогове значення T , фіксується потенційна аномалія. Для умов польових випробувань було встановлено робоче значення $T=0,15$, однак у разі складних фонів або змін освітленості поріг може обчислюватися адаптивно.

Для того, щоб уникнути реакції на випадкові та короточасні збурення, такі як відблиски або рух листя, було введено часовий критерій стабільності. Аномалія вважається істотною лише тоді, коли відхилення зберігається протягом певної кількості вимірювань (наприклад, не менше трьох кадрів) або мінімального інтервалу часу T_{min} . Таким чином система ігнорує шумові коливання та реагує лише на сталі об'єкти.

Додатковим рівнем перевірки є просторове підтвердження. Для того, щоб уникнути помилкових спрацювань на одиничні пікселі, відхилення має бути виявлене відразу в кількох суміжних елементах, наприклад, у локальному вікні розміром 3×3 . Лише у випадку, якщо кількість сусідніх відхилених точок перевищує певний мінімум, система підтверджує наявність аномалії. Така комбінація спектрального, часового та просторового критеріїв дозволяє значно знизити кількість хибнопозитивних детекцій і забезпечує стабільну роботу навіть у змінних природних умовах [11]-[14].

Для підвищення ефективності аналізу можуть також застосовуватись інші колірні простори, зокрема HSV або CIELAB. Вони краще відображають перцептивні особливості зору людини та дозволяють чіткіше виокремлювати відмінності у відтінках при схожих рівнях яскравості [12].

5.6. Математичні та алгоритмічні засоби

У перспективі система може бути вдосконалена за рахунок впровадження методів машинного навчання. Зокрема, можливо використання алгоритмів k -найближчих сусідів (k -NN) або методу опорних векторів (SVM — один із базових алгоритмів машинного навчання для розділення класів) для точнішої класифікації фонових і нефонових пікселів. Такі моделі можуть навчатися на підготовлених вибірках і краще адаптуватися до нових або складних сцен [18].

Алгоритмічна частина системи побудована так, щоб поєднати простоту реалізації з достатньою стійкістю до шумів. На першому кроці застосовується ковзне середнє для згладжування повільних змін у часових рядах, а також медіанна фільтрація у вікні з п'яти

вимірювань, яка ефективно усуває імпульсні викиди і водночас зберігає контурні характеристики [15], [16]. У результаті формується стабільніший потік даних для подальшої обробки.

Наступним кроком у процесі аналізу відхилень є визначення порогів відхилення. У запропонованій системі використано адаптивний підхід, коли значення T змінюється залежно від дисперсії фонових профілів. Це дозволяє автоматично підлаштовувати чутливість системи до різних умов сцени: підвищувати поріг у шумних ділянках і знижувати його в однорідних. Класифікація пікселів (або зон спостереження) як фонових здійснюється шляхом обчислення ΔF^* та перевірки його відносно порогового значення. Такий метод є обчислювально простим, тому може виконуватися у реальному часі навіть на малопотужних мікроконтролерах [17].

У перспективі передбачається розширення алгоритмічної частини за рахунок впровадження методів машинного навчання. Зокрема, можливим є використання класифікаторів k -NN або SVM, які здатні навчатися на експериментальних даних та покращувати якість відокремлення фону від об'єктів. Проте у поточному прототипі ці методи не застосовувалися, оскільки головним завданням було створення максимально простої та енергоефективної системи [18].

5.7. Особливості роботи з камуфляжем

Камуфляжні засоби спрямовані на наближення кольорових і текстурних характеристик об'єкта до навколишнього середовища; втім, ідеальної відповідності досягти неможливо, тому при коректній обробці сигналу такі зразки можна виявити. У запропонованій системі аналіз здійснюється в нормалізованому кольоровому просторі RGB – тобто на основі відстані ΔF^* між поточним та фоновим векторами – та доповнюється просторовим і текстурним аналізом.

Перший тип ознак, на якій орієнтується алгоритм, – локальні контурні розбіжності: у місцях переходу «об'єкт–фон» спостерігаються різкі градієнти кольору або яскравості, які проявляються як локальні підвищення ΔF^* і корелюють з високою величиною градієнта інтенсивності.

Другий тип ознак – текстурні відмінності: штучні візерунки зазвичай мають інші локальні статистики (наприклад, інші значення другого або третього моменту, іншу енергію у високочастотній частині спектра), що дозволяє відокремити повторювані штучні патерни від природних текстур.

Третій тип – спектральні «чужорідні» вкраплення, які виникають при використанні синтетичних барвників або матеріалів; такі вкраплення дають локальні аномалії у нормалізованих відтінках і можуть бути виявлені як стійкі точки або малі області з підвищеним ΔF^* .

Для посилення надійності детекції колориметричний аналіз комбінується з просторовими фільтрами та просторово-частотними статистиками; у складних випадках доцільне залучення додаткових каналів (наприклад, інфрачервоного або мультиспектрального діапазону), що істотно підвищує ймовірність виявлення. Практично, параметри алгоритму (поріг T , мінімальна площа кластера, локальний поріг на текстурну енергію) підбираються експериментально для кожного типу місцевості; також слід враховувати обмеження апаратури (чутливість датчика, дальність спостереження, робочі умови освітлення) і включати в протокол калібрування відповідні тестові сценарії.

6. Результати натурних випробувань

6.1. Загальні особливості натурних випробувань

Після створення дослідного зразка системи було проведено серію натурних випробувань у різних природних умовах з метою перевірки точності виявлення

камуфльованих об'єктів і стійкості алгоритму до зовнішніх чинників. Експерименти здійснювалися у трьох типах місцевості: листяний ліс, відкрита степова зона та кам'яниста місцевість. Для кожного середовища моделювалися умови наявності камуфляжних елементів, максимально наближених за кольоровими характеристиками до фону.

Випробування проводилися згідно з алгоритмом, описаним у розділі 5: нормалізація RGB-компонентів у відсотковій формі, обчислення вектора кольорового відхилення ΔF^* , використання порога чутливості $T=0,15$ у нормованій шкалі, медіанна фільтрація у часовому вікні з п'яти кадрів і перевірка часової стабільності відхилення протягом мінімум трьох послідовних кадрів.

Для забезпечення відтворюваності параметри середовища були стандартизовані. Спостереження проводилися лише у денний час (10:00–16:00), коли природне освітлення характеризується найвищою стабільністю. Умови освітлення обмежувалися відсутністю хмарності, опадів або туману, щоб виключити додаткові флуктуації фону. Дистанція між датчиком і цільовим об'єктом змінювалася в межах 1–3 м, що відповідає ближнім умовам розвідки.

Кожна серія експериментів включала 100 тестових циклів спостереження. У кожному циклі аналізувалося 300 кадрів із роздільною здатністю 640×480 пікселів. Камуфльовані об'єкти займали у середньому 10–20 % площі кадру. У системі реєструвалися три типи результатів: істинно-позитивне виявлення (об'єкт присутній і система його правильно зафіксувала), хибнопозитивне спрацювання (система сигналізувала за відсутності об'єкта, False Positive, FP) та пропуск (об'єкт присутній, але система його не виявила).

В експерименті використовувалися такі типи камуфляжу:

- «мультикам» (MultiCam) – універсальний багатоспектральний камуфляжний візерунок, адаптований для ефективного маскування в широкому діапазоні ландшафтів, зокрема у пустельних, лісових та урбанізованих середовищах;
- «піксель» – маскувальний шаблон із виразною піксельною структурою, створений на основі принципів цифрової дискретизації зображення, що сприяє розмиттю контурів об'єкта в полі зору спостерігача;
- однотонний оливковий камуфляж – монохромний варіант маскування, призначений переважно для експлуатації в середовищі з щільною лісовою рослинністю, де переважають зелені тони спектру.

Зібрані дані фіксувалися для подальшого статистичного аналізу. Точність виявлення обчислювалася як відсоткове співвідношення правильно зафіксованих об'єктів до загальної кількості спроб, згідно з методикою, наведеною у підрозділі 6.2.

6.2. Методика випробувань

Методику випробувань було побудовано на основі алгоритмічної схеми, описаної в розділі 5.

Кожен експериментальний сценарій складався зі 100 кадрів, усього для кожного типу камуфляжу проведено 5 серій (загалом 500 спостережень).

Обробка даних проводилася у такій послідовності.

Крок 1. Нормалізація RGB-значень у відсотковій формі;

Крок 2. Обчислення відхилення ΔF^* від фонового профілю;

Крок 3. Використання порога чутливості $T=0,15$ у нормалізованій шкалі RGB з можливістю адаптації $\pm 0,05$ залежно від освітлення;

Крок 4. Медіанна фільтрація у часовому вікні з 5 кадрів для приглушення шумів;

Крок 5. Умова часової стабільності: відхилення мало зберігатися щонайменше протягом $T_{min}=3$ послідовних кадрів.

Точність було визначено за класичними метриками [12]–[15]:

$$Accuracy = \frac{TP}{TP + FN + FP} \times 100\%, \quad (4)$$

де TP – кількість істинно-позитивних спрацювань (зона аномалії збігалася з фактичним розташуванням об'єкта, підтвердженом візуально); FN – кількість пропусків; FP – хибнопозитивні спрацювання.

Наприклад, у випадку з цифровим камуфляжем «піксель» у лісі із 100 тестів 86 були правильно ідентифіковані системою, 11 – пропущені, а 3 – хибнопозитивні. Наведені за результатами дослідження значення (86 %, 78 %, 65 %) відповідають середнім результатам, отриманим за формулою (4) для сукупності всіх серій. Таким чином, точність становила:

$$Accuracy = \frac{86}{86 + 11 + 3} \times 100\% \approx 86\%. \quad (5)$$

Результати експериментальних досліджень зведено у таблиці 1.

Таблиця 1

Результати експериментальних досліджень з виявлення камуфляжу

Тип камуфляжу	Фон	Виявлення, %	Хибні спрацювання, %
«Піксель»	листяний ліс	86 %	7 %
«Мультикам»	відкрита степова зона	78 %	11 %
Однотонний	кам'яниста зона	65 %	13 %

Найвищу ефективність система показала при виявленні камуфляжу типу «піксель» у листяному лісі. Це пояснюється кращим контрастом дрібномасштабної структури пікселя зі статистичним фоном, що фіксується RGB-датчиком. Тип «мультикам», адаптований під широкий спектр природних умов, виявився складнішим для розпізнавання у відкритій степовій зоні через близькість колірному профілю до середовища. Найгірші результати спостерігалися у випадку однотонного камуфляжу на кам'янистому фоні. Відсутність різких контрастів і варіативна колірна структура фону знижували ефективність алгоритму виявлення, заснованого лише на RGB-аналізі.

Порівняльний аналіз технічних характеристик трьох типів систем виявлення об'єктів (табл. 2) демонструє, що запропонована система суттєво переважає за показниками вартості, енергоефективності та мобільності, хоча поступається альтернативним рішенням у можливості нічного функціонування.

Як видно з таблиці 2, запропонована система виявлення кольорових аномалій на основі RGB-датчика демонструє низку суттєвих переваг у порівнянні з традиційними технологіями, такими як тепловізори та інфрачервоні камери з елементами штучного інтелекту. Найочевиднішою перевагою є надзвичайно низька вартість (< \$ 20), що робить можливим масове впровадження такої системи у військових або охоронних застосуваннях із обмеженим бюджетом.

Таблиця 2

Порівняльний аналіз технічних характеристик систем виявлення об'єктів

Параметр	Запропонована система	Тепловізор	Інфрачервона камера + AI
Вартість	< \$ 20	> \$ 1000	> \$ 300
Споживання енергії	дуже низьке	високе	середнє
Робота вночі	ні	так	так
Мобільність	висока	середня	середня
Залежність від погодних умов	помірна	висока	помірна

Важливою є також мінімальна енергозалежність, що дає змогу використовувати систему у мобільних або автономних застосуваннях, зокрема у безпілотних платформах або сенсорних сітках. Хоча запропонована система не функціонує у нічний час за відсутності зовнішніх джерел освітлення, її ефективність у денний час при стабільних погодних умовах є достатньою для виконання завдань виявлення камуфляжних об'єктів у ближній зоні.

За мобільністю запропонована система значно перевершує більшість аналогів завдяки мініатюрним габаритам та низькому енергоспоживанню, що розширює спектр можливих сценаріїв застосування. Хоча залежність від погодних умов залишається помірною, вона є прийнятною в контексті задач денного моніторингу при хорошій видимості.

7. Обговорення результатів дослідження

Експериментальні дослідження підтвердили, що ефективність запропонованої RGB-системи істотно залежить від типу камуфляжу та фонового середовища. Найвищий показник точності (86 %) отримано для цифрового камуфляжу в умовах листяного лісу; нижчий результат (78 %) – для «мультикама» у степовій зоні; мінімальне значення (65 %) зафіксовано для однотонного камуфляжу на кам'янистій місцевості. Наведені значення точності відповідають усередненим результатам серій експериментів, описаних у підрозділі 6.3. Середній рівень FP становив 3-7 %.

Зміни в освітленні (перехід від прямого сонячного до частково хмарного) та поява локальних тіней збільшували кількість FP, що свідчить про потребу в удосконаленні адаптивних механізмів компенсації освітленості. Порівняння з тепловізорами та інфрачервоними камерами показало: запропонована система значно дешевша (вартість сенсорного вузла становить < 5 % від вартості інфрачервоної камери), споживає у 8-10 разів менше енергії й має меншу масу, але обмежена роботою у видимому діапазоні [20]-[22].

Попри задовільні показники точності, система має низку обмежень. Вона не функціонує у нічний час, залежить від рівня природного освітлення та чутливості сенсора до змін температури. Крім того, дальність дії обмежена кількома метрами, що зумовлено оптичними характеристиками датчика TCS34725.

Отримані результати підтверджують, що RGB-системи доцільно застосовувати як складову багатоспектральних платформ у завданнях ближнього виявлення при обмежених ресурсах.

8. Висновки

У ході роботи розроблено прототип апаратно-програмної системи виявлення камуфляжу на основі RGB-датчика. Структура системи побудована з використанням доступних компонентів – сенсора TCS34725 із вбудованим інфрачервоним фільтром і мікроконтролера ESP32. Така конфігурація забезпечила низьке енергоспоживання,

компактність та можливість автономної роботи в польових умовах.

Розроблено алгоритм виявлення кольорових аномалій у просторі RGB, який включає формування фонового профілю, нормалізацію колірних компонентів, обчислення вектора відхилення ΔF^* , медіанну фільтрацію та адаптивне порогове визначення. Застосування адаптивного підходу дало змогу підвищити стійкість системи до змін освітлення та зменшити кількість FP.

Проведено натурні випробування у трьох типах фонових середовищ — лісовому, степовому та кам'янистому. Результати показали, що точність виявлення камуфляжу становить: 86 % для цифрового «пиксельного» камуфляжу, 78 % для «мультикаму» та 65 % для однотонного оливкового. Отримані дані підтверджують працездатність запропонованої системи у денних умовах і при стабільному освітленні.

Порівняння з тепловізійними та інфрачервоними системами показало, що запропоноване RGB-рішення має істотні переваги за вартістю, енергоефективністю та мобільністю, однак обмежене роботою у видимому діапазоні. Система не функціонує у нічний час та є чутливою до різких змін освітлення, що визначає напрями подальшої оптимізації.

У подальшому передбачається інтеграція інфрачервоних каналів і перевірка системи на більших відстанях та за змінних умов освітлення.

