

Міністерство освіти і науки України
Чорноморський національний університет ім. Петра Могили
Факультет комп'ютерних наук
Кафедра інтелектуальних інформаційних систем



Інтелектуальні інформаційні системи

*Матеріали всеукраїнської науково-практичної
конференції молодих вчених, аспірантів і
студентів*

ТЕЗИ

4 – 5 грудня 2025 року

Миколаїв

Інтелектуальні інформаційні системи : матеріали всеукр. наук.-практ. конф. молодих вчених, аспірантів, студентів : 4–5 грудня 2025 р., м. Миколаїв : тези / М-во освіти і науки України ; ЧНУ ім. Петра Могили ; Ф-т комп. наук ; Каф. інтелект. інформ. систем. – Миколаїв : Вид-во ЧНУ ім. Петра Могили, 2025. – 180 с.

Голова оргкомітету:

КОНДРАТЕНКО Юрій, доктор технічних наук, професор, Заслужений винахідник України, кафедра Інтелектуальних інформаційних систем ЧНУ імені П. Могили.

Співголова оргкомітету:

ГОЖИЙ Олександр, доктор технічних наук, професор, кафедра Інтелектуальних інформаційних систем ЧНУ імені П. Могили.

ОРГАНІЗАЦІЙНИЙ КОМІТЕТ

БОЙКО Анжела, кандидат технічних наук, доцент, декан факультету комп'ютерних наук ЧНУ імені П. Могили.

ЖУРАВСЬКА Ірина, доктор технічних наук, професор, завідувач кафедри Комп'ютерної інженерії ЧНУ імені П. Могили.

КАЛІНІНА Ірина, доктор технічних наук, професор, кафедра Інтелектуальних інформаційних систем ЧНУ імені П. Могили.

КОЗЛОВ Олексій, доктор технічних наук, професор, кафедра Інтелектуальних інформаційних систем ЧНУ імені П. Могили.

ДАВИДЕНКО Євгеній, кандидат технічних наук, доцент, завідувач кафедри інженерії програмного забезпечення ЧНУ імені П. Могили.

МАЛАХОВ Євгеній доктор технічних наук, професор, завідувач кафедри математичного забезпечення комп'ютерних систем ОНУ імені І. Мечнікова.

СІДЕНКО Євгеній, кандидат технічних наук, доцент, завідувач кафедри Інтелектуальних інформаційних систем ЧНУ імені П. Могили.

ТРУНОВ Олександр, доктор технічних наук, професор, кафедра Автоматизації та комп'ютерно інтегрованих технологій ЧНУ імені П. Могили.

ВІДПОВІДАЛЬНИЙ СЕКРЕТАР

ДИМО Валерій

dymovalery@gmail.com

Телеграм: @dymovalerii

ТЕХНІЧНИЙ СЕКРЕТАР

САЛЮТІН Максим

hayakawa010101@gmail.com

ІНТЕЛЕКТУАЛЬНІ ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ

УДК 004.65+004.738

*Антипова П.Б., Болюбаш Н.М.
Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

АНАЛІЗ ДИНАМІКИ ТА СТАБІЛЬНОСТІ ЕЛЕКТРОННОГО ВРЯДУВАННЯ КРАЇН СВІТУ ІЗ ВИКОРИСТАННЯМ АЛГОРИТМУ FUZZY C-MEANS

Електронне врядування (англ. e-Government) є одним із ключових елементів цифрової трансформації державного управління, забезпечуючи підвищення ефективності, прозорості та інклюзивності державно-громадських послуг. Це обумовлює потребу у виявленні факторів, які впливають на рівень розвитку цифрових державних сервісів та динаміку їх змін у часі. Ефективним інструментом, який дозволяє виявити закономірності між показниками розвитку електронного врядування та механізми їх підвищення, є групування країн із використанням алгоритмів кластерного аналізу й оцінка стабільності кластерів у динаміці.

Такий аналіз було проведено із використанням алгоритму кластеризації Fuzzy C-means, який дозволяє визначити ступінь належності кожної країни до кількох кластерів одночасно, що є доцільним при аналізі складних соціально-економічних систем. Для аналізу використано відкриті статистичні дані ООН за 2003–2024 роки [1]. Дослідження проведено на наборі країн із даними загального індексу розвитку електронного врядування EGDI (англ. E-Government Development Index). Для кожної країни було враховано складові індексу: OSI (англ. Online Services Index) – індекс онлайн-послуг; ТІІ (англ. Telecommunications Infrastructure Index) – індекс телекомунікаційної інфраструктури; НСІ (англ. Human Capital Index) – індекс людського капіталу.

Нечітка кластеризація виконана по рокам із відстеженням переходів країн між кластерами. Динамічний аналіз рівня розвитку електронного урядування передбачає визначення належності країни до певного кластера у наступні роки та зміни структури кластерів у часі [2]. Аналіз

стабільності виконувався за матрицями центрів ваги кластерів $V_t = [v_j(t)]$ для кожного року t , де $v_j(t)$ – центр ваги j -го кластера. Розрахунок V_t реалізовано з використанням матриць належностей $U_t = [u_{ij}(t)]$ розміром $k \times n$ (k – кількість кластерів, n – кількість країн), де $u_{ij}(t)$ – коефіцієнти належності, рівні ймовірності, з якою j -та країна відноситься до i -го кластера.

Зміна належності кожної країни i до кластерів між двома періодами розраховувалась за формулою:

$$\Delta_i = \frac{1}{k} \sum_{j=1}^k \left| u_{ij}(t+1) - u_{ij}(t) \right|. \quad (1)$$

Оцінювання стабільності країн у кластерах здійснювалося наступним чином: якщо Δ_i близьке до 0 $\rightarrow$ країна стабільна (її позиція у кластері не змінюється); якщо Δ_i велике $\rightarrow$ країна переходить до іншого кластера або втрачає чітку належність.

Індекс глобальної стабільності кластерної структури визначався за формулою:

$$S_{global} = 1 - \frac{1}{n} \sum_{j=1}^n \Delta_i. \quad (2)$$

Значення $S_{global} \in [0,1]$: близькі до 1 відображають стабільну кластерну структуру, а близькі до 0 свідчать про суттєві зміни у кластеризації.

Запропонований підхід до динамічної нечіткої кластеризації країн за рівнем розвитку e-Government дозволив врахувати перехідні стани між рівнями цифрового розвитку та оцінити стабільність кластерної структури у часі (рис. 1). Аналіз динаміки зміни індексу глобальної стабільності кластерів показав достатньо високі значення кластерної стабільності по рокам. Зниження індексу спостерігалось у період фінансової кризи (2005-2008 роки). На початку пандемії COVID також відбулося спочатку падіння, а потім зростання індексу, що було обумовлено прискоренням цифровізації в умовах карантину.

За результатами проведеного дослідження виділено три стійкі групи країн, віднесені до кластерів із різними рівнями розвитку електронного врядування: Lagging – низький рівень (країни Африки та частини Азії, EGDI < 0,4); Developing – середній рівень (країни Центральної та Східної Європи, EGDI 0,4-0,7); Leaders – високий рівень (Країни Північної Європи, Південна Корея, Сінгапур включно з Україною, EGDI > 0,75).

На прикладі України аналіз стабільності виявив зростання індексу змін належності до кластерів у періоди уповільнення та підвищення рівня розвитку е-Government (рис. 2). У 2004 році складова OSI демонструвала регрес, зниження EGDI відбувалося на фоні політичної нестабільності та слабкої інтеграції державних ІТ-систем. Різке зростання індексу у період 2018–2020 років корелює зі стрибком показника OSI, суттєвим був також вплив показника ТП. Що підтверджує ефективність державної стратегії цифровізації, в результаті якої країна здійснила перехід до кластера лідерів.

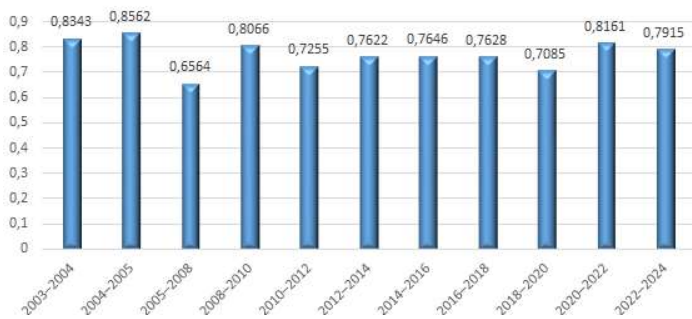


Рисунок 1 – Динаміка зміни індексу глобальної стабільності кластерів

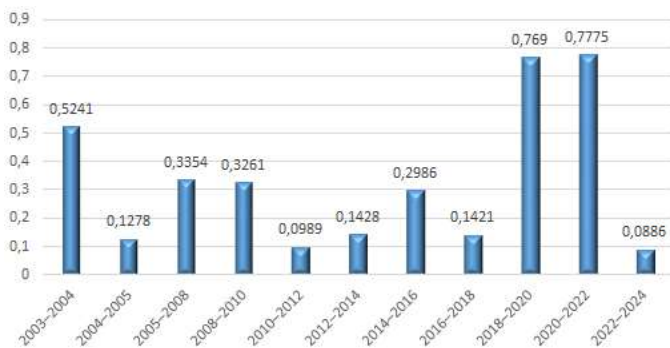


Рисунок 2 – Динаміка зміни індексу належності України до кластерів

Таким чином, запропонований підхід дозволив виявити високу стабільність кластерної структури, яка свідчить про те, що стабільність є закономірною і відображає реальні стійкі відмінності між країнами, обумовлені рівнем їх цифрового та економічного розвитку. Висока нестабільність належності України до кластерів протягом останніх років супроводжувалася її переходом до країн із високим рівнем розвитку електронного урядування.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. United Nations. E-Government Knowledgebase.
<https://publicadministration.un.org/egovkb/en-us/data-center>
2. Royna Daisy V., Nirmala S. Stability-integrated Fuzzy C means segmentation for spatial incorporated automation of number of clusters. *Sādhanā*. 2018. Vol. 43, 40. <https://doi.org/10.1007/s12046-018-0802-5>

УДК 004.8

*Артим В. І., Козлов О. В.
Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ІНТЕЛЕКТУАЛЬНА СИСТЕМА БАГАТОКАНАЛЬНОЇ ПІДТРИМКИ КЛІЄНТІВ У СФЕРІ ПОДОРОЖЕЙ

В роботі розглядається концепція інтелектуальної системи багатоканальної (омніканальної) підтримки клієнтів у сфері подорожей, яка поєднує можливості чат-ботів і голосових помічників з механізмами маршрутизації до живих операторів і підсистемою персоналізації. Описано архітектуру рішення (канальні адаптери, менеджер сесій і контексту, NLP-шар, бізнес-логіка і рекомендації, оркестратор ескаляції), механізми передачі контексту між каналами, підходи до класифікації намірів і витягування сутностей, а також сегменти тестування прототипу з метриками автоматичного вирішення запитів, точності намірів, часу відповіді і задоволеності користувачів.

Пропонована система базується на модульній архітектурі, що забезпечує незалежне масштабування компонентів і швидку інтеграцію нових каналів. Інтерфейсні шлюзи дозволяють підключати веб-чат, месенджери (Telegram, WhatsApp, Facebook Messenger), мобільні SDK і голосові канали (IVR, голосові асистенти). Менеджер сесій централізує збереження контексту — історії діалогу, станів транзакцій і метаданих користувача — щоб при переході між каналами або при ескаляції до оператора зберігався повний контекст спілкування. NLP-платформа поєднує два підходи: трансформерні моделі для розпізнавання складних намірів і витягування сутностей, а також легковагові класифікатори для сценаріїв з жорсткими SLA (короткий час відповіді). Поєднання правильних шаблонних флю для транзакційних дій (бронювання, скасування, повернення коштів) і генеративних компонентів для вільних консультацій дає баланс між точністю та гнучкістю відповідей.