Перелік посилань:

1. Sun Y., Chen G., Zhou T., Zhang Y., Liu N. Context-aware Cross-level Fusion Network for Camouflaged Object Detection (C2F Net). *Proceedings of the 30th International Joint Conference on Artificial Intelligence (IJCAI-21)*. 2021. С. 1025–1031. URL: <https://doi.org/10.24963/ijcai.2021/142>.
2. Jiang X., Cai W., Zhang Z., Jiang B., Yang Z., Wang X. MAGNet: Camouflaged Object Detection Network Simulating the Observation Effect of a Magnifier. *Entropy*. 2022. Vol. 24, № 12. С. 1804. URL: <https://doi.org/10.3390/e24121804>.
3. Sun Y., Wang S., Chen C., Xiang T.-Z. Boundary-Guided Network (BGNet) for Camouflaged Object Detection. *arXiv preprint*. URL: <https://arxiv.org/abs/2207.00794> (дата звернення: 15.07.2025).
4. Ji G.-P., Fan D.-P., Chou Y.-C. et al. Deep Gradient Learning for Efficient Camouflaged Object Detection (DGNet). *Machine Intelligence Research*. 2023. Vol. 20. p. 92–108. URL: <https://doi.org/10.1007/s11633-022-1365-9>.
5. Zheng D., Zheng X., Yang L.T. et al. MFFN: Multiview Feature Fusion Network for Camouflaged Object Detection. *arXiv preprint*. URL: <https://arxiv.org/abs/2210.06361> (дата звернення: 15.07.2025).
6. Guo Y., Huang H. CLAD: Contrastive Learning and Data Augmentation for Camouflaged Object Detection. *Sensors*. 2024. Vol. 24, № 6. p. 2581. URL: <https://doi.org/10.3390/s24062581>.
7. Popa A., Botezatu N., Moraru L. Development of a Colorimetric Sensor for Autonomous, Networked, Real-Time Application. *Sensors*. 2020. Vol. 20, № 20. p. 5857. URL: <https://doi.org/10.3390/s20205857>.
8. Hakim G.P.N., Darmawan A., Fadillah M. Benchmarking In Microcontroller Development Board Power Consumption For Low Power IoT WSN Application. *Jurnal Teknologi Elektro dan Komputer*. 2022. Vol. 11, № 2. p. 105–112. URL: <https://doi.org/10.31294/jtekin.v11i2.12546>.
9. Paredes M., Gamarra A., Alvarado J. Prototype of an IoT-Based Low-Cost Sensor Network for Hydrological Monitoring in Remote Areas. *Sensors*. 2023. Vol. 23, № 3. p. 1340. URL: <https://doi.org/10.3390/s23031340>.
10. Hermawan R.A., Purwanto A. Smart Home Monitoring System Using ESP32 Microcontrollers. *IntechOpen*. URL: <https://www.intechopen.com/chapters/77491> (дата звернення: 15.07.2025).
11. Yan L.Y., Xu X.Q., Zhang M., Luo S.X. Enhanced Spatio-Temporal Attention Mechanism for Video Anomaly Event Detection. *Advances in Computer Engineering*. 2025. Vol. 117. URL: <https://doi.org/10.54254/2755-2721/2025.20838>.
12. Li Y., Liu H., Wang Y. Evaluation of Color Anomaly Detection in Multispectral Images Using Statistical Thresholding. *Engineering*. 2023. Vol. 3, № 4. p. 541–553. URL: <https://doi.org/10.3390/eng3040038>.
13. Zhao Y., Li C., Zhang J. Improved Spatio-Temporal Graph Convolutional Networks for Fine-Grained Video Anomaly Detection. *Opto-Electronic Engineering*. 2024. Vol. 51, № 2. Article ID: 240034. URL: <https://doi.org/10.12086/oee.2024.240034>.
14. Tran M., Nguyen T. Deep BiLSTM Attention Model for Spatial and Temporal Anomaly Detection in Surveillance Videos. *Sensors*. 2025. Vol. 25, № 1. Article № 251. URL: <https://doi.org/10.3390/s25010251>.

15. Romero-González G., Santana-Cedrés D., Díaz-Díaz N., Lorenzo-Navarro J. Background Subtraction for Dynamic Scenes Using Gabor Filter Bank and Statistical Moments. *Algorithms*. 2024. Vol. 17, № 4. Article № 133. URL: <https://doi.org/10.3390/a17040133>.
16. Pambudi S., Wahyudi I., Astuti N.I. Impact of Wolf Thresholding on Background Subtraction for Human Motion Detection. *Compiler*. 2024. Vol. 13, № 1. p. 37–45. URL: <https://doi.org/10.28989/compiler.v13i1.2116>.
17. Minaee S., Boykov Y., Porikli F., Plaza A., Kehtarnavaz N., Terzopoulos D. Image Segmentation Using Deep Learning: A Survey. *IEEE Transactions on Pattern Analysis and Machine Intelligence*. 2022. Vol. 44, № 7. p. 3523–3542. URL: <https://doi.org/10.1109/TPAMI.2021.3059968>.
18. Li X., Wang Y., Yang J., Zhou Y., Li Z. Transformer-Based Visual Segmentation: A Survey. *arXiv preprint*. 2023. arXiv:2304.09854. URL: <https://doi.org/10.48550/arXiv.2304.09854>.
19. Хорошайло Ю. Є., Семенов С. Г., Лимаренко В. В. Цифровий датчик для вимірювання кольору : пат. на корисну модель 107317 Україна; ХНУРЕ. 2016.
20. BOSON® LWIR OEM Thermal Camera Module Datasheet. *Teledyne FLIR*. URL: <https://f.hubspotusercontent10.net/hubfs/20335613/flir-boson-datasheet.pdf> (дата звернення: 15.07.2025).
21. FLIR Lepton® Series Datasheet (Lepton 3 / 3.5). *Teledyne FLIR*. URL: https://d3mm496e6885mw.cloudfront.net/manufacture_product/652edc2b527692120f62e958/specsheet/spec_sheets/original/LeptonSeriesDatasheet1.pdf (дата звернення: 15.07.2025).
22. MLX90640 32×24 IR Array Datasheet. *Melexis*. URL: <https://www.melexis.com/-/media/files/documents/datasheets/mlx90640-datasheet-melexis.pdf> (дата звернення: 15.07.2025).

Надійшла до редколегії 10.09.2025 р.

Крицький Дмитро Миколайович, кандидат технічних наук, доцент, декан факультету літакобудування, Національний аерокосмічний університет «Харківський авіаційний інститут», м. Харків, Україна, e-mail: d.krickiy@khai.edu, ORCID: <https://orcid.org/0000-0003-4919-0194>.

Оніщук Денис Анатолійович, аспірант кафедри ІВТ ХНУРЕ, м. Харків, Україна, e-mail: denys.onishchuk@nure.ua, ORCID: <https://orcid.org/0009-0003-3080-7710>.

Валюженич Олександр Олексійович, аспірант кафедри ІВТ, ХНУРЕ, Харків, Україна, e-mail: oleksandr.valiuzhenych@nure.ua, ORCID: <https://orcid.org/0009-0000-3695-1658>.

УДК 004.8:004.9

DOI: 10.30837/0135-1710.2025.186.095

*С.Ф. ЧАЛИЙ, І.О. ЛЕЩИНСЬКА***МЕТОД ОЦІНКИ НЕГАТИВНИХ АСПЕКТІВ РІШЕННЯ ІНТЕЛЕКТУАЛЬНОЇ СИСТЕМИ В ЗАДАЧАХ ПОБУДОВИ ПОЯСНЕНЬ**

Об'єктом дослідження є процес оцінки рішення в інтелектуальній системі. Метою є розробка методу оцінки негативних аспектів рішень інтелектуальних систем з тим, щоб забезпечити формування збалансованих пояснень на основі інтеграції суттєвих негативних характеристик у ментальні моделі користувачів. Розроблено гібридний підхід до побудови ментальної моделі, який поєднує аналіз важливості та контекстуальний аналіз релевантності рішень з метою комплексної оцінки негативних властивостей вибраного користувачем рішення. Запропоновано метод оцінки негативних аспектів користувацького рішення при побудові ментальної моделі користувача, який інтегрує SHAP-аналіз важливості негативних властивостей з контекстним аналізом їхньої релевантності, що дає можливість автоматизованого виявлення та упорядкування негативних характеристик рішення для формування зрозумілих пояснень та підвищення довіри користувачів.

1. Вступ

Сучасні інтелектуальні системи (ІС) формують ефективні рішення для користувачів, використовуючи методи машинного навчання. Такі системи аналізують великі обсяги даних при формуванні персоналізованих рекомендацій у широкому спектрі предметних областей – від сфери електронної комерції й до області медичної діагностики. Проте наслідком використання технологій машинного навчання є виникнення проблеми непрозорості алгоритму прийняття рішень, коли користувачі ІС не мають інформації щодо логіки прийняття рішень інтелектуальною системою [1], [2].

Незрозумілість процесу формування рішення знижує довіру користувачів до ІС та обмежує їхню готовність використовувати рішення такої системи [3]. Це приводить до виникнення потреби у розробці методів пояснювального (зрозумілого) штучного інтелекту з тим, щоб забезпечити зрозумілість алгоритмічних рішень [4]–[6].

Одним із ключових підходів до вирішення проблеми непрозорості ІС у науковому напрямку пояснювального штучного інтелекту є побудова ментальних моделей користувачів [7]. Такі моделі відображають внутрішнє уявлення користувача ІС про процеси функціонування системи та її логіку прийняття рішень [8]. Однак традиційні підходи до побудови ментальних моделей орієнтовані переважно на представлення позитивних аспектів рішень ІС, не приділяючи достатньо уваги її негативним характеристикам [9]. Зокрема, підходи до інтерпретації важливості ознак вхідних даних не враховують аспект негативного сприйняття рішення користувачами [10]–[12]. Проте аналіз негативних аспектів рішень ІС має суттєве значення для формування збалансованого представлення про можливості та обмеження процесу формування рішення в такій системі. Користувачі ІС мають оцінити інформацію як про переваги запропонованого рішення, так і про його потенційні недоліки. Такий комплексний підхід дає можливість користувачеві підвищити довіру до інтелектуальної системи [13].

Таким чином, актуальною є проблема оцінки негативних аспектів рішень в ІС з тим, щоб забезпечити формування збалансованих пояснень в інтелектуальній системі на основі включення важливих негативних характеристик рішення до ментальної моделі користувача.

2. Аналіз сучасних наукових публікацій і постановка проблеми дослідження

Пояснювальний штучний інтелект орієнтований на розробку методів підтримки інтерпретації рішень ІС, що були отримані з використанням машинного навчання. Даний

напрямок започатковано програмою XAI DARPA [3]. Сучасні дослідження з пояснювального штучного інтелекту присвячені розробці методів як локальної, так і загальної інтерпретації рішень ІС [10]-[12], [14].

Представлений в [12] метод LIME призначений для локальної інтерпретації рішень ІС шляхом побудови спрощених моделей пояснень для конкретних прикладів рішень. Даний підхід виділяє ключові ознаки, що впливають на конкретне рішення інтелектуальної системи. Однак даний метод не враховує обмеження щодо глобальної узгодженості пояснень, а також не враховує контекст сприйняття рішень ІС користувачами [10].

Метод побудови пояснень SHAP (SHapley Additive exPlanations) використовує теорію кооперативних ігор і, на цій основі, забезпечує обґрунтований розподіл важливості ознак вхідних даних [10], [11]. Проте традиційний метод SHAP не враховує обмеження, пов'язані із негативним сприйняттям користувачами окремих аспектів рішення ІС.

Сучасні дослідження у сфері пояснювального штучного інтелекту орієнтовані на розробку методів, що враховують психологічні аспекти сприйняття пояснень користувачами. Останні представляються ментальними моделями [13], [14]. Ряд досліджень враховує також важливість адаптації пояснень до персональних особливостей користувачів ІС та контексту використання рішень [5], [15]-[17]. Класичні дослідження показали, що ментальні моделі відіграють ключову роль у взаємодії людини з технічними системами [18].

У контексті інтелектуальних систем ментальні моделі допомагають користувачам прогнозувати поведінку системи та приймати обґрунтовані рішення [8], [19]. Концептуальне представлення ментальної моделі щодо пояснення в ІС забезпечує структуровану репрезентацію знань про логіку прийняття рішень [8]. Ієрархічне, каузальне та функціонально-темпоральне представлення ментальних моделей в комплексі забезпечують можливість врахувати динаміку взаємодії користувача з ІС, наслідки альтернативних рішень [14], [20]. Принцип уточнення ментальних моделей на основі доповнення вхідних даних з урахуванням негативних характеристик рішення забезпечує підвищення релевантності пояснень [13], оскільки дозволяє врахувати позитивні, негативні та нейтральні аспекти, які впливають на вибір користувача [21], [22]. Виявлення негативних аспектів рішення дає можливість зрозуміти контекст незадоволення користувачів у сенсі обмежень щодо його використання і адаптувати пояснення [23].

Таким чином, існуючі підходи до побудови ментальних моделей та пояснень в ІС орієнтовані переважно на поясненні логіки рішення системи, не інтегруючи у пояснення важливих негативних характеристик цього рішення. Передумовою такої інтеграції є розробка методів виявлення та оцінки негативних характеристик рішень ІС на основі частково структурованих даних користувачів.

Зазначені обмеження існуючих методів свідчать про важливість розробки гібридного підходу, який би поєднував можливості існуючих методів ХАІ з аналізом негативних аспектів рішення для формування ментальних моделей користувачів ІС.

3. Мета і задачі дослідження

Метою дослідження є розробка методу оцінки негативних аспектів рішень інтелектуальних систем з тим, щоб забезпечити формування збалансованих пояснень на основі інтеграції суттєвих негативних характеристик у ментальні моделі користувачів.

Для досягнення поставленої мети вирішуються такі задачі: розробити гібридний підхід до побудови ментальної моделі, який поєднує аналіз важливості та контекстуальний аналіз релевантності рішень з метою комплексної оцінки негативних властивостей вибраного користувачем рішення; розробити метод оцінки негативних аспектів користувацького рішення при побудові ментальної моделі користувача для автоматизованого виявлення та упорядкування негативних характеристик цього рішення.

4. Гібридний підхід до побудови ментальної моделі з урахуванням негативних оцінок рішень інтелектуальної системи

Об'єктом дослідження є процес оцінки рішення в інтелектуальній системі. Предметом дослідження є принципи та методи оцінки негативних аспектів рішень інтелектуальних систем при формуванні користувацьких ментальних моделей для створення пояснень.

Запропонований підхід інтегрує результати аналізу важливості впливу властивостей вхідних даних з результатами контекстного аналізу релевантності цих властивостей для формування збалансованих ментальних моделей користувачів. Підхід ґрунтується на запропонованому в [20] принципі доповнення вхідних даних, згідно з яким доповнена j -та ментальна модель M_j представляється у вигляді множини позитивних (релевантних) властивостей рішення V_j^+ та множини негативних (нерелевантних) властивостей V_j^- , які обмежують використання рішення: $M_j = V_j^+ \cup V_j^-$.

Збалансована ментальна модель формує комплексне представлення користувача про рішення інтелектуальної системи, враховуючи як переваги, так і потенційні обмеження останнього. Таке рішення зменшує невизначеність при прийнятті рішень користувачем та, відповідно, створює умови для підвищення довіри до ІС.

Підхід реалізує послідовну оцінку негативних властивостей рішень ІС. На першому рівні виконується статистичний аналіз важливості характеристик даних для виявлення значущості останніх на основі їхнього внеску у формування загального сприйняття рішення. На другому рівні виконується контекстний аналіз релевантності рішення в сенсі оцінки суб'єктивної важливості для користувача негативних аспектів цього рішення для конкретних умов його використання. Процес формування ментальної моделі з урахуванням негативних аспектів рішення включає три основні етапи, представлені в табл. 1. Як вхідні дані при реалізації даного підходу використовується набір неструктурованих відгуків користувачів у текстовому форматі, дані щодо властивостей рішення ІС (предметна область, категорія, параметри, умови використання тощо). В рамках підходу як проміжні дані використовуються структуровані набори позитивних та негативних властивостей рішення ІС з кількісними оцінками частоти їх згадування у відгуках. На першому етапі розробленого підходу виконується підготовка даних для аналізу. На другому етапі формується комплексна оцінка важливості негативних аспектів на основі аналізу їхніх значущості та релевантності. На третьому етапі формується кінцева ментальна модель, що містить найсуттєвіші негативні характеристики рішення.

Представлений підхід забезпечує оцінку збалансованості ментальної моделі на основі співвідношення позитивних та негативних аспектів рішення ІС.

5. Метод оцінки негативних аспектів рішення інтелектуальної системи

Розроблений метод базується на гібридному підході до побудови ментальної моделі.

Метод враховує негативні оцінки рішень інтелектуальної системи та включає такі етапи.

Етап 1. Попередня обробки неструктурованих вхідних даних. На даному етапі здійснюється стандартизація текстових даних, видалення шуму та виконання настроєного аналізу для виявлення емоційного забарвлення цих даних. Результатом етапу є множина структурованих даних, кожен елемент якої містить компоненти вхідного тексту та його настроєно-оцінку.

Етап 2. Класифікація властивостей рішення. На даному етапі виділяються підмножини позитивних V_j^+ та негативних V_j^- властивостей рішення за результатами настроєного аналізу на попередньому етапі.

Таблиця 1

Основні етапи гібридного підходу

Етапи	Задача етапу	Реалізація етапу	Результат етапу
1. Структуризація вхідних даних	Перетворення неструктурованих відгуків користувачів у табличний формат	Сентимент-аналіз, класифікація негативних/позитивних аспектів	Множини структурованих властивостей вхідних даних P^+ , P^-
2. Гібридна оцінка негативних властивостей рішення	Оцінка значущості та релевантності негативних аспектів рішення	– SHAP-аналіз значущості негативних аспектів рішення; – контекстний аналіз релевантності негативних аспектів рішення	Ранжовані негативні аспекти з інтегрованими оцінками
3. Формування моделі	Інтеграція важливих негативних аспектів у ментальну модель	Балансування ментальної моделі, відбір і валідація важливих негативних характеристик рішення	Доповнена ментальна модель

Етап 3. SHAP-аналіз важливості негативних властивостей. На даному етапі на основі теорії кооперативних ігор формується оцінка внеску кожної негативної характеристики у формування загального сприйняття рішення користувачами. Використовується метод SHAP. Оцінка кожної із властивостей визначається через обчислення її SHAP-значення [10] та подальшу нормалізацію цього значення. Результатом етапу є сума S_j^-

SHAP-значень $s_{i,j}^-$ всіх негативних властивостей: $S_j^- = \sum_i s_{i,j}^-$.