Академічні огляди підтверджують, що такий гібридний підхід ефективний у туристичній галузі, де одночасно потрібна автоматизація простих запитів і людська участь у складних випадках [1].

Ключовою особливістю системи є механізм «seamless handoff» — безшовна передача контексту від бота до оператора. Це досягається через формалізований transcript діалогу, структуровані сутності й індикатори стану транзакції, що дозволяють оператору одразу бачити попередні питання, робочі дії та ймовірні варіанти рішення. Розумна маршрутизація визначає порогові значення довіри (confidence thresholds) для автоматичної відповіді: якщо ймовірність розпізнавання наміру висока — бот обробляє запит самостійно; якщо середня — бот уточнює; якщо низька — відразу ескалює до оператора. Така стратегія зменшує ризики неправильної автоматичної обробки та оптимізує завантаження контакт-центру, що підтверджується результатами сучасних досліджень у сфері чат-ботів і CRM [2].

Персоналізація реалізується через інтеграцію з CRM та системами бронювання: профіль користувача, історія попередніх поїздок та уподобання використовуються для ранжування рекомендацій (готелі, маршрути, додаткові послуги) і для адаптації тональності відповіді. Безпека і відповідність регуляціям — важливі вимоги: дані шифруються у транзиті та збереженні, доступ до РІ захищений рольовою моделлю, а лог-дані анонімізуються для навчання моделей. Для збереження конфіденційності передбачено прозоре інформування користувачів про автоматичні рішення та можливість відмови від автоматизації у чутливих сценаріях [2].

Щодо експериментальної частини, для практичної демонстрації створено прототип омніканальної системи (вебчат + Telegram + IVR) з вбудованим класифікатором намірів на базі «BERT-lite» та простим механізмом маршрутизації до оператора. Прототип тестували протягом двох тижнів із залученням 30 добровольців, котрі ставили реальні запити: пошук рейсу, зміна броні, уточнення правил багажу, запити щодо в'їзних обмежень. Метрики показали, що система може ефективно обробляти більшість типових запитів: коефіцієнт автоматичного вирішення становив 60–65%, точність класифікації намірів — близько 86–90% на тестовому наборі, середній час відповіді бота — <1 секунда в вебчаті, а середній час підключення оператора після ескалації — близько 40–50 секунд. Оцінка задоволеності користувачів (шкала 1–5) показала середній бал 4.2 для сесій, вирішених ботом, і 4.5 для сесій із ескалацією, що вказує на позитивне сприйняття гібридної моделі. Аналіз помилок виявив, що основні слабкі місця — некоректне витягування сутностей у складних фразах і потреба

в кращій локалізації для специфічних подорожніх запитів (віза, митні правила, локальні COVID/санітарні обмеження) [1] (рис.1, а, б). На рис.1, б показано середній час відповіді: чат-бот відповідає приблизно 0.9 секунди, людина-оператор — у середньому 45 секунд.

Отримані результати узгоджуються з висновками систематичних оглядів про роль чат-ботів у туризмі: автоматизація добре працює для стандартних і повторюваних операцій, тоді як для запитів з високою семантичною неоднозначністю потрібні або покращені NLP-моделі, або швидка ескалація до людини. Рекомендації на основі пілотного тесту включають розширення тренувальних корпусів специфічними travel-прикладми, посилення модулів витягування сутностей (slot filling) і введення механізмів безперервного навчання за зворотним зв'язком від операторів [1].

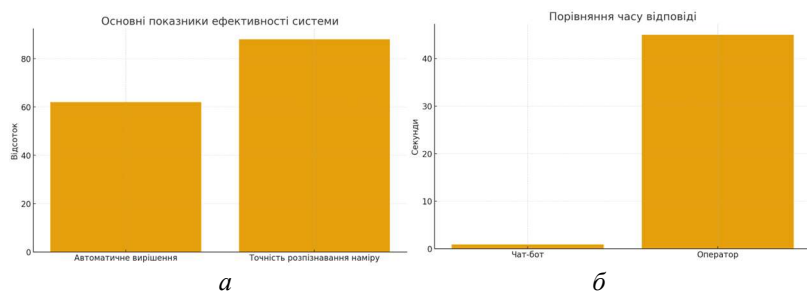


Рис.1 Основні показники ефективності системи: а) відсоток автоматичного вирішення запитів та точність класифікації; б) порівняння часу відповіді

Підсумком є те, що інтелектуальна багатоканальна система підтримки для індустрії подорожей є практично реалізованою і приносить значну операційну вигоду: скорочення навантаження на контакт-центр, прискорення відповіді й підвищення клієнтської задоволеності при збереженні контролю над складними випадками. Подальші дослідження доцільно зосередити на масштабуванні системи і розробці мультимодальних інтерфейсів (зображення, голосова біометрія).

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Ouaddi, C.; Hammoudi, S.; Bouslama, F. Exploring and analyzing the impact of chatbots in tourism. In: Proceedings of the 2024 ACM International Conference on Interactive Experiences for TV and Online Video (TVX '24). New York (NY): ACM; 2024. p. 1-10.
2. Khneyzer, Chadi; Boustany, Zaher; Dagher, Jean. AI-Driven Chatbots in CRM: Economic and Managerial Implications across Industries. *Administrative Sciences*. 2024, Vol. 14, No. 8, article 182.

УДК 621.398:004.71/.72

Банар А.Ю., Воробець Г.І.

*Чернівецький національний університет
ім. Юрія Федьковича, м. Чернівці, Україна*

ПОБУДОВА ТА УПРАВЛІННЯ SDN МЕРЕЖ З ВИКОРИСТАННЯМ МЕТОДІВ ШТУЧНОГО ІНТЕЛЕКТУ

Software-Defined Networking (SDN) поєднує централізацію управління та гнучкість конфігурації мережевих політик, що критично важливо для сучасних динамічних та гетерогенних систем. Розділення площини даних і керування [1] дає змогу швидко змінювати маршрутизацію, політики якості надання послуг та механізми безпеки без фізичного втручання.

Методи штучного інтелекту (ШІ) розширюють ці можливості, дозволяючи реалізувати адаптивний розподіл ресурсів, виявлення аномалій, прогнозування навантаження та оптимізацію мережевих ресурсів у режимі реального часу [2]. Інтеграція ШІ дозволяє підвищити ефективність SDN-контролерів і мінімізувати затримки прийняття рішень [3].

Створення експериментальної інтелектуальної SDN-платформи для побудови стало необхідною умовою для:

- відлагодження алгоритмів ШІ в умовах обмежених ресурсів;
- відтворюваності експериментів та перевірки нових моделей ШІ;
- спільної роботи дослідників і розробників над реальними мережевими сценаріями;
- швидкого тестування контролерів, топологій і модулів оптимізації.

Запропонована система (Рис. 1) складається з трьох функціональних рівнів:

Рівень даних - включає інфраструктурні елементи мережі (хости, комутатори, API-сервіси, мережеві ресурси).

Рівень управління SDN - представлений контролером Ryu, що здійснює централізоване управління потоками трафіку та взаємодіє з ІІІ-агентами через REST-інтерфейси.

Прикладний рівень - реалізує алгоритми машинного навчання для прогнозування навантаження, керування доступом до обмежених ресурсів (наприклад, rate-limited API), оптимізації маршрутів і балансування навантаження.

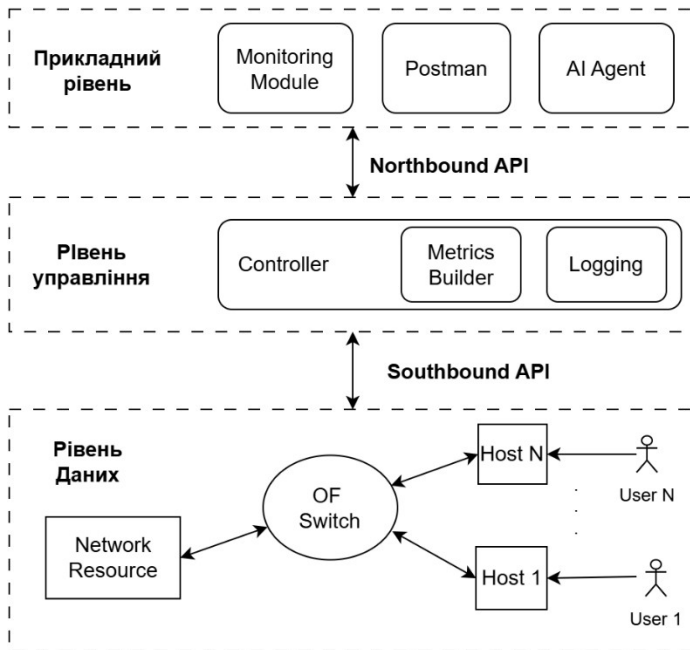


Рисунок 1 – Архітектура побудованої інтелектуальної системи

Запропонована інтелектуальна система дозволяє об'єднати можливості SDN та ІІІ для побудови адаптивних та керованих мереж. Інтегрований підхід забезпечує підвищення справедливості та точності розподілу обмежених ресурсів, зменшення кількості порушень лімітів та зниження затримок прийняття рішень у мережних сценаріях реального часу. Платформа є перспективним інструментом для подальших досліджень у напрямках оптимізації якості обслуговування, безпеки та енергоспоживання.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Jiménez, M.B., Fernández, D., Rivadeneira, J.E., Bellido, L. & Cárdenas, A.A. Survey of the Main Security Issues and Solutions for the SDN Architecture. *IEEE Access*, 2021, vol. 9, pp. 122016-122038. DOI: 10.1109/ACCESS.2021.3109564.
2. Банар А. Ю., Воробець Г. І. Алгоритми штучного інтелекту для оптимізації функціонування SDN: сучасні підходи та перспективи. *Зв'язок*. 2025. № 4. С. 11-18. DOI: 10.31673/2412-9070.2025.041241.
3. Adanza, D., Gifre, L., Alemany, P., Fernández-Palacios, J.-P., González-de-Dios, O., Muñoz, R., & Vilalta, R. Enabling Traffic Forecasting with Cloud-Native SDN Controller in Transport Networks. *Computer Networks*, 2024, vol. 250, article no. 110565. DOI: 10.1016/j.comnet.2024.110565.

УДК 004.42

Безощук С. С., Калініна І. О.

*Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ІНТЕЛЕКТУАЛЬНА СИСТЕМА МОНІТОРИНГУ ПАЛИВНО- МАСТИЛЬНИХ МАТЕРІАЛІВ

Ефективність використання паливно-мастильних матеріалів (ПММ) є критичним фактором морської логістики, оскільки витрати на паливо складають до 50-60% бюджету, а екологічні вимоги щодо викидів CO₂ постійно посилюються. Традиційні методи контролю на основі статичних норм втрачають ефективність у динамічних умовах, що обумовлює необхідність впровадження систем адаптивного моніторингу.