Етап 4. Контекстний аналіз негативних властивостей. На даному етапі враховується частота згадування $f_{i,j}^-$ кожної негативної властивості у відгуках користувачів, оцінка впливу властивості на можливість використання рішення $p_{i,j}^-$, що отримана за результатами сентимент-аналізу, а також оцінка впливу на загальний користувацький досвід $u_{i,j}^-$. При оцінці $p_{i,j}^-$ враховується вплив на функціональність (чи буде працювати рішення без відповідної функції), а також частота згадування про проблему у вхідних даних. При оцінці $u_{i,j}^-$ враховується частота виникнення проблеми, а також її вплив на досвід використання рішення. Оцінка $w_{i,j}^-$ для кожної властивості нормалізується зі з використанням змінної $n_{i,j}^- = \max_i (f_{i,j}^- p_{i,j}^- u_{i,j}^-)$:

$$w_{i,j}^- = \frac{f_{i,j}^- p_{i,j}^- u_{i,j}^-}{n_{i,j}^-}. \quad (1)$$

Результатом етапу є сумарна контекстна оцінка $W_j^- = \sum_j w_{i,j}^-$.

Етап 5. Формування комбінованої оцінки рішення ІС. На даному етапі визначається співвідношення між SHAP-оцінкою S_j^- та контекстною оцінкою W_j^- . Дане співвідношення порівнюється з пороговим значенням E :

$$\frac{S_j^-}{W_j^-} > E. \quad (2)$$

Такий підхід дає можливість оцінити вплив найважливіших негативних властивостей рішення ІС з урахуванням контексту їх використання. Тобто дане співвідношення дає можливість оцінити баланс представлення рішення з точки зору користувача.

Етап 6. Відбір ключових негативних аспектів рішення. На даному етапі формується підмножина рішень V_j^{--} , для яких співвідношення $\frac{S_{i,j}^-}{W_{i,j}^-}$ менше, ніж значення (2), тобто відбираються найсуттєвіші негативні властивості:

$$V_j^{--} = \left\{ v_{i,j}^- : \frac{S_j^-}{W_j^-} > \frac{S_{i,j}^-}{W_{i,j}^-} \right\}. \quad (3)$$

Такий підхід дає можливість сфокусувати увагу лише на ключових обмеженнях рішення ІС та уникнути інформаційного перевантаження користувача і тим самим підвищити довіру до такого рішення.

Етап 7. Формування доповненої ментальної моделі рішення. На даному етапі ментальна модель j -го рішення представляється через множини позитивних та ключових негативних властивостей, що створює збалансоване представлення користувача про рішення ІС:

$$M_j = V_j^+ \cup V_j^{--}. \quad (4)$$

Отримана ментальна модель відображає комплексне представлення рішення, що включає переваги та суттєві обмеження даного рішення. Відповідно, така модель створює умови щодо обґрунтованого вибору рішення користувачем.

6. Експериментальна перевірка розробленого методу

Експериментальну перевірку методу виконано з використанням користувацьких відгуків 2025 року про мобільний телефон на платформі систем електронної комерції. Експериментальна вибірка складалася із 65 відгуків. З них 45 відгуків містили інформацію про негативні аспекти даного телефону. Файл з відгуками містив ID користувача, дату публікації, оцінку телефону за п'ятибальною шкалою, текст відгуку, переваги та недоліки продукту, виділені користувачем властивості продукту.

Результати етапу 1 методу було представлено показниками негативних властивостей рішення. У відгуках було виділено 15 типів негативних аспектів, зокрема швидкість роботи, автономність батареї, якість камери, нагрів пристрою. Кожна негативна властивість отримала числовий показник відповідно до частоти згадування у множині відгуків та оцінки його важливості в контексті досвіду користувача. Діапазон оцінок властивостей визначався після нормалізації цих показників. При проведенні експерименту він становив 0,227–0,06. Показник з найвищим значенням характеризує властивість з найбільшою частотою негативних відгуків, що має найсуттєвіший вплив на загальну оцінку користувача. Показник з мінімальним значенням відповідає негативним властивостям, що мають найменше відгуків і, отже, мають незначний вплив на формування враження користувачів від продукту. На третьому етапі було отримано значення $S_j^- = 0,9$ сумарного показника негативних властивостей. Таке високе SHAP-значення свідчить про суттєвий внесок негативних властивостей у загальну оцінку. На етапі 4 отримано значення

контекстної оцінки $W_j^- = 0,346$. Дане значення враховувало ступінь впливу на користувацький досвід, частоту виникнення проблеми та рівень впливу на рішення про придбання цієї моделі. Зокрема, оцінка швидкості роботи системи становила 0,12 внаслідок її суттєвого впливу на повсякденне використання телефону. Оцінка нагріву телефону під навантаженням становила 0,05 внаслідок незначного впливу на повсякденне використання пристрою. Найсуттєвішими властивостями рішення, що мали найбільший негативний вплив, були швидкість роботи системи, час автономної роботи батареї, якість зйомки камерою при слабкому освітленні, нагрів телефону під навантаженням, співвідношення ціна-якість. Співвідношення $\frac{S_j^-}{W_j^-} = 2,6$ свідчить, що розроблений метод

сепарує найкритичніші негативні властивості, які мають найбільший вплив на досвід користувача, від всієї сукупності негативних факторів.

7. Обговорення результатів

Розроблений метод оцінки негативних аспектів рішення інтелектуальної системи при побудові користувацької ментальної моделі відрізняється від існуючих комбінуванням оцінок сумарного та контекстно-орієнтованого впливу негативних аспектів рішення на сприйняття цього рішення користувачем системи, що дає можливість побудувати в рамках ментальної моделі комплексне представлення рішення, що включає переваги та суттєві обмеження при його використанні. Така інтеграція властивостей підвищує релевантність пояснень в інтелектуальних системах, що забезпечує можливість обґрунтованого вибору рішення користувачем.

Порівняння розробленого методу з методом оцінки впливу значень вхідних даних LIME, який орієнтований на пояснення впливу властивостей вхідних даних на рішення ІС дало можливість зробити висновок про принципову відмінність між даними підходами. Метод LIME формує локальне пояснення для конкретного прикладу, щоб визначити вплив вхідних даних. Розроблений метод враховує вплив негативних характеристик на всьому наборі вхідних даних, тобто узагальнює особливості використання рішень багатьма користувачами. Тобто LIME оцінює алгоритм роботи моделі, а запропонований метод оцінює властивості рішення. Аналогічно, метод SHAP також пояснює причини, чому модель прийняла відповідне рішення і не дає можливість врахувати контекст використання рішення. На відміну від методу SHAP, який виконується на етапі 3 і дає можливість узагальнити вплив вхідних даних на рішення, розроблений метод оцінює результати SHAP у контексті використання рішення ІС.

Отримане при проведенні експерименту співвідношення $\frac{S_j^-}{W_j^-} = 2,6$ підтверджує

ефективність виділення ключових проблемних властивостей рішення. Тобто близько чверті від усіх негативних властивостей має найбільший вплив на сприйняття рішення користувачем. Таке значення співвідношення свідчить, що лише незначна кількість недоліків призводить до значної кількості скарг користувачів.

Обмеження на використання даного методу пов'язані із релевантністю і повнотою вхідних даних, оскільки відгуки можуть бути сформовані в результаті шилінг-атак. Останні направлені на примусову зміну рейтингів конкурентних рішень. неструктурованих вхідних даних

Напрямки подальших досліджень щодо оцінки негативного впливу характеристик рішення пов'язані із розробкою моделей оцінки загального та контекстно-орієнтованого

впливу негативних властивостей з використанням машинного навчання, що дасть можливість використовувати інші структури вхідних даних, в тому числі мультимодальні дані.

8. Висновки

Запропоновано гібридний підхід до побудови ментальної моделі на основі оцінки негативних властивостей рішення інтелектуальної системи, який використовує комбіновану оцінку негативних властивостей вхідних даних для відбору цих властивостей при побудові ментальної моделі рішення. Такий підхід забезпечує прийняття користувачем обґрунтованого рішення з урахуванням його ключових негативних властивостей.

Запропоновано метод оцінки негативних аспектів рішення інтелектуальної системи при побудові ментальної моделі користувача, який містить етапи попередньої обробки вхідних неструктурованих даних, SHAP-аналізу важливості негативних властивостей вхідних даних, контекстного аналізу негативних властивостей, формування комбінованої оцінки рішення ІС. Метод дає можливість виявляти та упорядковувати негативні властивості рішення інтелектуальної системи, що створює умови для побудови зрозумілих для користувача пояснень та підвищення довіри користувачів внаслідок представлення потенційних обмежень для рішень ІС.

Перелік посилань:

1. Miller T. Explanation in artificial intelligence: Insights from the social sciences. Artificial Intelligence. 2019. Vol. 267. P. 1-38.
2. Salih A.M., Amin R., Shanta S.S. A Perspective on Explainable Artificial Intelligence Methods: SHAP and LIME. arXiv preprint arXiv:2305.02012. 2023.
3. Gunning D., Aha D.W. DARPA's explainable artificial intelligence (XAI) program. AI Magazine. 2019. Vol. 40. № 2. P. 44-58.
4. Sun Q., Liu Y., Chen H. et al. Explainable Artificial Intelligence for Medical Applications: A Review. arXiv preprint arXiv:2412.01829. 2024.
5. Cheung J.C., Wong A., Liu H. The effectiveness of explainable AI on human factors in trust and decision making. Nature Scientific Reports. 2025. Vol. 15. Art. 4189.
6. Sheridan H., Clarke D., Murphy C. Exploring Mental Models for Explainable Artificial Intelligence. HCII 2023 Conference Proceedings. 2023.
7. Johnson-Laird P.N. Mental models and human reasoning. Proceedings of the National Academy of Sciences. 2013. Vol. 110. № 45. P. 18243-18250.
8. Чалий С.Ф., Лещинська І.О. Концептуальна ментальна модель пояснення в системі штучного інтелекту. Бюлетень Національного технічного університету "ХПІ". Серія: Системний аналіз, управління та інформаційні технології. 2023. № 1. С. 41-46.
9. Doshi-Velez F., Kim B. Towards a rigorous science of interpretable machine learning. arXiv preprint arXiv:1702.08608. 2017.
10. Salih A.M., Amin R., Shanta S.S. A Perspective on Explainable Artificial Intelligence Methods: SHAP and LIME. arXiv preprint arXiv:2305.02012. 2023.
11. Lundberg S.M., Lee S.I. A unified approach to interpreting model predictions. Proceedings of the 31st International Conference on Neural Information Processing Systems. 2017. P. 4768-4777.
12. 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. 2016. P. 1135-1144.
13. Чалий С.Ф., Лещинська І.О. Уточнення ментальної моделі рішення користувача інтелектуальної системи шляхом доповнення вхідних даних. Автоматизовані системи управління та прилади автоматики. 2024. № 188. С. 4-12.
14. Чалий С.Ф., Лещинський В.О., Лещинська І.О. Каузальна ментальна модель для пояснення в інтелектуальних системах. Автоматизовані системи управління та прилади автоматики. 2024. № 189. С. 12-20.
15. Sun Q., Liu Y., Chen H. et al. Explainable Artificial Intelligence for Medical Applications: A Review. arXiv preprint arXiv:2412.01829. 2024.
16. Jeong S., Miller T., Park J. AI Mental Models & Trust: The Promises and Perils of User Mental Models. EPIC Conference Proceedings. 2024. P. 234-251.
17. Dangol A., Peters C., Sharma K. Children's Mental Models of AI Reasoning. arXiv preprint

arXiv:2505.16031. 2025.

18. Norman D.A. Some observations on mental models. In D. Gentner & A. L. Stevens (Eds.), *Mental models* (pp. 7–14). Lawrence Erlbaum Associates. 1983.

19. Чалий С.Ф., Лещинська І.О. Функціонально-темпоральне представлення ментальної моделі користувача для генерації пояснень. *Автоматизовані системи управління та прилади автоматики*. 2025. № 190. С. 15-24.

20. Чалий С.Ф., Лещинська І.О. Принципи побудови ментальних моделей зовнішніх користувачів інтелектуальних систем. *Вісник Національного технічного університету "ХПІ"*. 2024. № 7. С. 89-95.

21. Liu B. Sentiment analysis and opinion mining. *Synthesis lectures on human language technologies*. 2012. Vol. 5. № 1. P. 1-167.

22. Ahmed M.T., Rhaman S., Kader A. A systematic review of explainable artificial intelligence for spectroscopic models. *Computers and Electronics in Agriculture*. 2025. Vol. 228. Art. 109462.

23. Ullah N., Khan J.A., El-Sappagh S. A novel explainable AI framework for medical image analysis. *Medical Image Analysis*. 2025. Vol. 97. Art. 103241.

Надійшла до редколегії 18.09.2025 р.

Чалий Сергій Федорович, доктор технічних наук, професор, професор кафедри ІУС ХНУРЕ, м. Харків, Україна, e-mail: serhii.chalyi@nure.ua; ORCID: <https://orcid.org/0000-0002-9982-9091>

Лещинська Ірина Олександрівна, кандидат технічних наук, доцент, доцент кафедри ПІ ХНУРЕ, м. Харків, Україна, e-mail: iryna.leshchynska@nure.ua, ORCID: <https://orcid.org/0000-0002-8737-4595>

УДК 004.8:004.9

DOI: 10.30837/0135-1710.2025.186.103

*О.В. ЧАЛА, О.М. БІТЧЕНКО, Л.Ф. САЙКІВСЬКА, М.Є. АЛФЬОРОВ, Д.Г. ГАНИШИН***ПРЕСКРИПТИВНА МОДЕЛЬ ТЕМПОРАЛЬНИХ ЗНАНЬ З OWL-ІНТЕГРАЦІЄЮ ДЛЯ ПІДТРИМКИ РІШЕНЬ В ІНФОРМАЦІЙНИХ СИСТЕМАХ**

Об'єктом дослідження є процес формування рекомендацій в інформаційних системах на основі інтеграції темпоральних правил із стандартизованими онтологічними конструкціями для підтримки управлінських рішень. Метою є розробка прескриптивної моделі темпоральних знань з інтеграцією OWL-онтологій для підтримки рішень в інформаційних системах, з тим, щоб забезпечити можливість автоматизованого формування рекомендацій з управління з урахуванням семантики управлінських рішень. Розроблено дворівневе представлення темпоральних знань, яке поєднує статистичні темпоральні правила із стандартизованими онтологічними конструкціями для створення двошарової архітектури знань. Запропоновано прескриптивну модель темпоральних знань, яка інтегрує множину темпоральних правил із відповідними онтологічними відношеннями через систему комбінованих ваг, що створює умови для автоматизованого формування управлінських рекомендацій на основі як історичного досвіду успішних рішень, так і семантичної узгодженості з існуючими знаннями предметної області, забезпечуючи перехід від дескриптивного аналізу до прескриптивного підходу до підтримки рішень.

1. Вступ

Підтримка управлінських рішень в інформаційних системах (ІС) передбачає обґрунтований вибір одного із множини можливих варіантів дій, орієнтованого на досягнення визначених цілей організації в умовах обмежених ресурсів та темпоральних обмежень. Процеси підтримки прийняття таких рішень реалізують складні багатоетапні процедури, що включають три основні фази: аналізу поточної ситуації, формування альтернативних варіантів рішень та вибору одного з варіантів. Підтримка управлінських рішень широко застосовується у фінансовій сфері для аналізу ризиків та запобігання шахрайству, у промисловості для управління взаємовідносинами з клієнтами та статистичного управління запасами, у телекомунікаціях при проєктуванні нових послуг на основі аналізу клієнтської поведінки [1], [2].

Процес підтримки рішень використовує знання як щодо предметної області, так і щодо виведення нових висновків, прогнозування розвитку ситуацій та генерації рекомендацій. У банківському секторі такі знання використовуються для аналізу кредитних ризиків та розробки VIP-програм клієнтського обслуговування, у виробничих підприємствах – для бюджетного планування, в задачах управління кадрами – для прогнозування потреб у персоналі [3], [4]. В контексті управлінських рішень суттєве значення має темпоральний аспект знань, оскільки темпоральні знання дають можливість врахувати послідовність подій, часові інтервали між діями та їхніми наслідками у процесі виконання управлінського рішення. Сучасні підходи до представлення темпоральних знань у системах підтримки рішень переважно базуються на описі вже реалізованих рішень, використовуючи темпоральні закономірності у даних, які характеризують такі рішення [5], [6]. Однак орієнтація на описові, дескриптивні підходи створює суттєвий розрив між наявними можливостями темпорального опису реалізованих управлінських рішень та практичними потребами у рекомендаціях щодо майбутніх дій для досягнення поставлених цілей при змінах не лише стану предметної області, а й задач управління [7], [8]. Даний розрив пов'язаний із обмеженими можливостями існуючих методів щодо автоматизованої генерації рекомендацій з управлінських дій на основі поєднання формалізованого досвіду минулих управлінських рішень у формі темпоральних залежностей

з семантичним описом актуального стану предметної області та задач управління.

Альтернативний, прескриптивний підхід відрізняється тим, що він орієнтований на формування рекомендацій щодо нових управлінських рішень, які необхідно виконати для досягнення цілей управління [9], [10]. Прескриптивність у контексті темпоральних знань означає здатність ІС не тільки інтерпретувати темпоральні залежності між подіями, але й використовувати ці знання для автоматизованого формування нових наборів послідовностей дій у складі управлінського рішення, орієнтованих на досягнення цільових станів системи відносно поточного стану [11], [12]. Тобто прескриптивний підхід орієнтований не лише на вибір одного із відомих рішень, а й на формування нових послідовностей дій управлінського рішення з урахуванням поточного семантичного опису предметної області.

Таким чином, в умовах підвищення складності задач, що вирішують ІС, та підвищення вимог до швидкості прийняття рішень виникає протиріччя між потребами практики в підтримці автоматизованої генерації управлінських рекомендацій та обмеженими можливостями існуючих дескриптивних методів, що і визначає актуальність проблеми розробки прескриптивних підходів до використання темпоральних знань для підтримки управлінських рішень.

2. Аналіз сучасних наукових публікацій і постановка проблеми дослідження

Дослідження в галузі систем підтримки прийняття рішень базуються на визначенні структури процесу прийняття рішень та обґрунтуванні принципів проєктування систем підтримки прийняття рішень як підкласу ІС [13].

Роботи у сфері управління знаннями [14]–[16], включаючи дослідження Nonaka І. та Takeuchi Н. [16], орієнтовані в першу чергу на процеси створення та трансформації знань в організаціях. Однак ці підходи не враховують специфіку темпоральних залежностей у процесах прийняття управлінських рішень та обмежуються дескриптивним характером аналізу.

Дослідження темпоральних аспектів знань представлено роботами, що розглядають логіку часу та її застосування у формальних системах [11], [14]. Зокрема у роботах Allen J.F. представлено інтервальну алгебру [14]. Проте більшість існуючих підходів зосереджується на теоретичних аспектах темпоральної логіки без практичної орієнтації на системи підтримки управлінських рішень. У ряді досліджень розглядається проблема формалізації знань у сучасних системах з ситуаційною обізнаністю на основі використання темпоральних онтологій, зокрема OWL-Time [17]. Аналіз практики використання онтологій у задачах підтримки рішень свідчить про актуальність стандартів W3C, зокрема OWL та RDF. OWL (Web Ontology Language) – це стандартизована мова для створення онтологій, розроблена консорціумом W3C (World Wide Web Consortium). Даний стандарт призначений для формального опису сутностей часу та темпоральних відношень на основі визначення понять моментів, інтервалів, періодів часу та відношень між ними [18]. У роботах [11], [19] досліджуються ефективні алгоритми для швидкої обробки великих масивів знань. Проте ці роботи орієнтовані в першу чергу на розгляд технічних аспектів логічного виведення, представлення та аналізу темпоральних відношень без прескриптивних варіантів застосування цих онтологій.

Огляд сучасних підходів до моделювання часу, в тому числі онтологічного моделювання, включаючи логіку Аллена та темпоральне розширення RDF, представлено в роботах [2], [8], [18]. Однак в цих роботах не розглядається можливість створення прескриптивних моделей з урахуванням формального опису часу.

Сучасні тенденції у розвитку бізнес-аналітики свідчать про перехід від дескриптивних до прескриптивних підходів [19]. Проте в цих роботах не розглядається специфіка інтеграції темпоральних знань із прескриптивними механізмами.

Дослідження [20]–[23] присвячені розробці комплексного підходу до управління

темпоральними базами знань [20] на основі формування та використання темпоральних правил [21], [22]. Проте в цих роботах не розглядається аспект семантики управлінських рішень з використанням онтологічного моделювання часу, зокрема з використанням стандарту OWL-Time.

Таким чином, існуючі дослідження що підтримки управлінських рішень орієнтовані в першу чергу на попереднє формування наборів альтернативних рішень і не враховують можливість вирішення задачі інтеграції опису реалізації рішення на основі темпоральних правил, що визначають послідовність дій рішення та темпоральної онтології, що визначає семантику цих дій. Проте така інтеграція має суттєву перевагу. Вона дає можливість в рамках визначеної темпоральними правилами послідовності дій, що мають виконуватись при реалізації управлінського рішення, сформувати декілька альтернативних варіантів рішення так, що вибір одного із варіантів виконується з урахуванням його семантики, представленої в рамках темпоральної онтології.

3. Мета і задачі дослідження

Метою дослідження є розробка прескриптивної моделі темпоральних знань з інтеграцією OWL-онтологій для підтримки рішень в ІС, з тим, щоб забезпечити можливість автоматизованого формування рекомендацій з управління з урахуванням семантики управлінських рішень.

Для досягнення даної мети необхідно вирішити такі задачі: розробити представлення темпоральних знань на основі інтеграції темпоральних правил та стандартизованих OWL-Time конструкцій; розробити прескриптивну модель темпоральних знань з можливостями семантичного відображення темпоральних правил на темпоральні онтологічні відношення.

4. Розробка представлення темпоральних знань на основі інтеграції темпоральних правил та OWL-Time онтологій

Об'єктом дослідження є процес формування рекомендацій в ІС на основі інтеграції темпоральних правил із стандартизованими онтологічними конструкціями для підтримки управлінських рішень. Предметом дослідження є методи формування рекомендацій в ІС на основі інтеграції темпоральних правил із стандартизованими онтологічними конструкціями для підтримки управлінських рішень.

Онтологія в контексті ІС є формальною моделлю знань, яка визначає базові поняття предметної області, їхні властивості та взаємозв'язки у вигляді, що забезпечує можливість автоматизованої обробки в ІС. Онтології створюють семантичну структуру, яка дає можливість оперувати зі змістом збереженої в цій структурі інформації і виводити нові знання на цій основі.

Розроблена онтологія темпоральних знань на основі стандартної онтології OWL-Time для підтримки рішень в ІС представлена на рис. 1. Базові класи OWL-Time описують темпоральні сутності (TemporalEntity), інтервали часу (time:Interval), моменти часу (time:Instant), а також протяжність інтервалу в часі (time:TemporalDuration).

Додаткові класи предметної області описують стани системи, що змінюються в часі (TemporalState), конкретні дії, що виконуються при зміні станів системи (ManagementAction), багатоетапні процеси з послідовностей станів та дій, що відображають переходи між станами на основі темпоральних відношень (Process).

Базові властивості в OWL-time онтології відображають початок (time:hasBeginning), завершення (time:hasEnd), тривалість (time:hasDuration) інтервалу, процесу. Розширені властивості відображають специфічні для предметної області характеристики, наприклад температуру (hasTemperature), управлінські рекомендації (hasRecommendation).

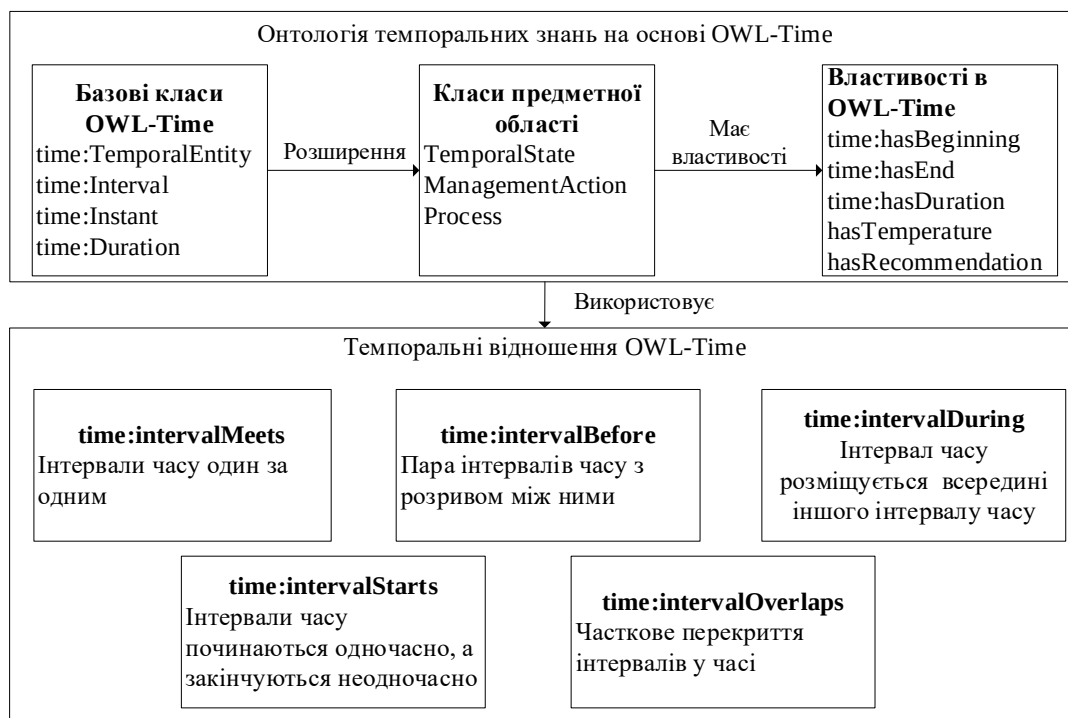


Рис. 1. Онтологія темпоральних знань на основі OWL-Time

Верхній рівень розробленого представлення темпоральних знань формується із правил NeXt, Future та Until (див. рис. 2). Правило NeXt відображає безпосереднє слідування подій у управлінському рішенні. Наприклад, послідовність «після завершення бюджетування виконується планування закупівель» представляється у вигляді правила NeXt(Бюджетування, План закупівель). Дане правило визначає таку залежність як обов'язкову послідовність робіт без проміжних дій.

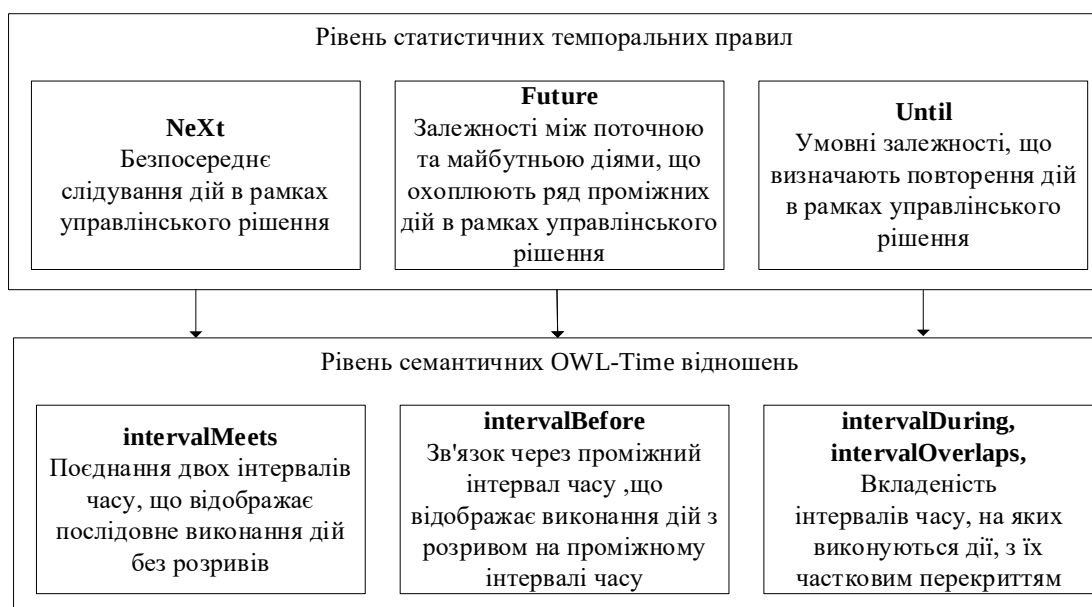


Рис. 2. Представлення темпоральних знань на основі поєднання темпоральних правил та OWL-Time онтологій

Темпоральне правило Future характеризує розділені у часі залежності між діями в рамках управлінського рішення. Наприклад, «розширення штату співробітників матиме як наслідок збільшення витрат на заробітну плату в майбутньому періоді» задається правилом Future(Розширення штату співробітників, Збільшення заробітної плати).

Правило Until задає умови, які перевіряються при виконанні дій рішення. Наприклад, правило Until(Моніторинг, Досягнення цілі, Уточнення стратегії) означає «потрібно продовжувати моніторинг показників до досягнення або зміни стратегії».

Нижній рівень розробленого представлення задається онтологічними темпоральними відношенням OWL-Time. Правило NeXt на даному рівні відображається як відношення безпосереднього стикування інтервалів часу «intervalMeets». Правило Future відображається на онтологічне відношення «intervalBefore» в OWL-Time, коли між заданими подіями існує інтервал часу й зберігається причинно-наслідковий зв'язок. Правило відображається на комбінацію онтологічних відношень «intervalDuring» та «intervalOverlaps». Перше відношення задає, що один процес відбувається в рамках іншого, а друге – що процеси частково перекриваються в часі.

Призначення обох рівнів даного представлення полягає в такому. Темпоральні правила визначаються на основі статистики щодо реалізації управлінських рішень. Розраховуються історичні дані про реалізовані послідовності подій, частоти переходів між станами і на цій основі виконується оцінка правил. Тобто правила обґрунтовують послідовність дій на основі статистичних даних.

Онтологічна складова відображає темпоральні правила через стандартизовані семантичні конструкції, що дає можливість виявляти логічні суперечності та виводити нові залежності на основі транзитивності темпоральних відношень.

5. Прескриптивна модель темпоральних знань

Формальний опис прескриптивної моделі темпоральних знань базується на дворівневому представленні таких знань, розглянутому у попередньому розділі.

Темпоральні правила $r_{j,m}$ верхнього рівня задають упорядкованість у часі між j -им та m -им станами об'єкта управління, що послідовно виникають внаслідок реалізації послідовності дій у складі управлінського рішення. Правила мають вагу $w_{j,m}^{temp}$, що визначається на основі відношення кількості реалізацій правила до загальної кількості можливих реалізацій всіх правил з використанням відомих даних про виконанні управлінські рішення. Відповідно, вища вага правил означає, що при минулих реалізаціях управлінських рішень ці правила частіше приводили до успішних результатів.

Онтологічна вага відображає коректність семантичного опису $o_{j,m}$ правила $r_{j,m}$ за умови $C(r_{j,m})$, що це правило не протирічить існуючим, тобто що не існує протиріччя з будь-яким іншим правилом $l_{j,m}$:

$$C(o_{j,m}) = \begin{cases} 1, & \text{if } \exists l_{j,m} = \neg o_{j,m}, \\ 0, & \text{otherwise.} \end{cases} \quad (1)$$

Вага визначається на базі кількості нових правил, які можуть бути побудовані в комбінації із правилом $o_{j,m}$ на основі відношення транзитивності. Тобто онтологічна вага $w_{j,m}^{ont}$ правила визначається на основі нормованої кількості відношень $n(l_{j,m}, o_{j,m})$ з іншими правилами $l_{j,m}$:

$$w_{j,m}^{ont} = n(l_{j,m}, o_{j,m}) | C(o_{j,m}) > 0. \quad (2)$$

Композиція правил $(l_{j,m}, o_{j,m})$ дає можливість отримувати нові правила, що визначають семантику шляху до повної реалізації рішення.

Загальна вага правил на двох рівнях становить зважену суму темпоральних правил $w_{j,m}^{temp}$ та відповідних онтологій $w_{j,m}^{ont}$ з коефіцієнтам a та b відповідно:

$$w_{j,m} = aw_{j,m}^{temp} + bw_{j,m}^{ont}. \quad (3)$$

Коефіцієнти a та b визначаються згідно з особливостями предметної області. Коефіцієнт a для темпоральних правил впливає на ефективність рішення, оскільки він відображає статистику успішних реалізацій управлінських рішень. Коефіцієнт b для онтологічних правил впливає на узгодженість правила з іншими правилами, що використовуються в даній предметній області. Тобто якщо ключовим у предметній області є накоплений досвід, то збільшується перший коефіцієнт, а якщо потрібно враховувати відсутність протиріч із існуючими знаннями про процеси в предметній області, то збільшується другий коефіцієнт.