Розробка інтелектуальної системи моніторингу ПММ «OptiFuel» здійснюється з метою підвищення ефективності експлуатації морського транспорту та удосконалення процесів контролю паливних ресурсів, що дозволяє мінімізувати втрати та забезпечити високоточне виявлення аномальних відхилень у реальному часі. Для досягнення поставленої мети проведено порівняльний аналіз методів машинного навчання (Machine Learning, ML), зокрема лінійної регресії (Linear Regression, LR), випадкового лісу (Random Forest, RF) та градієнтного бустингу (Gradient Boosting, GB). Застосування цих алгоритмів дозволило суттєво покращити якість моделювання, врахувавши складні нелінійні

залежності між споживанням палива та експлуатаційними параметрами (ефективність двигуна, погодні умови, маршрут).

Аналіз досліджень підтверджує актуальність теми: роботи [1, 2] демонструють ефективність моделювання витрат та оптимізації часу переходу для збалансованого планування операцій. Водночас, дослідження [3] наголошує на критичній важливості впровадження пояснюваного штучного інтелекту (XAI) для забезпечення прозорості рішень, що підкреслює доцільність комплексного підходу в концепції «Smart Shipping».

Для навчання та тестування моделей використано набір даних [4] телеметрії морських перевезень, що містить записи про рейси суден різних типів (танкери, буксири тощо). На етапі попередньої обробки даних (Data Preprocessing) виконано кодування категоріальних ознак (One-Hot Encoding, Ordinal Encoding) та циклічне перетворення часових міток (Month Sin/Cos). Виявлено високу мультиколінеарність між споживанням палива та викидами CO₂ ($r > 0.99$), що обумовило виключення останнього фактора для уникнення витоку даних (data leakage).

Архітектура розробленої системи «OptiFuel» базується на мікросервісному підході, що забезпечує масштабованість та надійність. Комплекс складається з клієнтської частини (React, Mantine UI), основного серверу бізнес-логіки (.NET 9 Web API) та сервісу машинного навчання (Python, FastAPI). Зберігання історії прогнозів та метаданих реалізовано на базі PostgreSQL. Важливою складовою є контейнеризація всіх компонентів за допомогою Docker, що спрощує розгортання системи (див. рис. 1).

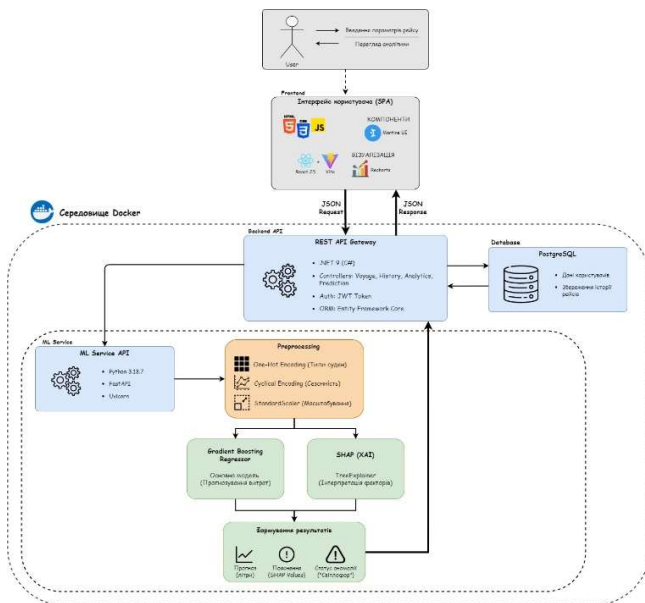


Рисунок 1 – Схема архітектури інтелектуальної системи моніторингу ПММ

У ході експериментальних досліджень було встановлено, що ансамблеві методи значно перевершують лінійні моделі. Базова модель LR показала коефіцієнт детермінації $R^2 = 94.71\%$ та $RMSE = 1192$ л. Використання моделі RF дозволило покращити точність (), однак найкращий результат продемонстрував алгоритм Gradient Boosting з $R^2 = 95.49\%$ та $RMSE = 1101$ л. На основі цієї моделі розроблено алгоритм детекції аномалій, який класифікує рейси за ступенем відхилення фактичного споживання від прогнозного («норма» < 3%, «попередження» 3 – 7%, «аномалія» > 7%). Додатково інтегровано модуль пояснюваного ШІ (XAI) на базі бібліотеки SHAP, що дозволяє візуалізувати вплив кожного фактора (наприклад, штормової погоди) на кінцевий прогноз.

Таким чином, розроблена система «OptiFuel» поєднує високу точність прогнозування на базі Gradient Boosting із сучасним веб-інтерфейсом та аналітичними інструментами. Це дозволяє автоматизувати контроль за витратами ПММ, мінімізувати ризики крадіжок та підвищити загальну енергоефективність транспортних перевезень.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Yijian Sang, Yu Ding, Jiarun Xu, Congbiao Sui (2023). Ship voyage optimization based on fuel consumption under various operational conditions. *Fuel*, vol 352, 129086, ISSN 0016-2361. URL: <https://doi.org/10.1016/j.fuel.2023.129086>.
2. Krzysztof Rudzki, Piotr Gomulka, Anh Tuan Hoang (2022). OPTIMIZATION MODEL TO MANAGE SHIP FUEL CONSUMPTION AND NAVIGATION TIME. *POLISH MARITIME RESEARCH 3 (115) 2022 Vol. 29; pp. 141-153*. URL: <https://sciendo.com/2/v2/download/article/10.2478/pomr-2022-0034.pdf>.
3. Mohamed Abuella, M. Amine Atoui, Slawomir Nowaczyk, Simon Johansson & Ethan Faghani (2023). Data-Driven Explainable Artificial Intelligence for Energy Efficiency in Short-Sea Shipping. *Machine Learning and Knowledge Discovery in Databases*, pp 226-241. URL: https://link.springer.com/chapter/10.1007/978-3-031-43430-3_14.
4. Ship fuel efficiency dataset. *Kaggle*. URL: <https://www.kaggle.com/code/oumaimagasm/ship-fuel-efficiency/input>

УДК 004.056.5

Горшколенов І. В., Сіденко Є. В.

*Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ПРОГНОЗУВАННЯ ПОТРЕБ У РЕСУРСАХ ТА ОПТИМІЗАЦІЯ ЛОГІСТИЧНИХ МАРШРУТІВ НА ОСНОВІ НЕЙРОННИХ МЕРЕЖ

Управління ланцюгами постачання в умовах сучасних бойових дій, що визначаються концепцією «оспорованої логістики» (Contested Logistics), характеризується критично високим рівнем невизначеності та динамічності середовища. Традиційні детерміновані моделі планування втрачають ефективність через неможливість врахування стохастичних факторів, таких як раптові зміни метеоумов, стан дорожнього покриття та накопичена втома техніки. Це вимагає впровадження адаптивних інтелектуальних систем, здатних прогнозувати нелінійні зміни у потребах ресурсів та оптимізувати маршрути руху колон у режимі реального часу.

Розробка гібридної інтелектуальної системи прогнозування потреб та оптимізації маршрутів здійснюється з метою підвищення адаптивності та ефективності логістичного планування шляхом

інтеграції динамічних прогнозів споживання ресурсів з механізмом прийняття рішень. Для досягнення поставленої мети використано підхід Neuro-Symbolic AI, що поєднує методи глибокого навчання (Deep Learning) для роботи з часовими рядами та нечіткої логіки (Fuzzy Logic) для прийняття рішень. Зокрема, застосовано рекурентні нейронні мережі архітектури Gated Recurrent Unit (GRU) для прогнозування попиту та метод багатокритеріальної оптимізації Fuzzy-SAW (Simple Additive Weighting) для ранжування маршрутів. Таке поєднання дозволило врахувати приховані стани середовища та адаптувати вагові коефіцієнти критеріїв безпеки залежно від прогнозованої складності умов.

Аналіз предметної області свідчить про обмеженість існуючих підходів до планування військової логістики. Стандартні методики, описані в NATO Logistics Handbook [1], базуються на фіксованих нормативах споживання (L/100km), які ігнорують кумулятивний вплив зовнішніх факторів. Водночас дослідження J. Chung et al. [2] підтверджують перевагу GRU над LSTM у задачах моделювання послідовностей на пристроях з обмеженими ресурсами («тактичний край»), а роботи L. Zadeh [3] обґрунтовують застосування нечітких множин для формалізації експертних знань в умовах невизначеності. Це підтверджує доцільність інтеграції предиктивної аналітики в контур управління логістикою для виявлення нелінійних залежностей, які не фіксуються лінійними моделями.

Для навчання та оцінювання ефективності розробленого модуля прогнозування було створено генератор синтетичних даних, що імітує логістичні процеси протягом 1095 діб (3 роки). Генератор базується на ланцюгах Маркова та враховує вплив «прихованих станів» (Hidden States), таких як індекс в'язкості ґрунту (Mud Index) та індекс зносу техніки (Fatigue Index), які мають інерційний характер. Вхідний вектор ознак включає метеорологічні дані, темп операцій та циклічне кодування часу ($\sin/\cos$ time), що забезпечило репрезентативність вибірки та здатність моделі виявляти сезонні патерни.

Архітектура розробленої системи базується на двохетапній обробці даних. Перший модуль – прогнозування – реалізовано на базі стекової нейромережі GRU (64+32 нейрони) з регуляризациєю Dropout, що перетворює вхідні часові ряди на прогноз абсолютної потреби в пальному. Другий модуль – прийняття рішень – трансформує отриманий прогноз у динамічний коефіцієнт «паливної інтенсивності» (intensity factor, л/км), який подається на вхід алгоритму Fuzzy-SAW разом із параметрами доступних маршрутів. Система автоматично перераховує рейтинг маршрутів, змінюючи пріоритети між критеріями

часу, витрат та безпеки (Risk Index) залежно від прогнозованого рівня навантаження (див. рис. 1).

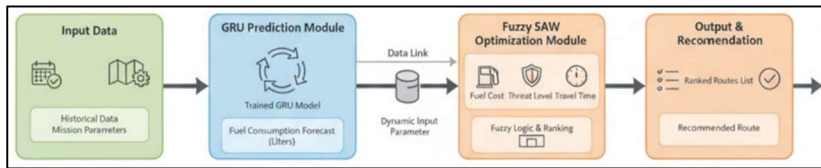


Рисунок 1 – Структурна схема гібридної нейро-символьної системи прогнозування та оптимізації логістики

У процесі навчання моделі GRU на тестовій вибірці було досягнуто показника середньої абсолютної похибки (MAE) на рівні 289.02 літрів для логістичного підрозділу. Детальний аналіз показав, що дана величина не є помилкою навчання, а відображає «шумовий поріг» (Noise Floor) бойового середовища – стохастичний компонент, спричинений мікро-подіями (пробуксовка, маневрування), який неможливо детермінувати (див. рис. 2).

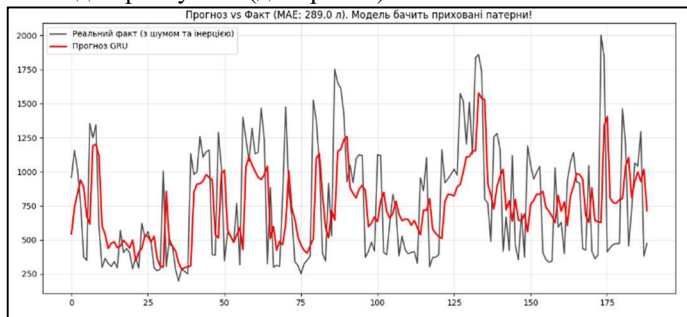


Рисунок 2 – Порівняння прогнозу моделі з реальним споживанням

Вирішено використовувати значення MAE як математично обґрунтовану величину страхового запасу (Safety Stock). Валідацію системи проведено в середовищі імітаційного моделювання Arma 3 із застосуванням масштабування дистанцій 1:10. У сценарії «Шторм» контрольна група (без СППР) обрала найкоротший маршрут, що призвело до аварії та втрати вантажу. Експериментальна група, дотримуючись рекомендації Fuzzy-SAW (об'їзний маршрут), забезпечила 100% успішність місії (Mission Success Rate), незважаючи на планове збільшення витрат на 18%.

Таким чином, реалізована інтелектуальна система продемонструвала здатність виявляти приховані нелінійні залежності у споживанні ресурсів та адаптивно керувати ризиками. Поєднання рекурентної нейромережі GRU з алгоритмом Fuzzy-SAW дозволило перейти від реактивного планування до предиктивного, забезпечуючи баланс між економією ресурсів та безпекою особового складу в умовах невизначеності.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. NATO Logistics Handbook. Brussels : NATO Headquarters, 2012. 238 p. (URL: https://www.nato.int/docu/logi-en/logistics_hndbk_2012-en.pdf).
2. Chung J., Gulcehre C., Cho K., Bengio Y. Empirical evaluation of gated recurrent neural networks on sequence modeling. arXiv preprint arXiv:1412.3555. 2014. DOI: <https://doi.org/10.48550/arXiv.1412.3555>.
3. Zadeh L. A. Fuzzy sets. Information and Control. 1965. Vol. 8, No. 3. P. 338–353.

УДК 004.896

Денисенко Н. В.

*Одеський національний університет
імені І. І. Мечникова, м. Одеса, Україна*

МОДЕЛЮВАННЯ ІНТЕЛЕКТУАЛЬНОЇ СИСТЕМИ КЕРУВАННЯ МІКРОКЛІМАТОМ ТЕПЛИЧНОГО ГОСПОДАРСТВА

Інтелектуальне та автоматизоване керування мікрокліматом тепличного господарства сприяє підвищенню ефективності цих господарств. Для росту рослин найважливішими параметрами є: температура, вологість, освітлення та концентрація вуглекислого газу. Саме ці параметри будуть зчитуватись датчиками системи, аналізуватись та регулюватись в разі потреби.

Така система буде складатись з чотирьох елементів:

1. Виконавчі пристрої, які змінюють певні параметри
2. Сенсори для зчитування параметрів (в ґрунті та зовнішні) [1]
3. Програмне забезпечення з використанням штучного інтелекту або нечіткої логіки
4. Інтерфейс користувача, через який він може відслідковувати поточні параметри та статистику.

Параметри самої системи також взаємопов'язані: підвищення температури знижує вологість, фотосинтез залежить від світла та вуглекислого газу. Такі тонкощі також треба буде враховувати при написанні програмного коду. [2]

Наразі існує декілька видів систем керування окрім ручної: автоматичні системи на реле і таймерах та сучасні цифрові системи на мікроконтролерах. Такі системи не вміють самонавчатися, побудовані на жорсткому алгоритмі if-then та погано адаптуються до зміни у погоді чи сезонності. Впровадження інтелектуальної системи може виправити ці недоліки. Використовуючи нейронні мережі можна отримати найточніші дані про конкретну рослину, як за нею доглядати, які параметри для неї будуть найбільш оптимальні. Завдяки цьому вирощування рослини буде найефективнішим. Інтелектуальна система до того ж може слідувати за часом доби, сезоном та погодою. [3]

Для ефективної реалізації такої системи було спочатку проведено моделювання. Було побудовано структурну схему IDEF3, змодельовано fuzzy-правила та відображено взаємодії всіх апаратних компонентів за допомогою Cisco Packet Tracer. Було виконано не тільки програмну частину, але й апаратну. Для цього використано платформу для створення електронних пристроїв Arduino. Ця платформа має великий вибір апаратних рішень, а також власну середу розробки та мову програмування C/C++ спрощену за допомогою бібліотеки "Wiring".

У результаті реалізовано програмне забезпечення та мініатюрну версія апаратного рішення для демонстрації роботи, схему наведено на рисунку 1.

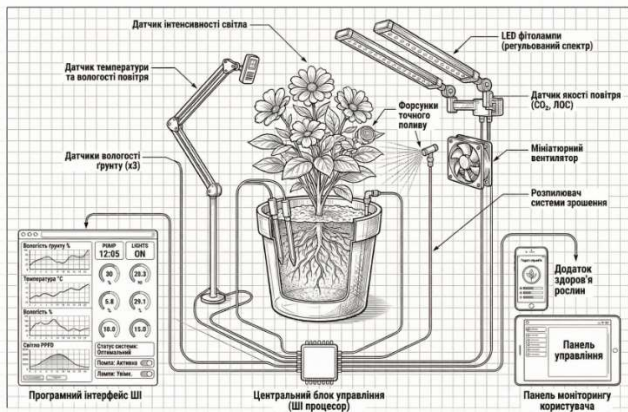


Рис. 1 – схема апаратного рішення для демонстрації (зображення створено за допомогою ШІ)

В роботі показано весь процес створення такої системи, починаючи з створення структурних діаграм, схем підключення, написання коду, апаратної реалізації та тестування. Інтеграція інтелектуальної системи дозволила додати автоматизацію процесів догляду за рослиною та надавати користувачу інформацію у зручному форматі

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Гайдукевич С. В., Семенова Н. П., Леськів Я. А. Автоматизована система керування електрообладнанням у спорудах закритого ґрунту. *Вісник Черкаського державного технологічного університету*. 2021. № 1. С. 20–31. DOI: 10.24025/2306-4412.1.2021.216915
2. Канарський Є., Орехов О., Желтухін О. Розробка автоматизованої системи керування тепличним господарством. *Молодий вчений*. 2022. № 3 (103). С. 5–12. DOI: 10.32839/2304-5809/2022-3-103-2
3. Грачов О. Plant monitoring system: ші-система для розумного моніторингу рослин. *Information technology and society*. 2025. № 1 (16). С. 59–64. DOI: 10.32689/maup.it.2025.1.7

УДК 004.85; 004.89

Кисса А. Ф., Гунченко Ю.О.

*Одеський національний університет
імені І.І.Мечникова, м. Одеса, Україна*

АНАЛІЗ ТА ПОРІВНЯННЯ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ РОЗПІЗНАВАННЯ ДОРОЖНІХ ЗНАКІВ

Дорожні знаки відіграють життєво важливу роль на дорозі. Зі зростанням соціально-економічного розвитку та попиту на ефективний і безпечний транспорт розвиваються системи автономного водіння та інтелектуальні транспортні системи, а розпізнавання дорожніх знаків (TSR) є ключовим компонентом таких систем, через що TSR привертає значну увагу в останні роки [1-3].

Попри значний прогрес у сфері розпізнавання дорожніх знаків, сучасні системи залишаються вразливими до впливу освітлення, погодних умов та оклюзій, що знижує їхню точність у реальних сценаріях. Тому актуальним є пошук рішень, які поєднують високу продуктивність і компактність моделі [1].

Щоб підвищити точність і стабільність розпізнавання, у сучасних дослідженнях застосовуються різні моделі, що поєднують традиційні та

глибокі алгоритми, мультимодальні дані та оптимізовані архітектури нейронних мереж. До найефективніших належать [2]:

1. Згорткова нейронна мережа (CNN) для одночасного виявлення та оцінки меж – метод, заснований на згорткових нейронних мережах, забезпечує одночасне виявлення дорожніх знаків і точне визначення їх меж. Архітектура має дві гілки: одна відповідає за класифікацію та передбачення рамок, інша — за уточнення контурів знака. Завдяки використанню Німецького еталонного набору даних для розпізнавання дорожніх знаків (GTSRB) та прийомів доповнення зображень модель демонструє високу стійкість до змін освітлення, погодних умов та часткових перекриттів.

2. Підхід, який поєднує зображення дорожніх знаків із додатковою контекстною інформацією (GPS, погодою, часом доби) – Мультимодальна модель на основі трансформерів (MTSRB). Модель використовує трансформер як основну архітектуру, що дозволяє ефективно обробляти кілька джерел даних одночасно. Двоступний механізм «Think Twice» спочатку розпізнає знак лише за зображенням, а потім уточнює результат, якщо є неоднозначність. Такий підхід підвищує точність у складних умовах, де одного зображення недостатньо.

3. Гібридна 2D–3D CNN на основі трансферного навчання поєднує попередньо навчені 2D CNN (такі як ResNet або VGG16) для розпізнавання дорожніх знаків і 3D CNN для аналізу дорожньої сцени. Це дозволяє враховувати просторову інформацію про дорогу та покращує загальну якість розпізнавання. Модель навчається на кількох наборах даних та показує високу точність як у класифікації знаків, так і у визначенні дорожньої розмітки, що робить її корисною для систем допомоги водію (ADAS).

4. Легка система розпізнавання на основі LeNet використовується як компактне та швидке рішення для базових завдань виявлення дорожніх знаків. Модель складається з простих згорткових шарів, що робить її придатною для пристроїв з обмеженими ресурсами. Попередня обробка зображень — фільтрація, сегментація кольору та виділення ознак — допомагає підвищити точність розпізнавання.

5. Моделі на основі YOLO (You Only Look Once) демонструють потенціал для розпізнавання дорожніх знаків у реальному часі завдяки об'єднанню виявлення й класифікації в одній мережі. Оптимізовані версії YOLOv5 та YOLOv8 застосовують механізми уваги (CBAM, Coordinate Attention), вдосконалені функції втрат (SIoU, WIoU) та додаткові шари ознак для поліпшення виявлення малих цілей. Такі моделі дозволяють ефективно розпізнавати дрібні та частково перекриті

знаки, зберігаючи швидкість обробки, що критично для вбудованих систем автономного водіння і ADAS [3].

Таким чином, сучасні підходи до розпізнавання дорожніх знаків поєднують високоточні нейронні мережі, оптимізовані алгоритми обробки зображень та використання мультимодальних даних. Завдяки цьому досягається стабільне та швидке виявлення навіть малих або частково перекритих об'єктів, що особливо важливо для систем автономного водіння та ADAS.

Ефективність таких моделей може бути оцінена за різними показниками точності, що дозволяє порівнювати їх між собою та обирати найбільш придатні для практичного застосування.

Для порівняння ефективності різних підходів до розпізнавання дорожніх знаків наведено відсоткову точність основних моделей (рис. 1) на основі аналізу [1-3].

В роботі проведено аналіз методів та моделей розпізнавання дорожніх знаків з підвищеною точністю і стабільністю розпізнавання, що є ключовим елементом інтелектуальних транспортних систем та систем автономного водіння, забезпечуючи безпечне керування транспортними засобами.

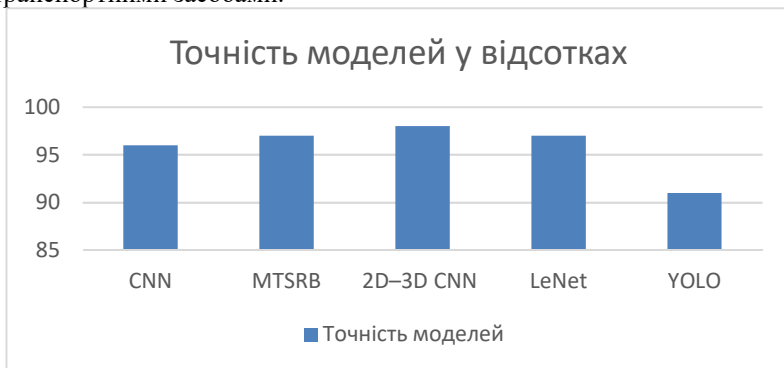


Рис. 1 – Порівняння точності моделей розпізнавання дорожніх знаків

Сучасні моделі TSR, такі як CNN, MTSRB, 2D-3D CNN, LeNet та оптимізовані YOLO, демонструють високу точність і швидкість роботи, зокрема при виявленні малих або частково перекритих знаків. Подальший розвиток TSR повинен фокусуватися на підвищенні стабільності розпізнавання в реальних умовах, оптимізації нейронних мереж і інтеграції мультимодальних даних для забезпечення надійної роботи автономних транспортних засобів і систем допомоги водію (ADAS).