Прескриптивна модель M темпоральних знань включає в себе множину $R = \{r_{j,m}\}$ темпоральних правил типів NeXt, Future та Until, множину відповідних онтологічних відношень O та множину ваг $W = \{w_{j,m}\}$:

$$M = (R, O, W). \quad (4)$$

Комбінація ваг (3) дає можливість вибирати такі дії в рамках управлінського рішення, які, з одного боку, приводили до успіху в минулому (тобто підтверджено їхню ефективність), а з іншого боку, відповідають формальним правилам щодо узгодженості з існуючою системою знань. Тобто для заданих вагових коефіцієнтів a та b можуть бути вибрані дії з максимальною вагою, що свідчить про їхню ефективність у минулому та поточну узгодженість із знаннями щодо предметної області. Такий підхід дає можливість реалізувати зворотний зв'язок щодо вибору рекомендацій. Результати поточного вибору правил впливають на майбутні ваги $w_{j,m}^{temp}$, що приводить до подальшого уточнення послідовності дій в рамках управлінського рішення.

Узагальнену послідовність побудови рекомендацій з управління на основі моделі наведено на рис. 3

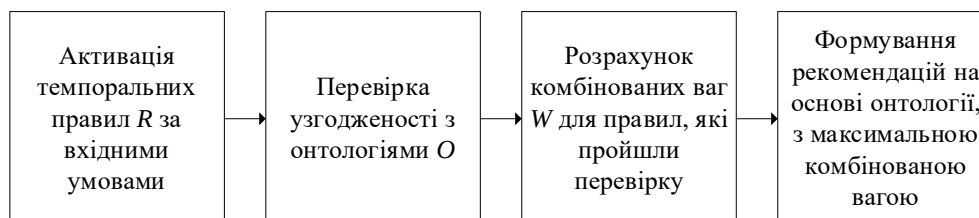


Рис. 3. Послідовність побудови рекомендацій з використанням темпоральних правил та OWL-Time онтологій

Розглянемо приклад перевірки узгодженості темпорального правила з онтологією. Якщо темпоральне правило визначає, що при перевищенні температури необхідно зупинити обладнання через 10 хвилин, а згідно з онтологією зупинка потребує 20-хвилинного інтервалу для даного обладнання, то дане темпоральне правило блокується.

6. Експериментальна перевірка властивостей моделі темпоральних знань

Експериментальна перевірка можливостей розробленої моделі виконана з використанням двох наборів даних.

Перший набір відображає зміну стану підшипника у процесі його роботи на промисловому обладнанні. Використані показники таких IoT-сенсорів:

- температура підшипника (змінюється діапазоні 25-85°C);
- вібрації (0,05-15,0 g);
- швидкість обертання (1200-2000 об/хв);
- сила струму двигуна (5-25 А).

Для даного набору виділено чотири категорії станів підшипника (за значенням ваги правил): нормальний (normal); потребує попередження (warning); аварійний (alarm); критичний (critical).

Пороги переходів між станами встановлено згідно з рекомендаціями стандартів IEEE 1349-2001 та ISO 13373-1:2002. Згідно з цими рекомендаціями, порогове значення температури для переходу до попередження про стан підшипників було визначено як 65°C. Порогове значення вібрації було визначено на рівні 2.0g.

Датасет містить 1933 переходи між цими станами.

Приклади отриманих темпоральних правил NeXt та Future та відповідних OWL-Time відношень наведено в табл. 1.

Таблиця 1

Приклади темпоральних правил та відповідних OWL-Time відношень

Темпоральне правило	Статистична характеристика правила, частота обертання (оборотів/сек.)	OWL-Time відношення
NeXt(warning → normal)	0,136	intervalMeets(MaintenanceAction, NormalOperation)
Future (warning → critical)	0,15	intervalMeets(MaintenanceAction, NormalOperation)

Приклад правила-попередження, отриманого з використанням логічного виведення за допомогою OWL-reasoner, представлено на рис. 4.

```

HighTemperature ≡ ∃hasTemperature.value >65°C
HighVibration ≡ ∃hasVibration.value >2.0g
HighRiskState ≡ WarningState ∧ HighTemperature ∧ HighVibration
HighRiskState → hasRecommendation.ImmediateInspection

```

Рис. 4. Правило переходу до стану попередження

На основі таких правил модель генерує рекомендації щодо управлінських дій. При переході між станами normal→warning представлено на рис. 4 правило-попередження

генерує рекомендацію ImmediateInspection – провести інспекцію стану відповідного підшипника. Вибір правила серед альтернатив обумовлений його вагою.

Порівняльний аналіз переходів між станами показав, що покриття ситуацій переходів правилами становить 88 % за умови 100 % логічної цілісності.

Другий набір даних Online Retail Dataset із репозиторію UCI Machine Learning Repository містить інформацію про покупки в інтернет-магазині подарункових наборів 4372 клієнтами. Приклад дворівневого представлення темпоральних знань для даного набору представлено на рис. 5.

```
rdfs:label "Future правило: Кошик → Нагадування"@uk
```

```
ecom:hasTemporalWeight "0.73",  
ecom:hasOntologicalWeight "0.88",  
ecom:hasCombinedWeight "0.727"  
ecom:hasAbandonmentThreshold "PT24H"  
ecom:hasReminderDelay "PT6H"
```

Рис. 5. Приклад представлення знань в системі електронної комерції

RDF-конструкція `ecom:FutureRule_CartReminder` є представленням темпорального правила Future-типу в рамках прескриптивної моделі темпоральних знань. Ідентифікатор `com:FutureRule` є підкласом базового класу `ecom:TemporalRule`. Імплементація правила Кошик → Нагадування відображає послідовність дій у часі. Запуск правила визначається параметром `ecom:hasAbandonmentThreshold "PT24H"`, тобто товар в кошику має перебувати 24 години перед запуском правила. Темпоральна вага `ecom:hasTemporalWeight "0.73"` відображає відсоток успішних реалізацій правила, як було представлено в розділі 4. Онтологічна вага `ecom:hasOntologicalWeight "0.88"` відображає ступінь узгодженості правила з існуючою онтологією згідно з (2). Комбінована вага `ecom:hasCombinedWeight "0.727"` розраховується згідно з (3) з коефіцієнтами 0.6 та 0.4, що відображають важливість інформації про успішне виконання правил. Ідентифікатор `ecom:hasReminderDelay "PT6H"` задає затримку для відправки нагадування. Представлене на рис. 5 правило задає нагадування користувачеві системи e-commerce щодо відібраних ним у корзину товарів. Сукупність таких темпоральних правил, зокрема «Until: Сесія активна → оновлення рекомендації з інтервалом T секунд», «NeXt: Користувач на сторінці → рекомендації згідно зі змістом сторінки», представлених OWL-Time онтологіями, дає можливість сформувати процес підтримки рішень щодо взаємодії з клієнтом системи електронної комерції.

Таким чином, експериментальна перевірка показала, що дана модель на відміну від існуючих підходів дає можливість динамічно формувати рекомендації з урахуванням як досвіду імплементації рішень, так і онтологічного опису предметної області.

7. Обговорення результатів

Запропонована модель інтегрує статистичні та онтологічні підходи до представлення темпоральних знань. Статистична складова забезпечує підтримку успішного досвіду управлінських рішень. Онтологічна складова темпоральних знань забезпечує відсутність протиріч із описом предметної області, що дає можливість динамічно, в залежності від поточного стану об'єкта управління, формувати рекомендації з урахуванням семантики дій в рамках управлінського рішення.

Використання комбінованих ваг (3) створює умови для адаптації режиму формування рекомендацій в залежності від типу задач, які вирішує ІС. Збільшення ваги традиційних темпоральних правил дає можливість використати наявний досвід успішної реалізації

управлінських рішень. Збільшення ваги онтологічних правил дає можливість формувати рекомендації у відповідності до знань про предметну область.

Обмеження даного підходу пов'язані із суттєвими витратами на побудову онтологій спеціалізованих предметних областей. Відсутність опису предметної області без протиріч суттєво обмежує онтологічний рівень моделі і така модель спрощується до традиційних темпоральних правил.

Суттєва перевага даного підходу полягає у можливості еволюційного уточнення темпоральної складової моделі, оскільки вибір рішення приводить до зміни статистичних характеристик темпоральних правил. Така зміна, в свою чергу, може привести до подальшого уточнення рекомендацій.

Подальший розвиток моделі пов'язаний із використанням методів пояснювального штучного інтелекту для забезпечення прозорості процесу формування рекомендацій. В практичному плані інтеграція можливостей пояснювального штучного інтелекту може бути реалізована на основі доповнення (3) компонентом, який враховує можливість пояснення прийнятого рішення, що створить умови для обґрунтованого вибору людиною, що приймає рішення.

8. Висновки

Розроблено дворівневе представлення темпоральних знань для підтримки рекомендацій щодо управлінських рішень в ІС. Розроблене представлення комбінусє можливості багатоваріантного опису послідовності дій управлінського рішення з можливостями його доповнення з урахуванням онтологічного опису предметної області. В практичному аспекті розроблене представлення створює умови для автоматизованого формування рекомендацій щодо дій управлінського рішення.

Запропоновано прескриптивну модель темпоральних знань, яка містить множину темпоральних правил, множину темпоральних онтологій, а також множину комбінованих ваг, що визначають вплив темпоральних правил та темпоральних онтологій на формування рекомендацій щодо дій управлінського рішення. Модель дає можливість формувати рекомендації щодо майбутніх дій управлінського рішення на основі композиції правил з перевіркою відповідності цих правил існуючим знанням щодо предметної області.

Перелік посилань:

1. Turban, E., Sharda, R., & Delen, D. (2020). Decision support and business intelligence systems (11th ed.). Pearson.
2. Power, D. J. (2021). A brief history of decision support systems. *Journal of Decision Systems*, 30(2), 73–92. <https://doi.org/10.1080/12460125.2021.1893339>
3. Bloomfire. (2025). How knowledge management enhances decision-making. <https://bloomfire.com/blog/how-knowledge-management-enhances-decision-making/>
4. Abacademies. (2021). The impact of decision-making models and knowledge management practices on performance. *International Journal of Economics and Business Modeling*, 12(1), 45–61. <https://doi.org/10.31961/ijebm.v12i1.10558>
5. EliXR-TIME Consortium. (2022). EliXR-TIME: A temporal knowledge representation for clinical process mining. *Journal of Biomedical Informatics*, 135, Article 104103. <https://doi.org/10.1016/j.jbi.2022.104103>
6. Hou, Z., Jin, X., Li, Z., Bai, L., Guan, S., Zeng, Y., Guo, J., & Cheng, X. (2023). Temporal knowledge graph reasoning based on N-tuple modeling. In *Findings of the Association for Computational Linguistics: EMNLP 2023* (pp. 1090–1100). Association for Computational Linguistics. <https://aclanthology.org/2023.findings-emnlp.77>
7. Bertsimas, D., & Kallus, N. (2020). From predictive to prescriptive analytics. *Management Science*, 66(3), 1025–1044. <https://doi.org/10.1287/mnsc.2018.3253>
8. Lepenioti, K., Bousdekis, A., Apostolou, D., & Mentzas, G. (2020). Prescriptive analytics: Literature review and research challenges. *International Journal of Information Management*, 50, 57–70. <https://doi.org/10.1016/j.ijinfomgt.2019.09.003>
9. Davenport, T. H., & Harris, J. G. (2024). *Competing on analytics 2.0: Prescriptive, cognitive, and augmented intelligence*. Harvard Business Review Press.

10. EliXR-TIME Consortium. (2024). Automated generation of prescriptive process recommendations using temporal patterns. *Data & Knowledge Engineering*, 150, Article 103012. <https://doi.org/10.1016/j.datak.2023.103012>
11. EliXR-TIME Consortium. (2024). Scalable temporal reasoning algorithms for large-scale knowledge graphs. In *Proceedings of the 2024 International Conference on Knowledge Graphs*, 54–63. <https://doi.org/10.1145/3618000.3618012>
12. Shoush, M., Dumas, M., La Rosa, M., & Maggi, F. M. (2024). White box specification of intervention policies for prescriptive process monitoring. *Data & Knowledge Engineering*, 149, Article 102271. <https://doi.org/10.1016/j.datak.2023.102271>
13. Simon, H. A. (1977). *The new science of management decision* (Revised ed.). Prentice-Hall.
14. Allen, J. F. (1983). Maintaining knowledge about temporal intervals. *Communications of the ACM*, 26(11), 832–843. <https://doi.org/10.1145/358434.358439>
15. Чалий, С. Ф., & Прибильнова, І. Б. (2021). Представлення узгоджених знань з урахуванням їх логічної несуперечливості для задачі побудови пояснень в інтелектуальних інформаційних системах. *Інформаційні технології та комп'ютерні системи*, 28, 45–54. <https://doi.org/10.33099/itcs.2021.28.045>
16. Nonaka, I., & Takeuchi, H. (2020). The knowledge-creating company revisited: A dynamic model of knowledge creation. *Knowledge Management Research & Practice*, 18(3), 279–302. <https://doi.org/10.1080/14778238.2020.1747980>
17. W3C OWL Working Group. (2012). *OWL 2 Web Ontology Language Document Overview* (2nd ed.). W3C Recommendation. <https://www.w3.org/TR/owl2-overview/>
18. W3C Working Group. (2017). *Time Ontology in OWL*. W3C Recommendation. <https://www.w3.org/TR/owl-time/>
19. Gruber, T. R. (2021). Toward principles for the design of ontologies used for knowledge sharing. *International Journal of Human-Computer Studies*, 122, 25–37. <https://doi.org/10.1016/j.ijhcs.2018.02.004>
20. Levykin, V., & Chala, O. (2018). Method of automated construction and expansion of the knowledge base of the business process management system. *EUREKA: Physics and Engineering*, 4, 29-35. <https://doi.org/10.21303/2461-4262.2018.00676>
21. Chala, O. (2018). Construction of temporal rules for the representation of cause-effect dependencies in business processes. *Advanced Information Systems*, 2(3), 54-59. <https://doi.org/10.20998/2522-9052.2018.3.09>
22. Chala, O. (2018). Models of temporal dependencies for a probabilistic knowledge base. *Econtechmod. An International Quarterly Journal*, 7(3), 53-58. <https://yadda.icm.edu.pl/yadda/element/bwmeta1.element.baztech-9d6a1102-180c-4298-b6b4-998723111507>
23. Chala, O. (2019). Development of information technology for the automated construction and expansion of the temporal knowledge base in the tasks of supporting management decisions. *Technology Audit and Production Reserves*, 1(2), 9-14. <https://doi.org/10.15587/2312-8372.2019.160205>

Надійшла до редколегії 18.09.2025 р.

Чала Оксана Вікторівна, доктор технічних наук, доцент, завідувачка кафедри РТІКС ХНУРЕ, м. Харків, Україна, e-mail: oksana.chala@nure.ua; ORCID: <https://orcid.org/0000-0002-9982-9091>

Бітченко Олександр Миколайович, кандидат технічних наук, доцент, доцент кафедри РТІКС ХНУРЕ, м. Харків, Україна, e-mail: oleksandr.bitchenko@nure.ua; ORCID: <https://orcid.org/0000-0002-4561-1046>

Сайківська Лілія Федорівна, кандидат технічних наук, доцент, доцент кафедри РТІКС ХНУРЕ, м. Харків, Україна, e-mail: liliia.saikivska@nure.ua; ORCID: <https://orcid.org/0000-0002-4139-7732>

Алфьоров Микола Євгенович, старший викладач кафедри РТІКС ХНУРЕ, м. Харків, Україна, e-mail: mykola.alforov@nure.ua, ORCID: <https://orcid.org/0000-0002-1590-3902>

Ганшин Дмитро Геннадійович, старший викладач кафедри РТІКС ХНУРЕ, м. Харків, Україна, e-mail: dmytro.hanshyn@nure.ua; ORCID: <https://orcid.org/0000-0002-3294-4220>

РЕФЕРАТИ / ABSTRACTS

УДК 004.93

Розробка комбінованого методу аналізу емоційної забарвленості текстів / К. Е. Петров, І. П. Боков, І. В. Кобзев. АСУ та прилади автоматики. 2025. № 186. С. 5-16.

Однією з ключових задач обробки природної мови (NLP) є аналіз емоційної забарвленості тексту, який відіграє важливу роль у численних прикладних сферах, зокрема в маркетингу, соціології, психології, аналізі громадської думки та інформаційній безпеці. Системи аналізу забарвленості текстової інформації дозволяють оперативно отримувати структуровану інформацію про емоційні настрої суспільства, прогнозувати реакцію на певні події, а також виявляти потенційні загрози чи деструктивний контент.