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Zhu Y., Yan W. Q. Traffic sign recognition based on deep learning. *Multimedia tools and applications*. 2022. URL: <https://doi.org/10.1007/s11042-022-12163-0> (дата звернення: 25.11.2025).
2. Asritha P. A survey on traffic sign recognition and detection: trends, techniques, applications. *International journal of research publication and reviews*. 2024. Т. 5, № 11. С. 6791–6799.
3. Zhu Y., Yan W. Q. Traffic sign recognition based on deep learning. *Multimedia tools and applications*. 2022. URL: <https://doi.org/10.1007/s11042-022-12163-0> (дата звернення: 25.11.2025).

УДК 004.8

Кісельов Б. Г.

*Криворізький національний університет
м. Кривий Ріг, Україна*

МЕТОДИ КЛАСИФІКАЦІЇ РУДОВИХ МАТЕРІАЛІВ ЗА СЕНСОРНИМИ ПОКАЗНИКАМИ ПРИ НИЗЬКІЙ КОНТРАСТНОСТІ ДАНИХ

Сенсорне сортування руд функціонує в умовах низького контрасту даних, коли сенсорні ознаки рудного та нерудного матеріалу значною мірою перекриваються через неоднорідність структури породи, варіативність форми шматків і шумові впливи вимірювальної апаратури. За таких умов лінійне розділення класів є недостатнім, а коректна класифікація потребує здатності моделі моделювати складні залежності між ознаками.

Ситуацію ускладнює обмежений обсяг навчальної вибірки, оскільки достовірне маркування зразків ґрунтується на лабораторних хімічних аналізах, що є ресурсомісткими та не можуть виконуватися у великих кількостях. Вибір методів класифікації повинен забезпечувати високу стійкість до шуму та мінімальні вимоги до кількості навчальних даних [1,2].

У Таблиці 1 розглянуті й порівняні основні підходи до класифікації сенсорних даних: статистичні методи (LDA/QDA), методи машинного навчання (SVM, ансамблеві методи), а також нейромережеві моделі. Порівняння здійснено за ключовими критеріями доцільності використання у середовищі сенсорного сортування руди за низького контрасту ознак та малої вибірки.

Таблиця 1 – Оцінювання методів класифікації для сенсорного сортування руди [2,3]

Метод	Потреба у навчальних даних	Стійкість до шуму та варіативності форми	Нелінійне розділення при перекритті класів	Узагальнена характеристика ефективності
LDA (Linear Discriminant Analysis)/QDA (Quadratic Discriminant Analysis)	Низька	Середня	Обмежене	Використовуються як базові моделі-орієнтири, але недостатні при складних розподілах
Naive Bayes	Дуже низька	Низька	Обмежене	Простий baseline; не враховує кореляцію ознак сенсорних сигналів
Логістична регресія	Середня	Середня	Обмежене	Лінійна модель для ймовірнісної класифікації; ефективність обмежена при суттєвому перекритті класів
k-NN (k-Nearest Neighbor)	Висока	Низька	Обмежено –середнє	Чутливий до шуму; нестабільний на малих вибірках
SVM (Support Vector Machine)	Середня	Висока	Високе	Оптимальний вибір для умов low-data та low-contrast; формує складні межі
Ансамблеві методи	Середня	Висока	Високе	Стійкі до шуму, придатні для багатовимірних табличних ознак
Нейромережі	Висока	Середня–висока	Дуже високе	Перспективні, але потребують суттєвого розширення вибірки

Проведений аналіз свідчить, що найбільш придатними для сенсорного сортування руди за умов низької контрастності та дефіциту даних є SVM та ансамблеві методи, які забезпечують стійке розділення класів навіть при значному перекритті розподілів. Статистичні моделі можуть застосовуватися як початкові орієнтири, а неймережеві підходи – як перспективний напрям розвитку за умови збільшення вибірки.

Оскільки властивості рудної сировини змінюються під час промислової експлуатації, доцільним є впровадження інкрементального навчання, яке дає змогу моделі адаптуватися до нових умов без переривання технологічного процесу, підтримуючи стабільно високу точність класифікації [4].

ВИСНОВОК

Запропоноване порівняння методів класифікації показало, що для сенсорного сортування рудних матеріалів у складних умовах низької контрастності ознак та дефіциту навчальних даних найбільш ефективними є методи опорних векторів (SVM) та ансамблеві методи, які забезпечують стійке розділення класів і високу завадостійкість. Враховуючи динамічність властивостей рудної сировини, перспективним напрямом розвитку є впровадження адаптивних стратегій навчання, що дозволяють моделі уточнювати свої параметри на основі нових виробничих даних, підтримуючи стабільність і результативність класифікації в умовах промислової експлуатації.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Zeng D. A novel magnetite ore refined sorting method based on magnetic induction and CNN-SK-BiLSTM network. *Gospodarka surowcami mineralnymi - mineral resources management*. 2025. Т. 41, № 1. С. 123–139. URL: <https://doi.org/10.24425/gsm.2025.153171> (дата звернення: 24.11.2025).
2. Xu Y. Development of methodological framework for XRF sensor based ore sorting utilizing machine learning approaches : PhD thesis. 2024. URL: <https://doi.org/10.14288/1.0444099> (дата звернення: 25.11.2025).
3. Choosing the best machine learning classification model and avoiding overfitting. *MathWorks - Maker of MATLAB and Simulink*. URL: <https://www.mathworks.com/campaigns/offers/next/choosing-the-best-machine-learning-classification-model-and-avoiding-overfitting.html> (дата звернення: 26.11.2025).
4. Awan A. A. What is incremental learning?. *Datacamp*. URL: <https://www.datacamp.com/blog/what-is-incremental-learning> (дата звернення: 27.11.2025).

ІНТЕЛЕКТУАЛЬНА СИСТЕМА КЛАСИФІКАЦІЇ ЕКОЛОГІЧНИХ ПОКАЗНИКІВ З УРАХУВАННЯМ ДИСБАЛАНСУ КЛАСІВ

Екологічний моніторинг є ключовим для забезпечення екологічної безпеки, оскільки стан природних ресурсів впливає на здоров'я населення та стабільність екосистем. Зростання обсягів екологічної інформації ускладнює її ручний аналіз і потребує застосування інтелектуальних методів. Особливо складними є задачі класифікації через дисбаланс класів, оскільки міноритарні класи часто відповідають критичним станам і є ключовими для управлінських рішень [1].

Метою дослідження є розробка інтелектуальної системи класифікації, здатної компенсувати дисбаланс класів і забезпечувати високу точність прогнозів. В експериментальній частині дослідження використано два набори даних – Water Potability [2] та Oil Spill Detection [3], які відрізняються ступенем дисбалансу.

Для обох наборів проведено аналіз і попередню обробку даних: аналіз структури даних, імпутацію пропусків, нормалізацію числових ознак, кодування категоріальних змінних та балансування класів (методи: Random Oversampling, Random Undersampling, SMOTE, ROSE), що покращило чутливість моделей до рідкісних подій [4].

Підсистема моделювання включає алгоритми, чутливі до дисбалансу – Logistic Regression, Support Vector Machine, Decision Tree, Random Forest, XGBoost – із вагами класів, а також ансамблеві методи: Balanced Bagging, RUSBoost, SMOTEBoost, EasyEnsemble та Balance Cascade [5]. Для оцінювання моделей використано метрики якості Precision, Recall, F1-score, Specificity, Balanced Accuracy, G-mean, AUC-ROC та MCC. Це дало можливість об'єктивно оцінити продуктивність моделей, незалежно від нерівномірного розподілу класів.

Отримані значення в табл. 1 демонструють суттєві відмінності в поведінці моделей залежно від ступеня дисбалансу. На першому наборі методи з oversampling і вагами класів покращили Recall міноритарного класу, проте не всі з них зберегли високі значення специфічності. Наприклад, SVM та XGBoost показали збалансовані результати, тоді як Logistic Regression мала найнижчий MCC.

Таблиця 1 – Порівняльна таблиця методів навчання чутливого до втрат для першого та другого наборів даних

Модель / № набору	Precision	Recall	F1	Specificity	Balanced Accuracy	G-mean	MC C	AUC ROC
XGBoost / 1	0.688	0.735	0.710	0.666	0.700	0.700	0.402	0.760
Random Forest / 1	0.730	0.718	0.724	0.735	0.726	0.810	0.453	0.810
SVM / 1	0.661	0.728	0.693	0.626	0.677	0.733	0.356	0.733
Logistic Regression / 1	0.505	0.494	0.500	0.516	0.505	0.508	0.010	0.508
Decision Tree / 1	0.625	0.478	0.541	0.713	0.595	0.611	0.196	0.611
XGBoost / 2	0.504	0.998	0.667	0.02	0.501	0.041	0.029	0.747
Random Forest / 2	0.509	0.990	0.672	0.045	0.518	0.211	0.107	0.716
SVM / 2	0.503	0.998	0.669	0.12	0.505	0.108	0.062	0.743
Logistic Regression / 2	0.846	0.506	0.633	0.908	0.707	0.678	0.452	0.813
Decision Tree / 2	0.534	0.977	0.691	0.149	0.563	0.381	0.223	0.8

Ансамблеві методи в табл. 2 на першому наборі показали високий Recall, але частина моделей втратила специфічність. Balanced Bagging продемонстрував найкращий компроміс між чутливістю та точністю, а Balance Cascade забезпечив найстабільнішу поведінку. На другому наборі, що має сильний дисбаланс, усі ансамблеві методи досягли високої специфічності, але Recall суттєво знизився. Найбільш рівномірні результати показали EasyEnsemble та Balanced Bagging.

Таблиця 2 – Порівняльна таблиця ансамблевих методів навчання для першого та другого наборів даних

Модель / № набору	Precision	Recall	F1	Specificity	Balanced Accuracy	G-mean	MC C	AUC ROC
RUSBoost / 1	0.500	0.980	0.667	0.002	0.501	0.041	0.029	0.747
SMOTEBoost / 1	0.509	0.990	0.672	0.045	0.518	0.211	0.107	0.716
EasyEnsemble / 1	0.503	0.998	0.669	0.012	0.505	0.108	0.062	0.743
Balanced Bagging / 1	0.846	0.506	0.633	0.908	0.707	0.678	0.452	0.813
Balance Cascade / 1	0.534	0.977	0.691	0.149	0.563	0.381	0.223	0.800
RUSBoost / 2	0.444	0.333	0.381	0.981	0.657	0.572	0.361	0.887
SMOTEBoost / 2	0.474	0.350	0.385	0.984	0.661	0.564	0.368	0.899
EasyEnsemble / 2	0.5	0.417	0.455	0.981	0.699	0.639	0.434	0.883
Balanced Bagging / 2	0.571	0.333	0.421	0.989	0.661	0.574	0.418	0.915
Balance Cascade / 2	0.5	0.250	0.333	0.989	0.619	0.497	0.334	0.940

Порівняння двох наборів підтвердило, що зі зростанням дисбалансу ефективність моделей змінюється: алгоритми, що добре працюють на більш збалансованих вибірках, втрачають продуктивність на складних. Це підкреслює необхідність адаптації алгоритмів та технік балансування до структури даних.