Попри досягнуті значні успіхи в галузі NLP, існуючі методи визначення емоційної забарвленості текстів мають ряд обмежень, які знижують їхню ефективність. Зокрема, традиційні методи часто не враховують контекстуального значення слів, що є критично важливим для точного розпізнавання емоційної забарвленості. Крім того, деякі методи мають труднощі при аналізі багатозначних слів, сарказму, іронії та сленгових виразів. Тому актуальним завданням є подальше вдосконалення методів аналізу тональності тексту, зокрема через поєднання кількох методів та використання моделей глибокого навчання.

Метою дослідження є підвищення точності класифікації емоційної забарвленості природномовних текстів за рахунок використання лексиконних, статистичних і контекстуальних методів, які дозволять врахувати як поверхневі лексичні ознаки, так і глибокі семантичні зв'язки у тексті.

В запропонованому комбінованому методі поєднуються статистичне (TF-IDF) та контекстуальне (BERT) векторні представлення тексту. Таке поєднання дозволяє враховувати як частотні закономірності, так і глибокі семантичні залежності між словами. Використання ансамблевого класифікатора Random Forest дозволило побудувати стійку модель, яка здатна ефективно класифікувати короткі англійські тексти з високим рівнем точності.

Результати експериментів показали, що запропонований комбінований метод має вищу точність класифікації (89 %) текстів, у порівнянні з базовими – TF-IDF + RF та BERT + RF (78 % і 82 % відповідно).

Використання комбінованого методу дозволить підвищити ефективність аналізу контексту, розпізнавання складних мовних конструкцій, що робить його перспективним для аналізу громадської думки в соціальних мережах, медіа та чат-ботах; для застосування у службах підтримки клієнтів; при визначенні емоцій користувачів веб-сервісів, сайтів та веб-додатків.

Ключові слова: емоційна забарвленість, природна мова, лексичний аналіз, машинне навчання, глибока нейронна мережа, трансформер, механізм уваги, класифікація текстів.

Табл. 0. Іл. 5. Бібліогр.: 23 назви.

UDC 004.93

Devising a combined method for analyzing the emotional coloring of texts / K. E. Petrov, I. P. Bokov, I. V. Kobzev. Management Information System and Devices. 2025. № 186. P. 5-16.

One of the key tasks of natural language processing (NLP) is to analyze the emotional coloring of text, which plays an important role in numerous application areas, including marketing, sociology, psychology, public opinion analysis and information security. Text sentiment analysis systems allow you to quickly obtain structured information about the emotional mood of society, predict reactions to certain events and identify potential threats or destructive content.

Despite significant progress in the field of NLP, existing methods for determining emotional coloring of text have a number of limitations that reduce their effectiveness. In particular, traditional methods often do not take into account the contextual meaning of words, which is critical for accurate emotional coloration recognition. In addition, some methods have difficulty analyzing polysemous words, sarcasm, irony, and slang expressions. Therefore, it is an urgent task to further improve methods for analyzing text tone, in particular by combining several methods and using deep learning models.

The aim of the study is to increase the accuracy of classifying the emotional coloring of natural

language texts by using lexical, statistical, and contextual methods that will allow taking into account both superficial lexical features and deep semantic connections in the text.

The combined method, which proposed in this paper combines statistical (TF-IDF) and contextual (BERT) vector representations of the text. This combination allows taking into account frequency patterns and deep semantic dependencies between words. The use of the Random Forest ensemble classifier allowed to build a robust model, that can effectively classify short English texts with a high level of accuracy.

Experimental results have shown, that the proposed combined method has a higher classification accuracy (89 %) of texts compared to the basic – TF-IDF + RF and BERT + RF (78 % and 82 % respectively).

The use of the combined method will increase the efficiency of context analysis and recognition of complex language structures, which makes it promising for analyzing public opinion in social networks, media, and chatbots; for use in customer support services; for determining the emotions of users of web services, websites and web applications.

Keywords: emotional coloring, natural language, lexical analysis, machine learning, deep neural network, transformer, attention mechanism, text classification.

Tab. 0. Fig. 5. Ref.: 23 items.

УДК 681.5:004.89

Модель прогнозування використання ресурсів у хмарних обчисленнях з використанням архітектури Informer / І. В. Михайліченко, О. С. Ляшенко. АСУ та прилади автоматики. 2025. № 186. С. 17-28.

Об'єктом дослідження є процес прогнозування та моніторингу навантаження хмарної інфраструктури. Визначено, що одним з ефективних способів його реалізації є застосування нейронних мереж трансформерного типу для обробки багатовимірних часових рядів телеметрії. Класичні архітектури, зокрема Informer, забезпечують високу точність прогнозування, але потребують значних обчислювальних ресурсів і тривалого навчання, що ускладнює інтеграцію у системи моніторингу реального часу.

Метою дослідження є розробка та експериментальна оцінка модифікованої архітектури Informer, оптимізованої для швидшого навчання та ефективнішого використання ресурсів при збереженні прийнятної точності. Запропонована модель враховує сезонні закономірності та реалізує ефективну самоувагу з позиційними й сезонними embedding-представленнями. Модель протестовано у системі із замкненим циклом МАРЕ (Моніторинг–Аналіз–Планування–Виконання), що дозволило оцінити її роботу в умовах автоматичного управління ресурсами Kubernetes.

Проведено порівняльний аналіз із класичною архітектурою Informer. Оптимізована модель забезпечила восьмиразове прискорення навчання та зменшення кількості параметрів у 10 разів при збереженні понад 88 % пояснюваності варіації даних і близької точності прогнозування метрик CPU, пам'яті, мережових і дискових операцій. Це робить її придатною для розгортання у промислових системах моніторингу та управління хмарними ресурсами в режимі реального часу.

Ключові слова: прогнозування навантаження, хмарна інфраструктура, багатовимірні часові ряди, Informer, трансформер, МАРЕ, оптимізація моделі, AIOps, Kubernetes.

Табл. 0. Іл. 7. Бібліогр.: 18 назв.

UDC 681.5:004.89

Resource usage forecasting model in cloud computing using Informer architecture / I. V. Mykhailichenko, O. S. Liashenko. Management Information System and Devices. 2025. No. 186. P. 17-28.

The object of the study is the process of forecasting and monitoring the load of cloud infrastructure. It is determined that one of the effective ways to implement this process is the use of transformer-based neural networks for processing multivariate telemetry time series. Classical architectures, in particular Informer, provide high forecasting accuracy but require significant computational resources and long training times, which complicates their integration into real-time monitoring systems.

The aim of the research is to develop and experimentally evaluate a modified Informer architecture optimized for faster training and more efficient resource usage while maintaining acceptable accuracy. The proposed model takes seasonal patterns into account and implements efficient self-attention with positional

and seasonal embedding representations. The model was tested in a system implementing a closed MAPE (Monitoring–Analysis–Planning–Execution) loop, which made it possible to assess its performance in automatic resource management for Kubernetes.

A comparative analysis with the classical Informer architecture was carried out. The optimized model achieved an eightfold acceleration in training and a tenfold reduction in the number of parameters, while maintaining over 88 % of the explained variance and comparable accuracy in forecasting CPU, memory, network, and disk metrics. This makes it suitable for deployment in industrial cloud resource monitoring and management systems in real-time operation.

Keywords: load forecasting, cloud infrastructure, multivariate time series, Informer, transformer, MAPE, model optimization, AIOps, Kubernetes.

Tab. 0. Fig. 7. Ref. 10 items.

УДК 004.358:681.518

Підтримка обмежень цілісності реляційної бази даних на етапах експлуатації та реінжинірингу у задачах інтелектуального аналізу / О. В. Золотухін, М. С. Кудрявцева, В. О. Філатов. АСУ та прилади автоматики. 2025. № 186. С. 29-39.

Об'єктом дослідження є реляційна база даних – база даних, заснована на реляційній моделі. Предметом дослідження є методи підтримки цілісності даних у реляційних системах, орієнтованих на інтелектуальний аналіз.

Розглядається задача виявлення нової інформації про взаємозв'язки між даними, які могли бути додані та внесені у процесі функціонування бази даних. Взаємозв'язки представляються у вигляді залежностей різних типів, які можна використовувати як вихідні дані для методів повторного проєктування (реінжинірингу) реляційної бази даних.

За підсумками аналізу особливостей проєктування інформаційних систем, які застосовують реляційну модель даних, розглянуто основні проблеми підтримки цілісності у разі багатозначних функціональних залежностей. Досліджено основні властивості багатозначних функціональних залежностей атрибутів та запропоновано метод підтримки цілісності засобами реляційних систем. Наведено низку прикладів, що пояснюють загальну проблему підтримки цілісності реляційних моделей даних, а також технологію специфічних модельних обмежень та метод вирішення поставленої задачі дослідження.

Отримані результати дозволяють в подальшому вирішити задачу розробки інформаційних систем орієнтованих на ефективну комплексну технологію – реінжиніринг реляційної бази даних в поєднанні з сучасними методами інтелектуального аналізу у застосунках користувача.

Ключові слова: реляційна модель, цілісність даних, функціональні залежності, теорія нормалізації, реінжиніринг, інформаційна система, база даних, інтелектуальний аналіз, ключові слова.

Табл. 0. Іл. 8. Бібліогр.: 14 назв.

UDC 004.358:681.518

Supporting relational database integrity constraints during operation and reengineering in business intelligence tasks / O. Zolotukhin, M. Kudryavtseva V. Filatov. Management Information System and Devices. 2025. № 186. P. 29-39.

The object of the study is a relational database - a database based on the relational model. The subject of the research is methods for maintaining data integrity in relational systems oriented towards intellectual analysis.

The task of identifying new information about the relationships between data that could be added and entered during the database operation is considered. Relationships are represented in the form of dependencies of various types, which can be used as input data for methods of redesign (reengineering) of a relational database.

Based on the analysis of the design features of information systems that use the relational data model, the main problems of maintaining integrity in the case of multivalued functional dependencies are considered. The main properties of multivalued functional dependencies of attributes are studied and a method of maintaining integrity using relational systems is proposed. A number of examples are given that

explain the general problem of maintaining the integrity of relational data models, as well as the technology of specific model constraints and the method of solving the research problem.

The results obtained allow us to further solve the problem of developing information systems focused on effective integrated technology - relational database reengineering in combination with modern methods of intelligent analysis in user applications.

Keywords: relational model, data integrity, functional dependencies, normalization theory, reengineering, information system, database, intelligent analysis, keywords.

Tab. 0. Fig. 8. Ref.: 14 items.

УДК 621.391:004.032.26

Моделювання та аналіз графових нейронних мереж для оптимізації маршрутизації в інфокомунікаційних мережах / С. В. Штангей, Л. І. Мельнікова, А. В. Марчук, О. В. Лінник, О. К. Соколов. АСУ та прилади автоматики. 2025. № 186. С. 40-54.

Об'єктом дослідження є процес побудови маршрутів у інфокомунікаційній мережі. Розроблено, реалізовано та експериментально досліджено модель маршрутизації на основі графової нейронної мережі з edge-level класифікацією, яка побудована на архітектурі GENConv.

Розглянуто архітектурні особливості графових нейронних мереж, зокрема механізми message passing, attention, агрегування та оновлення ознак. Проведено порівняння GCN, GAT і GENConv, обгрунтовано вибір останньої як базової архітектури для edge-level класифікації маршрутних ребер. Побудовано модель на основі GENConv з MLP-декодером. Проведено навчання на великій вибірці графів та оцінено її точність, середню затримку й відсоток успішно побудованих маршрутів.

Наведено порівняння з класичним алгоритмом за якістю рішень і часом виконання. Встановлено, що в режимі inference графова модель працює значно швидше, особливо на великих графах, і не потребує повторного перебору всього простору маршрутів при кожному запиті. Це робить запропонований підхід придатним до використання в реальному часі, у динамічних мережах, де швидкість прийняття рішень є критичною.

Ключові слова: графові нейронні мережі, маршрутизація, GENConv, edge-level класифікація, бази даних, інфокомунікаційні мережі, оптимізація, моделювання, Python, програмування.

Табл. 2. Іл. 11. Бібліогр.: 22 назви.

UDC 621.391:004.032.26

Modeling and analysis of graph neural networks for optimizing routing in infocommunication networks / S. V. Shtangey, L. I. Melnikova, A. V. Marchuk, O. V. Linnyk, O. K. Sokolov. Management Information System and Devices. 2025. № 186. P. 40-54.

The research object is the process of route construction in infocommunication networks. A routing model based on a graph neural network (GNN) with edge-level classification has been developed, implemented, and experimentally studied. The model uses the GENConv architecture.

Architectural features of GNNs are analyzed, including message passing, attention mechanisms, feature aggregation, and updating techniques. A comparison of GCN, GAT, and GENConv architectures is conducted, and GENConv is justified as the base architecture for edge-level classification of routing edges. A model is built using GENConv with an MLP decoder. The model is trained on a large set of graphs and evaluated in terms of accuracy, average latency, and the percentage of successfully constructed routes.

The proposed approach is compared with classical algorithms based on decision quality and execution time. It is established that in inference mode, the graph model works significantly faster – especially on large graphs – and does not require exhaustive route space reprocessing for each query. This makes the proposed method suitable for real-time use in dynamic networks, where decision-making speed is critical.

Keywords: Graph Neural Networks, routing, GENConv, edge-level classification, databases, infocommunication networks, optimization, modeling, Python, programming

Tab. 2. Fig. 11. Ref.: 22 items.

УДК 004.75

Використання методів машинного навчання для виявлення атак на блокчейн-системи / В. В. Просолов, Г. З. Халімов, П. В. Шулік, А. О. Смірнов, Д. О. В'юхін. АСУ та прилади автоматики. 2025. № 186. С. 55-70.

Предметом дослідження є методи виявлення атак у мережах із консенсусом Proof-of-Stake (PoS). Мета роботи – експериментальне дослідження та аналіз ефективності класичних алгоритмів машинного навчання для виявлення шкідливих вузлів у блокчейн-системах. Задачі: аналіз вразливостей технології блокчейн, створення та використання спеціалізованого набору даних для мереж PoS, а також побудова й тестування моделей машинного навчання. Основна увага приділяється порівнянню трьох алгоритмів – Random Forest, Support Vector Machine та k-Nearest Neighbors – з метою визначення їхньої придатності для моніторингу активності вузлів та виявлення аномалій. Для вирішення поставлених задач застосовано методи: моделювання, емпіричні та математичні методи. Моделювання полягає у програмній реалізації вибраних алгоритмів і подальшому аналізі їхніх результатів із використанням метрик точності, повноти, F1-міри та матриці плутанини. Емпіричні методи реалізовано шляхом тестування моделей на напівсинтетичному датасеті, який містить понад 10000 записів про вузли та транзакції блокчейну. Математичні методи передбачають обчислення статистичних показників ефективності, а також аналіз інформативності ознак, що визначають поведінку вузлів.

Досягнуті результати: було апробовано датасет для блокчейнів PoS, що включає ключові операційні параметри транзакцій і вузлів, сформовано пропозиції щодо подальшого використання моделей машинного навчання, здійснено тестування моделей машинного навчання

Висновки. Доведено, що машинне навчання є ефективним інструментом для ідентифікації аномалій та шкідливої активності у блокчейн-системах. Отримані результати закладають підґрунтя для подальших досліджень, які можуть бути спрямовані на розширення ознакового простору, інтеграцію глибоких нейронних мереж, розробку ансамблевих підходів та адаптацію методів до різних типів блокчейнів.

Ключові слова: блокчейн, Proof-of-Stake, атаки, машинне навчання, виявлення аномалій.

Табл. 3. Іл. 2. Бібліогр.: 24 назви.

UDC 004.75

Application of machine learning methods for detecting attacks on blockchain systems / V. V. Prosolov, H. Z. Khalimov, P. V. Shulik, A. O. Smirnov, D. O. Viukhin. Management Information System and Devices. 2025. № 186. P. 55-70.

The subject of the research is methods for detecting attacks in networks with the Proof-of-Stake (PoS) consensus mechanism. The purpose of this experimental investigation and analysis is to evaluate the effectiveness of classical machine learning algorithms for detecting malicious nodes in blockchain systems. The tasks include the analysis of blockchain technology vulnerabilities, the creation and use of a specialized dataset for PoS networks, as well as the construction and testing of machine learning models. The main focus is placed on comparing three algorithms – Random Forest, Support Vector Machine, and k-Nearest Neighbors – in order to determine their suitability for monitoring node activity and detecting anomalies. To solve the tasks set, the following methods were implemented: modeling, empirical, and mathematical approaches were applied. Modeling consisted of software implementation of the selected algorithms and subsequent analysis of their performance using accuracy, recall, F1-score metrics, and confusion matrices. Empirical methods were realized through testing the models on a partially synthetic dataset containing more than 10,000 records of blockchain nodes and transactions. Mathematical methods involved the calculation of statistical performance indicators and the analysis of feature importance that characterizes node behavior.

The achieved results include the validation of a dataset for PoS blockchains that incorporates key operational parameters of transactions and nodes, the development of recommendations for further use of machine learning models, and the testing of selected models.