Система реалізована в RStudio із використанням бібліотек dplyr, caret, ROSE, smotefamily та pROC.

Практична цінність полягає у застосуванні для екологічного моніторингу, контролю якості екологічних показників та виявлення забруднень. Вона автоматизує аналіз великих вибірок та забезпечує

точне розпізнавання критичних станів, придатна для інтеграції в роботу лабораторій і дослідницьких центрів.

Наукова новизна дослідження полягає у поєднанні класичних та ансамблевих алгоритмів із різними стратегіями балансування даних, які використовують і на етапі препроцесінгу і та етапі моделювання.

Дослідження підтвердило, що врахування дисбалансу класів критично важливе для отримання якісних результатів при розв'язанні задачі класифікації. Розроблена інтелектуальна система демонструє високу точність, стійкість до шумів і може використовуватися для автоматизованого контролю природних ресурсів, раннього виявлення ризиків та підтримки ухвалення рішень.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Basharat A., Ali A., Mughal H., Mohamad M. M. Imbalanced data handling techniques for classification: a state-of-the-art review // *Procedia Computer Science*. – 2023. – URL: <https://pdf.elspublishing.com/paper/pj/open/pcs/v2/pcs2023020010.pdf>
2. Water Potability Dataset [Електронний ресурс] // Kaggle. – Режим доступу: <https://www.kaggle.com/datasets/adityakadiwal/water-potability>.
3. How to Develop an Imbalanced Classification Model to Detect Oil Spills [Електронний ресурс] // *AI Pro Blog*. – 2020. – URL: <https://www.aiproblog.com/index.php/2020/02/20/how-to-develop-an-imbalanced-classification-model-to-detect-oil-spills/>.
4. Zhao H., Li X., Wang S. A comprehensive survey of oversampling techniques for imbalanced data learning // *Expert Systems with Applications*. – 2021. – URL: <https://www.sciencedirect.com/science/article/pii/S0957417421005091>
5. Liu L. et al. Solving the class imbalance problem using ensemble algorithm: application of screening for aortic dissection // *BMC Medical Informatics and Decision Making*. – 2022. – URL: <https://bmcmmedinformdecismak.biomedcentral.com/articles/10.1186/s12911-022-01821-w>.

ІНФОРМАЦІЙНА СИСТЕМА РОЗПІЗНАВАННЯ ЕМОЦІЙ З ВИКОРИСТАННЯМ НЕЙРОННОЇ МЕРЕЖІ

Автоматичне розпізнавання людських емоцій, відоме як Facial Emotion Recognition (FER), є критично важливою функцією. Ці системи знаходять своє застосування в людино-машинній взаємодії, моніторингу та індивідуалізованих сервісах. Для того, щоб система FER була ефективною в умовах реального часу, вона повинна забезпечувати не лише високу точність, але й надзвичайну швидкість та компактність. Це створює виклик, оскільки класичні моделі глибокого навчання часто є занадто об'ємними та вимагають значних обчислювальних ресурсів. У цьому контексті, архітектура YOLOv8s (small) [1-2], як частина сімейства You Only Look Once, пропонує оптимальне рішення, забезпечуючи ідеальний баланс між продуктивністю та обчислювальною ефективністю для створення високошвидкісної та компактною системи розпізнавання емоцій.

Метою даної роботи було дослідження ефективності та розробка оптимізованої архітектури для розпізнавання восьми базових емоцій. До цих емоцій належать радість, смуток, злість, страх, здивування, огида, презирство та нейтральний стан. Досягнення цієї мети реалізовано через адаптацію легкої архітектури YOLOv8s, а також порівняння її ключових метрик швидкості та точності з класичними моделями згорткових нейронних мереж (CNN).

Для успішного вирішення задачі класифікації емоцій вирішальне значення має вибір якісного, збалансованого і надійного набору даних. Навчання проводилося на великому, розміченому наборі даних AffectNet [3], який був спеціально підготовлений у форматі YOLO, що включає зображення та відповідні текстові файли анотацій. Розподіл екземплярів за емоційними категоріями в датасеті відображено на рисунку 1. Для оптимізації процесу використовувався оптимізатор AdamW протягом п'ятдесяти епох. З метою оптимального використання ресурсів графічного процесора та запобігання перенавчанню, було застосовано автоматичний вибір розміру пакета (batch=-1) та механізм ранньої зупинки з терпінням у п'ятнадцять епох. Зображення були стандартизовані до розміру 640x640 пікселів, а для підвищення узагальнювальної здатності моделі до "диких" даних

AffectNet застосовувався потужний набір аугментацій. Сюди увійшли Mosaic-аугментація, випадкові геометричні перетворення, такі як повороти, зсуви, масштабування, та регулювання кольорових параметрів HSV. Крім того, був реалізований механізм відновлення навчання з останньої контрольної точки, що є критично важливим для забезпечення безперервності процесу та ефективного використання обчислювального часу, необхідного для інтеграції цієї високошвидкісної системи FER.

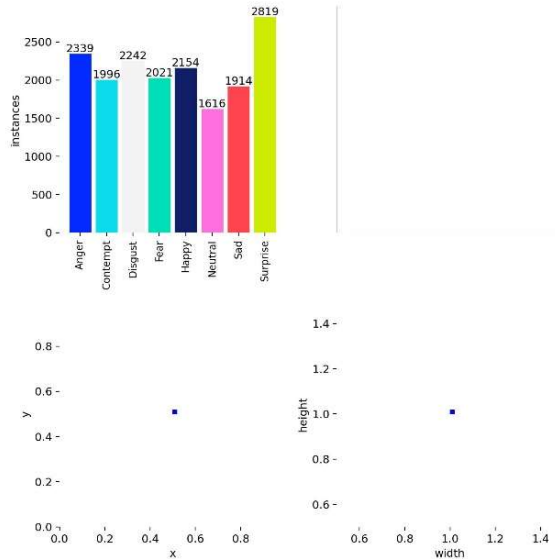


Рисунок 1 – Розподіл екземплярів за емоційними категоріями у навчальному наборі даних AffectNet

Адаптована модель на базі YOLOv8s продемонструвала високу ефективність на тестовому наборі даних. Було досягнуто загальної точності (Ассурасу) 0,82457 та збалансованого показника F1-Score на рівні 0,82457. Ключовим показником, що підтверджує придатність моделі для реального часу, є швидкість обробки 150 кадрів на секунду (FPS) на конкретній платформі. Ця швидкість на порядок перевищує швидкість інференсу класичних, глибоких CNN-моделей, таких як VGG16. Аналіз результатів класифікації показав, що найкращі результати отримано для емоцій радість та нейтральний стан, що свідчить про їхню високу візуальну відмінність. Водночас, найбільші помилки спостерігалися між емоціями страх та здивування (рисунок 2).



Рисунок 2 – Приклад роботи інформаційної системи: детекція обличчя та класифікація емоції з оцінкою впевненості

Таким чином, проведені дослідження довело, що архітектура YOLOv8s є високопродуктивною основою для створення ефективних систем розпізнавання емоцій. Вона забезпечує необхідну компактність та високу швидкість інференсу, що є критично важливими вимогами для її успішної інтеграції в ПС на мобільних пристроях або вбудованих системах. Основний внесок роботи полягає в демонстрації успішного трансферного навчання, переходу від задачі детекції об'єктів до задачі класифікації емоцій з досягненням високої ефективності.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. A. Goel, R. Karim, U. Singh and R. Kumar, "A Review On Emotion Identification Using YOLO and DeepSORT," 2024 IEEE International Conference on Computing, Power and Communication Technologies (IC2PCT), Greater Noida, India, 2024, pp. 430-434, doi: 10.1109/IC2PCT60090.2024.10486695.
2. Ultralytics YOLOv8. URL: <https://yolov8.com/> (date of access: 25.09.2025).
3. AffectNet Dataset. URL: <https://www.kaggle.com/datasets/mstjebashazida/affectnet> (date of access: 30.09.2025).

ІНФОРМАЦІЙНА СИСТЕМА АНАЛІЗУ ТЕНДЕНЦІЙ ЗМІНИ КЛІМАТУ ЗА СУПУТНИКОВИМИ ДАНИМИ

Зміна клімату є однією з найгостріших глобальних проблем сучасності, що проявляється у підвищенні середньорічних температур, зростанні кількості екстремальних погодних явищ, таненні льодовиків та піднятті рівня Світового океану. Для кількісної оцінки цих процесів широко використовуються супутникові кліматичні дані, які забезпечують регулярне глобальне покриття та дозволяють виявляти довгострокові тенденції й здійснювати науково обґрунтоване прогнозування.

Розробка інформаційної системи аналізу тенденцій зміни клімату за супутниковими даними здійснюється з метою автоматизованої обробки, візуалізації та прогнозування температурних аномалій поверхні Землі. Для дослідження використано відкритий набір даних Global Surface Temperatures з платформи Kaggle, що містить багаторічні глобальні та регіональні значення температурних аномалій. У межах роботи застосовуються методи статистичного аналізу часових рядів, трендового аналізу та прогнозування, зокрема моделі ARIMA, ETS та регресійні підходи. Реалізація інформаційної системи здійснюється мовою програмування R з використанням сучасних засобів обробки та візуалізації даних.

Аналіз аналогічних досліджень показав, що сучасні наукові підходи до аналізу кліматичних змін базуються на використанні довготривалих рядів супутникових та наземних спостережень у поєднанні з методами статистичного аналізу й прогнозування. Зокрема, у роботі J. Hansen [1] на основі набору даних GISTEMP встановлено, що середня глобальна температура поверхні Землі зросла приблизно на $1,1^{\circ}\text{C}$ з кінця XIX століття, при цьому темпи потепління після 1980 року істотно прискорилися. У Шостому звіті Міжурядової групи експертів зі зміни клімату (IPCC) [2] на основі супутникових та наземних спостережень підтверджено, що кожне з останніх чотирьох десятиліть було теплішим за будь-які попередні з 1850 року, а прогнозоване потепління до 2100 року може сягнути від $1,5^{\circ}\text{C}$ до $4,4^{\circ}\text{C}$ залежно від сценарію викидів парникових газів. У дослідженні A. Benestad [3] методи регресійного аналізу та ARIMA використано для прогнозування глобальної

температури, де встановлено середній темп зростання близько $0,18-0,22^{\circ}\text{C}$ за десятиліття. Крім того, у роботі S. Deo [4] із використанням експоненційного згладжування та нейронних мереж доведено, що моделі часових рядів забезпечують високу точність прогнозу ($\text{RMSE} < 0,15^{\circ}\text{C}$) на інтервалі 10-20 років. Це підтверджує доцільність застосування методів аналізу часових рядів і прогнозування для дослідження тенденцій глобального потепління.

Для реалізації практичної частини розроблено інформаційну систему аналізу кліматичних тенденцій мовою R. Опрацьовано помісячні глобальні температурні аномалії за період 1880-2023 років, у яких виявлено та відновлено 7 пропущених значень ($0,405\%$) методом лінійної інтерполяції. Після очищення сформовано безперервний часовий ряд і виконано первинний розвідувальний аналіз із побудовою графіка температурних аномалій та їх ковзних середніх, що візуально підтвердило наявність довгострокового висхідного тренду (рис. 1).