Conclusions. The study demonstrated that machine learning is an effective tool for identifying anomalies and malicious activity in blockchain systems. The obtained results lay the foundation for further research, which may focus on expanding the feature space, integrating deep neural networks, developing

ensemble approaches, and adapting methods to different types of blockchains.

Keywords: blockchain, Proof-of-Stake, attacks, machine learning, anomaly detection.

Tab. 3. Fig. 2. Ref.: 24 items.

УДК 004.75:519.8

Метод створення датасетів для оцінки алгоритмів розподілу валідаторів на основі механізму Proof of Stake / Є. Є. Деменко, І. В. Гребеннік, М. М. Колмиков. АСУ та прилади автоматики. 2025. № 186. С. 71-82.

Досліджено проблему відтворюваності експериментів з оптимізації розподілу валідаторів у блокчейн-мережах із консенсусом Proof of Stake, зокрема через відсутність стандартизованих датасетів та уніфікованих методів тестування, що ускладнює об'єктивне порівняння алгоритмів. Для вирішення цієї проблеми запропоновано метод побудови тестових наборів даних, що базуються на детермінованих генераторах псевдовипадкових послідовностей та характеристиках валідаторів, налаштованих за статистикою мережі Ethereum.

Кожен валідатор описано набором параметрів, що включає розмір стейку з мінімальною вимогою відповідно до стандартів Ethereum, продуктивність із рівномірним розподілом, надійність у високому діапазоні, мережеві затримки залежно від географічної близькості учасників, географічне розташування згідно з фактичною статистикою розподілу валідаторів по регіонах, якість мережевого з'єднання та історію штрафів відповідно до статистики порушень у Beacon Chain. Було створено три набори даних різного масштабу для малих, середніх та великих конфігурацій мереж із фіксованими початковими значеннями генераторів для забезпечення повної відтворюваності експериментів.

Розроблено систему багатокритеріальної оцінки, засновану на узагальненому показнику якості, що максимізує пропускну здатність системи та мінімізує дисбаланс навантаження і мережеві затримки з науково обґрунтованими ваговими коефіцієнтами. Протокол десятикратного тестування забезпечує статистичну достовірність результатів і зменшує вплив випадковості на висновки.

В експериментах було проведено порівняльний аналіз чотирьох алгоритмів розподілу: гібридного метаевристичного методу на основі оптимізації роєм частинок із локальним пошуком, випадкового розподілу з корекцією, адаптованого механізму перетасування Ethereum та жадібного алгоритму. Результати експериментів показали масштабно-залежну ефективність алгоритмів: гібридний метод забезпечує високу якість оптимізації на всіх досліджуваних масштабах, проте квадратичне зростання часу виконання обмежує його застосування періодичним офлайн-плануванням конфігурації мережі; механізм перетасування демонструє стабільні результати середньої якості при швидкому виконанні; випадковий метод характеризується помірною швидкістю з варіативними результатами; жадібний алгоритм показує максимальну швидкість із детермінованими результатами, але змінну ефективність залежно від масштабу мережі.

Запропонований метод формує основу для стандартизації експериментальних досліджень у системах консенсусу Proof of Stake та забезпечує об'єктивне порівняння нових алгоритмічних рішень для розподілу валідаторів у децентралізованих блокчейн-мережах.

Ключові слова: блокчейн, Proof of Stake, валідатори, експериментальні датасети, відтворюваність, оптимізація, розподіл.

Табл. 0. Іл. 4. Бібліогр.: 22 назви.

UDC 004.75:519.8

Method for creating datasets to evaluate validator allocation algorithms based on the Proof of Stake mechanism / Y. Y. Demenko, I. V. Grebennik, M. M. Kolmykov. Management Information System and Devices. 2025. № 186. P. 71-82.

The problem of reproducibility of experiments in optimizing validator allocation in blockchain networks with Proof of Stake consensus was investigated, in particular due to the absence of standardized datasets and unified testing methods, which complicates the objective comparison of algorithms. To tackle this issue, we propose a method for building test datasets that rely on deterministic pseudorandom sequence generators and validator profiles calibrated against Ethereum network statistics.

Each validator is described by a set of parameters that includes the stake size with the minimum requirement according to Ethereum standards, performance with a uniform distribution, reliability in a high range, network delays depending on the geographical proximity of participants, geographical location according to the actual statistics of validator distribution by regions, quality of network connection, and slashing history according to the violation statistics in the Beacon Chain. Three datasets of different scales were created for small, medium, and large network configurations with fixed initial values of the generators to ensure full reproducibility of experiments.

A multi-criteria evaluation system was developed based on a generalized quality indicator that maximizes system throughput and minimizes load imbalance and network delays with scientifically grounded weighting coefficients. The tenfold testing protocol ensures the statistical reliability of results and reduces the impact of randomness on conclusions.

The experiments conducted a comparative analysis of four allocation algorithms: a hybrid metaheuristic method based on particle swarm optimization with local search, random allocation with correction, an adapted Ethereum shuffling mechanism, and a greedy algorithm. The experimental results revealed scale-dependent efficiency of the algorithms: the hybrid method provides high optimization quality at all investigated scales, but quadratic growth of execution time limits its application to periodic offline planning of network configuration; the shuffling mechanism demonstrates stable medium-quality results with fast execution; the random method is characterized by moderate speed with variable results; the greedy algorithm shows maximum speed with deterministic results but variable efficiency depending on the network scale.

The proposed method forms a basis for standardizing experimental research in Proof of Stake consensus systems. It ensures the objective comparison of new algorithmic solutions for validator allocation in decentralized blockchain networks.

Keywords: blockchain; Proof of Stake; validators; experimental datasets; reproducibility; optimization; distribution.

Tab. 0. Fig. 4. Ref.: 22 items.

УДК 004.932.2:681.7.06

Система на основі RGB-датчика для виявлення камуфляжу у військовому середовищі / Д. М. Крицький, Д. А. Оніщук, О. О. Валуженич. АСУ та прилади автоматики. 2025. № 186. С. 83-94.

Предметом дослідження є розробка недорогої оптичної системи на основі RGB-датчика для виявлення замаскованих об'єктів у військовому середовищі. Метою дослідження є створення прототипу апаратно-програмної системи виявлення камуфляжу на основі RGB-датчика, що працює у видимому спектрі світла, з подальшою експериментальною перевіркою його ефективності у різних природних умовах. Система відзначається низьким енергоспоживанням, мобільністю та може бути інтегрована у портативні або безпілотні розвідувальні платформи. Для досягнення поставленої мети було вирішено такі задачі: розробка апаратної структури сенсорної системи на основі доступних компонентів; реалізація алгоритмів моделювання фону, виявлення кольорових аномалій та фільтрації шумів з використанням колориметричного аналізу у просторі RGB; проведення серії натурних експериментів у різних природних умовах – лісистій місцевості, степовій зоні та кам'янистому ландшафті; оцінка точності, стабільності та обмежень запропонованого рішення у порівнянні з тепловізорами та інфрачервоними камерами зі штучним інтелектом. Методологія базується на використанні RGB-датчика TCS34725 із вбудованим інфрачервоним фільтром та 16-бітним АЦП у поєднанні з мікроконтролером ESP32, який забезпечує обробку даних у реальному часі, бездротову передачу інформації та автономну роботу. Алгоритм виявлення ґрунтується на формуванні кольорового профілю фону, аналізі відхилень за евклідовою відстанню, нормалізації RGB-значень, медіанній фільтрації та адаптивному пороговому визначенні, що забезпечує стійкість до змін середовища. Результати випробувань показали, що запропонована система здатна з високою точністю ідентифікувати цифровий камуфляж «піксель» у лісистій місцевості (86 %), камуфляж «мультикам» у степових умовах (78 %) та однотонний оливковий камуфляж на кам'янистому фоні (65 %). Порівняння з тепловізійними та інфрачервоними системами підтвердило значні переваги RGB-рішення за вартістю (менше \$ 20),

енергоефективністю та мобільністю, хоча функціонування можливе лише у денний час. У висновках наголошується на доцільності застосування RGB-систем виявлення як економічно вигідного допоміжного інструменту для розвідки та охоронних завдань. Наукова новизна дослідження полягає в інтеграції простого й доступного RGB-сенсора з адаптивними алгоритмами моделювання фону та аналізу кольорових відхилень, що дозволяє створювати малопотужні платформи для виявлення камуфляжу. На відміну від традиційних підходів, запропонована система відкриває перспективи розгортання розподіленої мережі дешевих сенсорних вузлів або інтеграції у безпілотні літальні апарати для оперативного моніторингу великих територій.

Ключові слова: RGB-датчик, виявлення камуфляжу, військове застосування, низьковартісна система, аналіз кольорових аномалій, оптичний сенсор, ESP32, натурні випробування.

Табл. 2. Іл. 0. Бібліогр.: 22 назви.

UDC 004.932.2:681.7.06

System based on RGB sensor for detecting camouflage in military environments / D. Krytskyi, D. Onishchuk, O. Valiuzhenych. Management Information System and Devices. 2025. № 186. P. 83-94.

The subject of the article is the development of an inexpensive optical system based on an RGB sensor for de-tecting camouflaged objects in a military environment. The aim of the research is to create a prototype hardware and software system for detecting camouflage based on an RGB sensor operating in the visible light spectrum, with subsequent experimental verification of its effectiveness in various natural conditions. The system is characterized by low power consumption, mobility, and can be integrated into portable or un-manned reconnaissance platforms. To achieve the set goal, the following tasks were solved: development of the hardware structure of the sensor system based on available components; implementation of algorithms for background modeling, color anomaly detection, and noise filtering using colorimetric analysis in RGB space; conducting a series of field experiments in various natural conditions – wooded areas, steppe zones, and rocky landscapes; evaluating the accuracy, stability, and limitations of the proposed solution in comparison with thermal imagers and infrared cameras with artificial intelligence. The methodology is based on the use of a TCS34725 RGB sensor with a built-in infrared filter and a 16-bit ADC in combination with an ESP32 microcontroller, which provides real-time data processing, wireless data transmission, and autonomous operation. The detection algorithm is based on the formation of a color background profile, analysis of deviations by Euclidean distance, normalization of RGB values, median filtering, and adaptive threshold determination, which ensures resistance to environmental changes. Test results showed that the proposed system is capable of accurately identifying digital «pixel» camouflage in wooded areas (86 %), multicam camouflage in steppe conditions (78 %), and plain olive camouflage on rocky backgrounds (65 %). A comparison with thermal imaging and IR systems confirmed the significant advantages of the RGB solution in terms of cost (less than \$ 20), energy efficiency, and mobility, although it can only operate during daylight hours. The conclusions emphasize the feasibility of using RGB detection systems as a cost-effective auxiliary tool for reconnaissance and security tasks. The scientific novelty of the work lies in the integration of a simple and affordable RGB sensor with adaptive algorithms for background modeling and color deviation analysis, which allows the creation of low power platforms for camouflage detection. Unlike traditional approaches, the proposed system opens up prospects for deploying a distributed network of inexpensive sensor nodes or integrating them into unmanned aerial vehicles for operational monitoring of large areas.

Keywords: RGB sensor, camouflage detection, military applications, low-cost system, color anomaly analysis, optical sensor, ESP32, field testing.

Tab. 2. Fig. 0. Ref.: 22 items.

УДК 004.8:004.9

Метод оцінки негативних аспектів рішення інтелектуальної системи в задачах побудови пояснень / С. Ф. Чалий, І. О. Лещинська. АСУ і прилади автоматики. 2025. № 186. С. 95-102.

Предметом дослідження є процес оцінки негативних аспектів рішень інтелектуальних систем при формуванні користувацьких ментальних моделей для створення збалансованих пояснень. Метою є розробка методу до оцінки негативних аспектів рішень інтелектуальних систем з тим, щоб забезпечити формування збалансованих пояснень на основі інтеграції

суттєвих негативних характеристик у ментальні моделі користувачів. Задачі: розробити гібридний підхід до побудови ментальної моделі, який поєднує аналіз важливості та контекстуальний аналіз релевантності рішень з метою комплексної оцінки негативних властивостей вибраного користувачем рішення; розробити метод оцінки негативних аспектів користувацького рішення при побудові ментальної моделі користувача для автоматизованого виявлення та упорядкування негативних характеристик цього рішення. Використано принципи побудови ментальних моделей користувачів інтелектуальної системи, метод SHAP-аналізу, методи сентимент-аналізу, методи контекстного аналізу. Отримано такі результати. Запропоновано гібридний підхід до побудови ментальної моделі на основі комбінованої оцінки негативних властивостей рішення інтелектуальної системи. Розроблено метод оцінки негативних аспектів рішення інтелектуальної системи. Висновки. Наукова новизна отриманих результатів полягає в такому. Запропоновано метод оцінки негативних аспектів рішення інтелектуальної системи, який містить етапи попередньої обробки вхідних неструктурованих даних, SHAP-аналізу важливості негативних властивостей, контекстного аналізу негативних властивостей та формування комбінованої оцінки рішення. Метод забезпечує можливість обґрунтованого вибору рішення користувачем з урахуванням його ключових негативних властивостей та створює умови для побудови зрозумілих пояснень і підвищення довіри користувачів.

Ключові слова: негативні аспекти рішення, ментальна модель, інтелектуальна система, пояснювальний штучний інтелект, SHAP-аналіз, контекстний аналіз, гібридний підхід.

Табл. 1. Іл. 0. Бібліогр.: 23 назви.

UDK 004.8:004.9

Method for evaluating negative aspects of intelligent system solutions in explanation generation tasks / S. F. Chaly, I. O. Leshchynska. Management Information System and Devices. 2025. № 186. P. 95-102.

The subject of research is the process of evaluating negative aspects of intelligent system solutions when forming user mental models to create balanced explanations. The objective is to develop a method for evaluating negative aspects of intelligent system solutions to ensure the formation of balanced explanations based on the integration of significant negative characteristics into user mental models. Tasks: develop a hybrid approach to mental model construction that combines importance analysis and contextual analysis of solution relevance for comprehensive evaluation of negative properties of user-selected solutions; develop a method for evaluating negative aspects of user solutions in mental model construction for automated identification and organization of negative characteristics of these solutions. The following methods were used: principles of user mental model construction for intelligent systems, SHAP analysis method, sentiment analysis methods, contextual analysis methods. The following results were obtained. A hybrid approach to mental model construction based on combined evaluation of negative properties of intelligent system solutions is proposed. A method for evaluating negative aspects of intelligent system solutions is developed. Conclusions. The scientific novelty of the results obtained consists of the following. A method for evaluating negative aspects of intelligent system solutions is proposed, which includes stages of preprocessing input unstructured data, SHAP analysis of negative property importance, contextual analysis of negative properties, and formation of combined solution evaluation. The method provides the possibility for informed solution selection by users considering key negative properties and creates conditions for building understandable explanations and increasing user trust.

Keywords: negative aspects of solution, mental model, intelligent system, explainable artificial intelligence, SHAP analysis, contextual analysis, hybrid approach.

Tables 1. Figs. 0. Refs.: 23 titles.

УДК 004.8:004.9

Прескриптивна модель темпоральних знань з OWL-інтеграцією для підтримки рішень в інформаційних системах / О. В. Чала, О. М. Бітченко, Л. Ф. Сайківська, М. Є. Алфьоров, Д. Г. Ганшин. АСУ і прилади автоматики. 2025. № 186. С. 103-112.

Об'єктом дослідження є процес формування рекомендацій в інформаційних системах на основі інтеграції темпоральних правил із стандартизованими онтологічними конструкціями для переходу від дескриптивного до прескриптивного опису в задачах підтримки управлінських рішень. Метою є розробка прескриптивної моделі темпоральних знань з інтеграцією OWL-онтологій з тим, щоб забезпечити автоматизоване формування обґрунтованих управлінських рекомендацій на основі семантично узгоджених темпоральних залежностей у інформаційних системах. Завдання: розробити представлення темпоральних знань на основі інтеграції темпоральних правил та стандартизованих OWL-Time конструкцій; розробити прескриптивну модель темпоральних знань з можливостями семантичного відображення темпоральних правил на онтологічні відношення. Використано принципи темпоральної логіки Аллена, стандарти W3C OWL-Time, методи онтологічного моделювання, алгоритми логічного виведення. Отримано такі результати. Розроблено дворівневе представлення темпоральних знань, що поєднує емпіричний рівень темпоральних правил із семантичним рівнем онтологічних відношень через систему комбінованих ваг. Створено прескриптивну модель темпоральних знань, що включає множину темпоральних правил, множину онтологічних відношень та множину комбінованих ваг для автоматизованого формування управлінських рекомендацій. Висновки. Наукова новизна отриманих результатів полягає в такому. Запропоновано прескриптивну модель темпоральних знань з OWL-інтеграцією, яка містить множини темпоральних правил, онтологій, а також комбінованих ваг, що визначають вплив цих правил та онтологій на формування рекомендацій щодо дій управлінського рішення. Модель дає можливість формувати рекомендації щодо майбутніх дій управлінського рішення на основі композиції правил з перевіркою відповідності цих правил існуючим знанням щодо предметної області.

Ключові слова: прескриптивна модель, темпоральні знання, OWL-онтології, системи підтримки рішень, управлінські рекомендації, семантична інтеграція, двошарова архітектура, комбіновані ваги.

Табл. 1. Іл. 5. Бібліогр.: 23 назви.

UDC 004.8:004.9

Prescriptive model of temporal knowledge with OWL integration for decision support in information systems / O. V. Chala, O. M. Bitchenko, L. F. Saykivska, M. E. Alferov, D. G. Ganshin. Management Information System and Devices. 2025. № 186. P. 103-112.

The object of research is the process of recommendation formation in information systems based on integration of temporal rules with standardized ontological constructions for transition from descriptive to prescriptive description in tasks of supporting management decisions. The objective is to develop a prescriptive model of temporal knowledge with integration of OWL-ontologies in order to ensure automated formation of justified management recommendations based on semantically coordinated temporal dependencies in information systems. Tasks: to develop representation of temporal knowledge based on integration of temporal rules and standardized OWL-Time constructions; to develop a prescriptive model of temporal knowledge with capabilities of semantic mapping of temporal rules onto ontological relations. Used principles of Allen's temporal logic, W3C OWL-Time standards, methods of ontological modeling, algorithms of logical inference. The following results were obtained. A two-level representation of temporal knowledge was developed that combines the empirical level of temporal rules with the semantic level of ontological relations through a system of combined weights. A prescriptive model of temporal knowledge was created that includes a set of temporal rules, a set of ontological relations and a set of combined weights for automated formation of management recommendations. Conclusions. The scientific novelty of the obtained results consists in the following. A prescriptive model of temporal knowledge with OWL-integration is proposed, which contains sets of temporal rules, ontologies, as well as combined weights that determine the influence of these rules and ontologies on the formation of recommendations regarding actions of management decision. The model provides possibility to form recommendations

regarding future actions of management decision based on composition of rules with verification of correspondence of these rules to existing knowledge regarding the subject domain.

Keywords: prescriptive model, temporal knowledge, OWL ontologies, decision support systems, management recommendations, semantic integration, two-layer architecture, combined weights.

Tables 1. Figures 5. References: 23 titles.

**ПРАВИЛА ОФОРМЛЕННЯ СТАТЕЙ
У ВСЕУКРАЇНСЬКОМУ МІЖВІДОМЧОМУ НАУКОВО-ТЕХНІЧНОМУ
ЗБІРНИКУ
«АВТОМАТИЗОВАНІ СИСТЕМИ УПРАВЛІННЯ ТА ПРИЛАДИ АВТОМАТИКИ»**

1. Загальні вимоги

До розгляду приймаються раніше не опубліковані статті українською та англійською мовами. Статті англійською мовою подаються разом з українськомовним варіантом. Статті, перекладені англійською за допомогою комп'ютерного перекладача та не відредаговані належним чином, не розглядаються.

Наукова стаття, яка подається до розгляду, має бути структурована та містити всі основні частини, характерні для наукової статті:

- постановка проблеми у загальному вигляді та її зв'язок з важливими науковими та практичними задачами;
- аналіз останніх досліджень та публікацій, у яких розпочато вирішення даної проблеми та на які спирається автор, виділення невирішених раніше частин загальної проблеми;
- формулювання цілей статті (постановка задачі);
- подання основного матеріалу досліджень з повним обґрунтуванням отриманих результатів;
- висновки даного дослідження та перспективи подальших досліджень у даному напрямку;
- перелік посилань (References).

2. Вимоги до структури рукопису

Структурно матеріали статті поділяються на такі елементи:

- УДК;
- прізвища та ініціали авторів статті;
- заголовок статті;
- анотація до статті;
- основний текст статті;
- перелік посилань;
- дата надходження статті до редколегії збірника;
- відомості про авторів статті;
- реферати українською та англійською мовами.

Бажаний порядок та зміст розділів основного тексту статті:

а) розділ 1 «Вступ», в якому визначається проблема у загальному вигляді та її зв'язок з важливими науковими та практичними задачами;

б) розділ 2 «Аналіз сучасних наукових публікацій та визначення проблеми дослідження», в якому наводяться результати аналізу останніх досліджень та публікацій, де розпочато вирішення даної проблеми та на які спирається автор, виділяються невирішені раніше частини загальної проблеми дослідження та конкретизується головна проблема дослідження у даній статті;

в) розділ 3 «Мета і задачі дослідження», в якому наводяться описи мети дослідження та задач дослідження, вирішення яких дозволяє досягти визначеної раніше мети дослідження;

г) розділ 4 «Матеріали і методи дослідження», в якому наводяться описи формального апарату та раніше проведених експериментальних досліджень, які будуть використані у

подальшому тексті статті;

д) розділ 5 «Результаті дослідження», в якому структуровано наводяться результати вирішення сформульованих у розділі 3 окремих задач дослідження (теоретичних та експериментальних);

е) розділ 6 «Обговорення результатів дослідження», в якому наводяться: опис особливостей отриманих результатів дослідження та їхньої відмінності від результатів попередніх досліджень у відповідній галузі; опис переваг отриманих результатів перед існуючими; опис недоліків і обмежень, які утруднюють використання отриманих результатів дослідження; опис подальших перспектив проведення досліджень за цим напрямом;

ж) розділ 7 «Висновки», в якому наводяться стислі описи отриманих результатів вирішення окремих задач дослідження та загальний висновок про досягнення поставленої у розділі 3 мети дослідження.

Заголовки окремих розділів основного тексту статті можуть змінюватися відповідно до змісту конкретної статті.

Розділи основного тексту статті, перелік посилань, дата надходження статті до редколегії збірника та відомості про авторів статті відокремлюються один від одного одним порожнім рядком.

3. Вимоги до оформлення рукопису

До розгляду приймаються матеріали статей обсягом не менше 5 повних сторінок (з урахуванням рисунків і таблиць).

Матеріали статті повинні бути набраними у редакторі MS Word. Припустимі формати файлу з матеріалами статті – .doc або .docx.

Формат сторінки – А4 (210х297 мм). Поля знизу, зверху, справа, зліва – 3 см.

Основний текст статті набирається шрифтом Times New Roman, кегль 11, інтервал – 1,1, абзацний відступ – 8 мм, інтервали перед і після – 0 мм, вирівнювання по ширині.

Для УДК – шрифт Times New Roman, кегль 11, інтервал – 1,1, абзацний відступ – 8 мм, інтервал перед – 0 мм, інтервал після – 6 мм, вирівнювання по ширині.

Для прізвищ та ініціалів авторів статті – шрифт Times New Roman, кегль 11, інтервал – 1,1, абзацний відступ – 8 мм, інтервали перед і після – 6 мм, вирівнювання по ширині.

Для заголовка статті – шрифт Times New Roman, кегль 11, напівжирний, інтервал – 1,1, абзацний відступ – 8 мм, інтервали перед і після – 6 мм, вирівнювання по ширині.

Для анотації – шрифт Times New Roman, кегль 10, інтервал – 1,1, відступ зліва – 0,8 см, абзацний відступ – 8 мм, інтервал перед – 6 мм, інтервал після – 0 мм, вирівнювання по ширині.

Для заголовків таблиць – шрифт Times New Roman, кегль 10, інтервал – 1,1, абзацного відступу немає, інтервали перед і після – 0 мм, слово «Таблиця» та її номер – з вирівнюванням вправо, назва таблиці (якщо вона є) – з вирівнюванням по центру.

Для підписуваних підписів – шрифт Times New Roman, кегль 10, інтервал – 1,1, абзацного відступу немає, інтервали перед і після – 0 мм, вирівнювання по центру.

Для переліку посилань та відомостей про авторів – шрифт Times New Roman, кегль 9, інтервал – 1,1, абзацний відступ – 8 мм, інтервали перед і після – 0 мм, вирівнювання по ширині.

Для рефератів – шрифт Times New Roman, кегль 10, інтервал – 1,1, абзацний відступ – 8 мм, інтервали перед і після – 0 мм, вирівнювання по ширині.

Формули набираються у редакторі формул Microsoft Equation або MathType, розташовуються у центрі робочого поля, нумерація – з правої сторони поля. Для цього

необхідно весь рядок розташувати справа, а потім вирівняти формулу табуляціями так, щоб вона розташовувалася по центру. Відступ зверху і знизу – по 6 пунктів. Нумерація формул усередині кожної статті наскрізна.

Формули, а також їхні складові, присутні у тексті, набираються з такими параметрами (див. рис. 1).

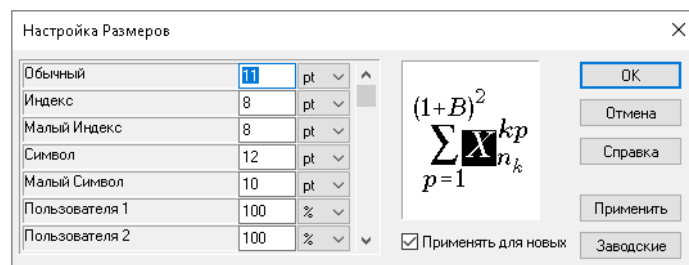


Рис. 1. Параметры настраивания размеров редактора формул MathType

Кожна таблиця виконується та розташовується в тексті одразу після посилання на неї. Усі таблиці у статті обов'язково нумеруються, незважаючи на їх кількість. Таблиця відокремлюється від попереднього та наступного тексту (таблиці, рисунку тощо) одним порожнім рядком.

Дані всієї таблиці набираються шрифтом розміром 10 пунктів, розміщуються по центру; у випадках, коли необхідно показати розрядність, – вирівнювання за знаком. Товщина сітки таблиці – 1 пункт. Приклад оформлення таблиці наведено на рис. 2.

Таблица 1

Множина описів сутностей функціональної задачі

ID	Найменування
1	Academic_load
2	Academic
3	Department
4	Individual_plan
5	Academic_section

Рис. 2. Приклад оформлення таблиці у тексті статті.

Бажаючи таблицю зі сторінки на сторінку не переносити. Якщо таблиця не може розміститися на сторінці, її поділяють на частини. У кожній частині таблиці повторюють її головку та боковик або заміняють їх відповідно номерами колонок або рядків, нумеруючи їх арабськими цифрами на першій частині таблиці. Слово «Таблиця» подається лише над першою її частиною. Над наступними її частинами праворуч друкується: «Продовження таблиці», а на останній – «Кінець таблиці», в усіх випадках вказується номер таблиці.

Кожен рисунок виконується та розташовується в тексті одразу після посилання на нього. Усі рисунки в статті обов'язково нумеруються, незважаючи на їх кількість. Необхідно вставляти рисунки у текст як графічні об'єкти (файли з розширенням .bmp, .jpg, .tiff чи .png, якість не менше 300 dpi), об'єкти MS Word або MS Visio.

Рисунок відокремлюється від попереднього та наступного тексту (таблиці, рисунку тощо) одним порожнім рядком.

Кожен рисунок повинен мати підпис, в якому вказується номер та, у

випадку необхідності, назва рисунку. Якщо рисунок займає менше 50 % ширини робочого поля, то можна зробити обтікання рисунку текстом, розташувавши його ліворуч або праворуч від робочого поля. Приклад рисунку з підписом наведений на рис. 3.

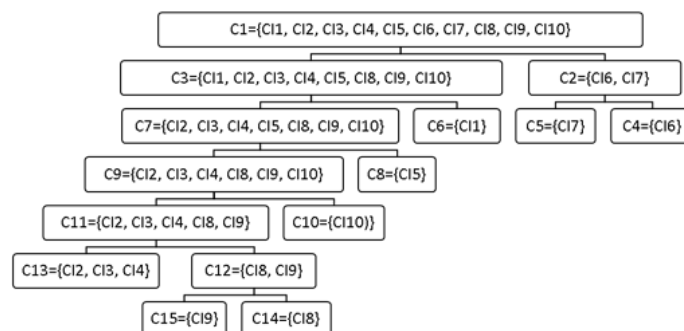


Рис. 3. Приклад виконання рисунку та підписового підпису

Посилання на літературні та електронні джерела у тексті статті позначаються у квадратних дужках [1]. До переліку посилань включаються тільки ті роботи, на які посилається автор статті. Посилання на неопубліковані роботи не допускаються.

Для оформлення переліку посилань слід використовувати один з таких шаблонів:

а) шаблон IEEE (автоматичне оформлення за шаблоном IEEE <https://www.citethisforme.com/ieee/source-type>);

б) положення ДСТУ 8302:2015 «Інформація та документація. Бібліографічне посилання. Загальні положення та правила складання» та ДСТУ 3582:2013 «Інформація та документація. Бібліографічний опис. Скорочення слів і словосполучень українською мовою. Загальні вимоги та правила».

Кожен з цих шаблонів слід використовувати для оформлення усіх елементів переліку посилань. Використання двох шаблонів для оформлення одного й того ж переліку посилань неприпустимо.

Кожне посилання у переліку посилань наводиться за порядком появи цих посилань у тексті статті.

У переліку посилань бажано використовувати посилання на сучасні публікації, вік яких не перевищує п'яти років у момент подачі статті до редакції. Крім того, під час формування переліку посилань статті необхідно дотримуватися такого розподілу: самоцититування – до 20 %, цитування зарубіжних публікацій – не менше 50%.

Відомості про авторів слід наводити українською та англійською мовами. У відомості про авторів слід включати: повні прізвища, ім'я та по-батькові; вчений ступінь (за наявності); вчене звання (за наявності); посаду; країну, місто; e-mail (вкрай бажано вказувати корпоративний e-mail, можна вказувати кілька e-mail, на які ви бажаєте отримувати повідомлення від редакції та читачів, які можуть зацікавитися вашою статтею); ORCID.

Реферат набирається українською та англійською мовами. Реферат повинен бути змістовним, дотримуватися логіки опису результатів у статті та давати можливість встановити її основний зміст. Реферат не повинен містити формул та рисунків. Необхідні символи в рефераті необхідно додавати через функцію вставки символів.

Реферат містить: УДК, назву статті (напівжирним шрифтом), ініціали та прізвища авторів (курсивом), текст (не менше 1800 друкованих знаків з пробілами та ключовими словами), ключові слова, кількість таблиць, рисунків та посилань у статті.

Ключові слова повинні містити до 10 слів, а не словосполучень, без використання аббревіатур, в іменному відмінку, розділятися крапкою з комою.

Реферати надаються до редколегії разом із статтею у вигляді окремого файлу.

4. Правила надсилання статей та подальшої взаємодії з редакційною колегією збірника

До редколегії збірника «АСУ та прилади автоматики» слід надсилати такі матеріали:

- файл у форматі .doc або .docx з текстом статті українською мовою;
- файл у форматі .doc або .docx з текстом статті англійською мовою (якщо автори бажають опублікувати статтю у збірнику англійською мовою);

- файл (у форматі .doc або .docx з текстами рефератів статті українською та англійською мовами;

- відскановану копію експертного висновку з дозволом опублікувати матеріали статті у відкритому друку. В разі потреби експертні висновки для авторів – співробітників (студентів, аспірантів тощо) ХНУРЕ можуть оформлюватися редколегією централізовано.

Матеріали статей надсилати електронною поштою – за адресою asu.devices@gmail.com.

Кожна надіслана в редакцію стаття після проходження рецензування і при позитивному рішенні редколегії буде надрукована в найближчому випуску збірника. Для цього авторам від імені редколегії надсилається ліцензійний договір, який закріплює право першої публікації статті у збірнику «АСУ та прилади автоматики». Автори статті повинні підписати цей ліцензійний договір та завірити свої підписи печаткою організації, в якій вони працюють. Підписаний ліцензійний договір автори статті надсилають на адресу редколегії збірника.

Відповідальний випусковий В.М. Левикін
Редактор О.Є. Неумивакіна
Комп'ютерна верстка М.В. Євланов, О.Є. Неумивакіна
Дизайн обкладинки номера за участю Є. Чех

Підп. до друку 29.09.2025. Формат 60х841/8. Умов. друк. арк.
Обл.-вид. арк. 15,34. Тираж 300 прим.
Зам. № 144. Ціна договірна.

Харківський національний університет радіоелектроніки (ХНУРЕ).
Україна, 61166, Харків, пр. Науки, 14

Оригінал-макет підготовлено у редакційно-видавничому відділі ХНУРЕ,
Україна, 61166, Харків, пр. Науки, 14

Збірник віддруковано в ТОВ «ДРУКАРНЯ МАДРИД»
61024, м. Харків, вул. Гуданова, 18
Тел.: +38(057)7565325
Свідоцтво суб'єкта видавничої справи
Серія ДК № 4399 від 27.08.2012 р.
www.madrid.in.ua e-mail:info@madrid.in.ua