Рисунок 1 – Середні річні глобальні температурні аномалії

У межах статистичного аналізу виконано агрегування даних за десятиліттями та роками, здійснено виявлення найхолоднішого та найтеплішого років за весь період спостережень. Встановлено, що найхолоднішим роком був 1909 рік із середньою температурною аномалією $-0,483^{\circ}\text{C}$, тоді як найтеплішим – 2020 рік із показником $1,021^{\circ}\text{C}$. Різниця між екстремальними значеннями становить $1,504^{\circ}\text{C}$, що кількісно відображає масштаб глобального потепління. Додатково виконано аналіз сезонності помісячних даних із побудовою boxplot-діаграм, графіка середніх значень за місяцями та сезонної декомпозиції STL, яка дозволила виокремити трендову, сезонну та випадкову складові температурного ряду.

Таблиця 1 – Результати аналізу глобальних температурних аномалій

Показник	Значення
Період спостережень	1880-01 – 2023-12
Кількість місячних спостережень	1728
Кількість пропущених значень (до інтерполяції)	7
Частка пропусків у часовому ряді (%)	0,405
Кількість десятиліть у вибірці	15
Зміна середньої температурної аномалії між першим та останнім десятиліттям (°C)	1,148
Найхолодніший рік	1909
Середня аномалія найхолоднішого року (°C)	-0,483
Найтепліший рік	2020
Середня аномалія найтеплішого року (°C)	1,021
Різниця між найтеплішим і найхолоднішим роком (°C)	1,504

Для кількісної оцінки тенденції зміни температури застосовано параметричні та непараметричні методи трендового аналізу. За результатами лінійної регресії встановлено, що середня глобальна температура зростає зі швидкістю $0,00791\text{ }^{\circ}\text{C}$ на рік, або $0,0791\text{ }^{\circ}\text{C}$ за десятиліття, при $p\text{-value} < 2,2 \cdot 10^{-16}$ та $R^2 = 0,7699$, що підтверджує статистично значущий характер тренду. Отриману залежність наочно відображає графік лінійного тренду глобальних температурних аномалій, побудований за річними середніми значеннями. Непараметричний тест Манна–Кендалла також підтвердив наявність сильного зростаючого тренду ($z = 12,809$, $\tau = 0,7208$), а робастна оцінка Sen's slope становить $0,00787\text{ }^{\circ}\text{C}$ на рік із довірчим інтервалом $[0,00700; 0,00869]$.

Розроблена інформаційна система аналізу тенденцій зміни клімату за супутниковими даними містить підсистему завантаження та зберігання кліматичних даних, підсистему попередньої обробки та очищення часових рядів, підсистему статистичного та трендового аналізу, підсистему візуалізації кліматичних показників, а також підсистему прогнозування температурних аномалій. Запропонована структура забезпечує автоматизовану обробку, наочну візуалізацію та науково обґрунтоване прогнозування температурних аномалій поверхні Землі.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Hansen, J., Ruedy, R., Sato, M., Lo, K. (2010). Global Surface Temperature Change. Reviews of Geophysics.
2. IPCC. (2021). Sixth Assessment Report (AR6): Climate Change 2021 – The Physical Science Basis. Intergovernmental Panel on Climate Change.

3. Benestad, A. E. (2013). Association between trends in daily rainfall percentiles and the global mean temperature. *Journal of Geophysical Research: Atmospheres*.

4. Deo, S., Şahin, M. (2015). Application of artificial neural networks for climate change prediction. *Atmospheric Research*.

УДК 004.89:004.932.2

Скиба О. М., Гожий О. П.

*Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ІНТЕЛЕКТУАЛЬНА СИСТЕМА ГЕНЕРАЦІЇ ЗОБРАЖЕНЬ ПО ТЕКСТОВОМУ ОПИСУ

Генерація зображень на основі текстового опису стала одним із найбільш актуальних напрямів розвитку сучасних систем штучного інтелекту, оскільки дозволяє автоматично створювати візуальний контент, інтерпретуючи природну мову. Такий підхід знаходить широке застосування у дизайні, медіавиробництві, науковій візуалізації, освітніх середовищах та автоматизації творчих процесів. Зростання можливостей дифузійних моделей та трансформерів актуалізує потребу в розробці інтелектуальних систем, здатних ефективно поєднувати мовні та візуальні представлення для отримання високоякісних зображень.

Метою розробки інтелектуальної системи генерації зображень за текстовим описом є створення програмного комплексу, який забезпечує обробку вхідних текстових даних, перетворення їх у векторні представлення, подальше генерування зображення за допомогою дифузійних моделей та можливість інтерактивної взаємодії користувача з результатами. Для досягнення цієї мети застосовано сучасні підходи глибинного навчання (Deep Learning), такі як трансформерні текстові енкодери, дифузійні моделі (Diffusion Models), варіаційні автокодери та модулі попередньої обробки зображень.

Аналіз наукових публікацій підтверджує актуальність і ефективність методів генеративного штучного інтелекту. Наприклад, автори OpenAI у роботі над моделлю DALL·E 2 [1] запропонували комбіновану архітектуру CLIP-текстового енкодера та дифузійного декодера, що забезпечила значне підвищення якості згенерованих зображень. У дослідженні від Google Research щодо Imagen [2] показано, що мовні моделі великого масштабу (Large Language Models, LLM) здатні

забезпечити високоточне семантичне кодування тексту, що прямо впливає на фотореалістичність результатів. Модель Stable Diffusion [3] продемонструвала ефективність латентних дифузійних процесів, дозволивши зменшити обчислювальні витрати та забезпечити можливість розгортання моделей на персональних комп'ютерах. Також дослідження Rombach et al. [4] підкреслює роль автокодерів у стисненні простору ознак без втрати ключової інформації, що робить генерацію більш масштабованою та швидкою.

Для навчання, тестування та перевірки працездатності інтелектуальної системи використовувалися відкриті датасети, такі як LAION-5B [5], що містить пару «опис – зображення» у великих обсягах та забезпечує різноманітність стилів і тематик. Додатково застосовано підмножини COCO Captions [6] і Conceptual Captions [7] для тестування узагальнювальної здатності текстового енкодера і якості відповідності між текстом і візуальними характеристиками. Сукупне використання кількох датасетів дозволило сформувати репрезентативне середовище для перевірки коректної роботи системи та стабільності моделі в реальних сценаріях.

Архітектура розробленої інтелектуальної системи побудована за модульним принципом, що забезпечує гнучкість, масштабованість та можливість розширення функціональності. Система складається з користувацького вебінтерфейсу (HTML, CSS, JavaScript), який дозволяє формувати запити, керувати чатами та переглядати результати; серверної частини на Flask (Python), що обробляє запити, запускає дифузійну модель та формує зображення; модулю текстового енкодера, який перетворює текст у латентний простір; дифузійного декодера, що виконує процес генерації; та додаткових модулів підвищення роздільності (super-resolution), кешування, управління сесіями та локального збереження результатів. Схема системи представлена на рис. 1.

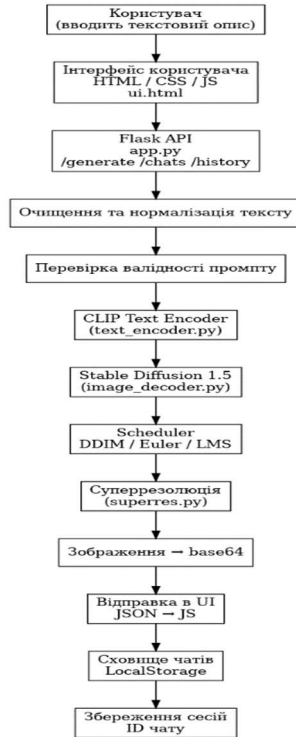


Рисунок 1 – Архітектура інтелектуальної системи генерації зображень за текстовим описом

Під час тестування системи було проведено оцінювання якості генерації зображень за критеріями семантичної відповідності, фотореалістичності, різноманітності та стабільності роботи моделі. Текстовий енкодер забезпечив коректне перетворення семантики запитів, мінімізуючи втрати змісту, а дифузійна модель стабільно генерувала зображення розміром 512×512 пікселів з високим рівнем деталізації. Було реалізовано механізми візуалізації прогресу генерації та систему багаточатовості, що дозволило підвищити зручність взаємодії користувача. Загалом експериментальні результати підтвердили ефективність розробленого підходу та відповідність системи поставленим вимогам щодо якості, продуктивності й інтерактивності роботи.

Таким чином, інтелектуальна система генерації зображень за текстовим описом демонструє високу функціональність і здатність

формувати реалістичні візуальні об'єкти на основі мовних запитів. Поєднання трансформерного текстового енкодера, дифузійного декодера та вебінтерфейсу забезпечує повний цикл обробки — від введення тексту до отримання готового зображення. Запропонована система може бути інтегрована у творчі, навчальні та виробничі процеси, виступаючи універсальним інструментом генеративного штучного інтелекту.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Goodfellow I., Bengio Y., Courville A. Deep Learning. Cambridge : MIT Press, 2016. 800 p.
2. Vaswani A., Shazeer N., Parmar N. et al. Attention Is All You Need // Advances in Neural Information Processing Systems. 2017. Vol. 30. URL: <https://arxiv.org/abs/1706.03762>
3. Rombach R., Blattmann A., Lorenz D. et al. High-Resolution Image Synthesis with Latent Diffusion Models // Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). 2022. P. 10684–10695.
4. Saharia C., Chan W., Saxena S. et al. Photorealistic Text-to-Image Diffusion Models with Deep Language Understanding // arXiv preprint. 2022. URL: <https://arxiv.org/abs/2205.11487>
5. Kingma D., Welling M. Auto-Encoding Variational Bayes // arXiv preprint. 2014. URL: <https://arxiv.org/abs/1312.6114>
6. Mirza M., Osindero S. Conditional Generative Adversarial Nets // arXiv preprint. 2014. URL: <https://arxiv.org/abs/1411.1784>
7. Karras T., Laine S., Aittala M. et al. Analyzing and Improving the Image Quality of StyleGAN // Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. 2020. P. 8107–8116.
8. Ho J., Jain A., Abbeel P. Denoising Diffusion Probabilistic Models // Advances in Neural Information Processing Systems. 2020. Vol. 33. P. 6840–6851.
9. Radford A., Kim J., Hallacy C. et al. Learning Transferable Visual Models From Natural Language Supervision (CLIP) // Proceedings of the 38th International Conference on Machine Learning. 2021. URL: <https://arxiv.org/abs/2103.00020>
10. Paszke A., Gross S., Massa F. et al. PyTorch: An Imperative Style, High-Performance Deep Learning Library // Advances in Neural Information Processing Systems. 2019. Vol. 32.

СИСТЕМА ІНТЕРПРЕТАЦІЇ МЕДИЧНИХ ДІАГНОСТИЧНИХ МОДЕЛЕЙ НА ОСНОВІ МЕТОДІВ ХАІ

Впровадження систем штучного інтелекту (ШІ) у медичну практику відкриває значні перспективи для підвищення якості діагностики, проте стикається з проблемою «чорного ящика» — непрозорості прийняття рішень складними моделями глибокого навчання. Відсутність інтерпретованості обмежує довіру лікарів до автоматизованих систем та ускладнює їх сертифікацію відповідно до сучасних регуляторних вимог (EU AI Act). Актуальність роботи полягає у необхідності створення діагностичних інструментів, які поєднують високу точність прогнозування зі зрозумілим поясненням результатів.

Розробка інтелектуальної системи медичної діагностики з пояснюваним штучним інтелектом (eXplainable AI, XAI) здійснюється з метою забезпечення прозорості та обґрунтованості діагностичних рішень при аналізі мультимодальних медичних даних. Для досягнення поставленої мети використано методи машинного навчання (Logistic Regression, Random Forest, XGBoost) для табличних даних та глибокого навчання (ResNet18) для аналізу зображень. Інтеграція алгоритмів XAI, таких як SHAP (SHapley Additive exPlanations), LIME (Local Interpretable Model-agnostic Explanations) та Grad-CAM (Gradient-weighted Class Activation Mapping), дозволила генерувати візуальні та текстові пояснення, що інтерпретують вплив окремих клінічних ознак та зон на зображеннях на кінцевий діагноз.

Аналіз сучасних досліджень підтверджує ефективність гібридних підходів до медичної діагностики. Так, роботи A. Rajkomar et al. [1] демонструють переваги глибокого навчання в аналізі електронних медичних карток, проте відзначають складність інтерпретації. Дослідження R. R. Selvaraju et al. [2], присвячені методу Grad-CAM, доводять можливість візуалізації уваги згорткових мереж, що є критичним для радіології. Водночас, S. M. Lundberg et al. [3] обґрунтовують використання значень Шеплі для уніфікованої оцінки важливості ознак. Це підтверджує доцільність поєднання різних методів ХАІ в єдиній системі для забезпечення всебічної інтерпретації.

Для навчання та оцінювання ефективності розробленої системи було використано відкриті верифіковані набори даних. Для задачі

діагностики раку грудей використано «Breast Cancer Wisconsin (Diagnostic) Data Set» [4] з репозиторію UCI ML, що містить 569 записів із 30 числовими характеристиками клітинних ядер. Для задачі виявлення пневмонії застосовано «Chest X-Ray Images (Pneumonia)» [5] з платформи Kaggle, який включає 5863 рентгенівських знімків грудної клітки. Використання цих даних дозволило навчити моделі з високою узагальнювальною здатністю та протестувати алгоритми пояснення на реальних клінічних випадках.

Архітектура розробленої системи базується на принципі «XAI Orchestrator», що забезпечує модульність та гнучкість обробки даних різних модальностей. Система реалізована мовою Python з використанням фреймворку Streamlit для створення вебінтерфейсу. Вона включає модулі завантаження та попередньої обробки даних (нормалізація, аугментація), модуль діагностичного ядра (навчені моделі ML/DL), модуль генерації пояснень (SHAP, LIME, Grad-CAM) та модуль візуалізації. Така структура дозволяє лікарю не лише отримати ймовірнісний прогноз захворювання, але й інтерактивно дослідити причини такого висновку через візуалізацію теплових карт та графіків важливості ознак (див. рис. 1).

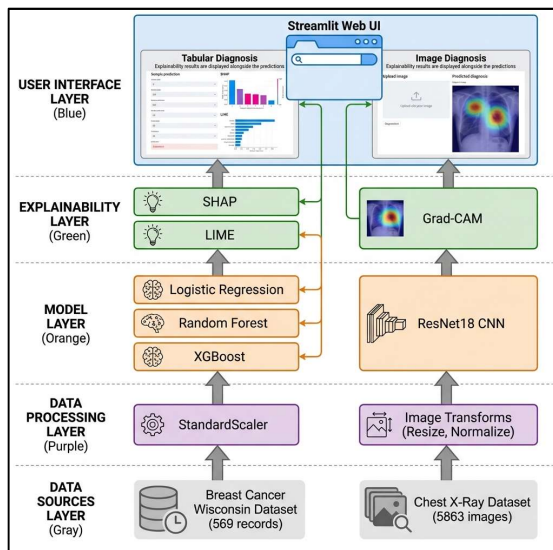


Рисунок 1 – Архітектура системи пояснюваної медичної діагностики

У процесі експериментальних досліджень було оцінено точність моделей та якість пояснень. Ансамбль моделей для табличних даних досяг точності 97% на тестовій вибірці, при цьому методи SHAP та LIME коректно ідентифікували ключові маркери злоякісності (наприклад, worst concave points). Згорткова нейромережа ResNet18 для аналізу рентгенівських знімків продемонструвала точність понад 90%. Застосування Grad-CAM дозволило виявити випадки «навчання на артефактах» (shortcut learning), коли модель фокусувалася на кісткових структурах замість легеневої тканини, що підтверджує критичну важливість ХАІ для валідації медичних моделей перед їх впровадженням.

Таким чином, отримані попередні результати підтверджують перспективність обраного підходу до побудови інтелектуальної системи медичної діагностики. Поеднання алгоритмів прогнозування з методами пояснюваного штучного інтелекту створює основу для забезпечення необхідного рівня прозорості. Подальша робота буде зосереджена на вдосконаленні архітектури, розширенні набору навчальних даних та проведенні детальної клінічної апробації прототипу системи, що є необхідним етапом перед її повноцінним впровадженням у медичну практику.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Rajkomar A., Dean J., Kohane I. Machine Learning in Medicine. New England Journal of Medicine. 2019. Vol. 380. P. 1347–1358. DOI: 10.1056/NEJMr1814259.
2. Selvaraju R. R. et al. Grad-CAM: Visual Explanations from Deep Networks via Gradient-based Localization. ICCV. 2017. P. 618–626. DOI: 10.1109/ICCV.2017.74.
3. Lundberg S. M., Lee S.-I. A Unified Approach to Interpreting Model Predictions. NIPS. 2017. P. 4765–4774.
4. Breast Cancer Wisconsin (Diagnostic) Data Set. UCI Machine Learning Repository. URL: [https://archive.ics.uci.edu/ml/datasets/Breast+Cancer+Wisconsin+\(Diagnostic\)](https://archive.ics.uci.edu/ml/datasets/Breast+Cancer+Wisconsin+(Diagnostic)).
5. Chest X-Ray Images (Pneumonia). Kaggle. URL: <https://www.kaggle.com/paultimothymooney/chest-xray-pneumonia>.

ІНТЕЛЕКТУАЛЬНА СИСТЕМА КЕРУВАННЯ ДОСТАВКОЮ ВАНТАЖІВ В ДИНАМІЧНОМУ СЕРЕДОВИЩІ ЗА ДОПОМОГОЮ БПЛА НА БАЗІ UNITY

У роботі розглядається концепція інтелектуальної системи автоматизованого керування доставкою вантажів за допомогою безпілотних літальних апаратів (БПЛА) у динамічному та невизначеному середовищі. Система реалізована у середовищі Unity 2022.3 та поєднує модулі навігації, сенсорики, прогнозування зовнішніх змін та оптимізації маршрутів. Описано архітектуру рішення, що включає модулі симуляції фізичного середовища, підсистему планування траєкторій, систему уникнення перешкод, блок керування флотом дронів, підсистему моніторингу та аналітики, а також інтерфейс оператора для відстеження та ручної корекції польотів.

Ефективне планування шляху є важливим елементом для операцій БПЛА, оскільки воно дозволяє БПЛА розрахувати найкращий можливий шлях від його поточного місцезнаходження до цільового пункту призначення, одночасно уникаючи перешкод та дотримуючись експлуатаційних обмежень [1]. На рис. 1 наведено блок-схему процесу планування траєкторії БПЛА в різних умовах. Як вказано на схемі, початок процесу планування шляху включає збір даних про навколишнє середовище за допомогою датчиків та карт, які формують основу для навігаційних рішень.

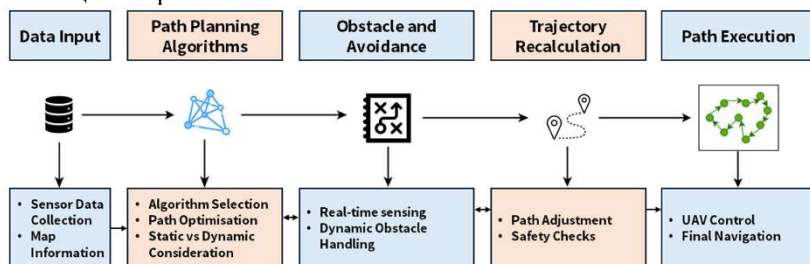


Рис.1 Блок-схема процесу планування траєкторії БПЛА в різних умовах [2]

Після цього БПЛА використовує алгоритми планування шляху для обчислення та оптимізації найефективнішого маршруту, враховуючи статичний або динамічний характер навколишнього середовища.

Уникнення та виявлення перешкод у режимі реального часу дозволяє БПЛА зосереджуватися на своєму оточенні та змінювати курс за потреби, щоб уникнути зіткнень. Виявляючи перешкоди або зміни умов, БПЛА коригує свою траєкторію, щоб забезпечити безпеку та оптимально підвищити продуктивність. Нарешті, БПЛА дотримується запланованого маршруту, постійно адаптуючись до будь-яких неочікуваних відхилень. Ця процедура забезпечує надійну, ефективну та безпечну навігацію на складній місцевості, за непередбачуваних обставин та з експлуатаційними обмеженнями[2].

Пропонована система побудована на модульній архітектурі, що забезпечує незалежний розвиток компонентів та можливість масштабування від локальних симуляцій до мережових багатодронових сценаріїв. Unity використовується як високоточне середовище візуалізації та фізичного моделювання, забезпечуючи роботу з тривимірною сценою, ореалістичними моделями колізій і можливістю симуляції погодних умов, вітру та динамічних перешкод. Підсистема планування маршруту базується на поєднанні алгоритмів A^* , D^* , та локальних методів ухилення[3], що забезпечує оптимальний компроміс між глобальною оптимізацією шляху та швидкою реакцією на зміни середовища. Такий підхід відповідає сучасним тенденціям у дослідженнях UAV-навігації: огляд останніх робіт підкреслює важливість комбінації глобального планування шляху та локальних методів уникнення перешкод в умовах динамічних і складних міських середовищ [2]. Передбачено підтримку «динамічної репланівки», коли дрон переглядає маршрут у режимі реального часу за умови появи нових перешкод — інших БПЛА, транспортних засобів, змін рельєфу чи погодних факторів. Аналогічні підходи, які дозволяють реалізувати керування в реальному часі з корекцією маршруту, описані в роботі [4].

Ключовою особливістю запропонованої системи є механізм динамічної перебудови маршруту в режимі реального часу. Завдяки постійному оновленню карти середовища, БПЛА здатний адаптуватися до раптових змін, наприклад, появи рухомих перешкод.

У прототипі даної системи реалізовано тестовий сценарій доставки вантажів у міському середовищі з моделюванням руху транспорту, зміною погоди та появою нових перешкод. Протягом експериментів оцінювалася якість ухилення від перешкод, стабільність маршруту та ефективність керування флотом у випадках конфліктів маршрутів. Отримані результати симуляційних експериментів демонструють, що система спроможна забезпечити автономне планування та стабільно уникати зіткнень за умов середньої щільності перешкод та скорочувати час доставки порівняно з базовими алгоритмами без динамічної

репланівки. У свою чергу, аналогічні результати сучасних досліджень показують, що при правильній організації avoidance- та replanning-механізмів можна досягти високої безпеки польотів навіть у щільному міському просторі [5].

Отримані результати підтверджують, що інтелектуальна система керування БПЛА, реалізована на базі рушія Unity, повністю здатна забезпечити автономну, безпечну та ефективну доставку вантажів у динамічному середовищі. Зокрема розроблена інтелектуальна система керування доставкою вантажів за допомогою БПЛА на базі Unity є перспективним рішенням для виконання міських логістичних задач, оскільки поєднує адаптивне планування, реалістичне моделювання, масштабованість та можливість безпечної взаємодії кількох дронів. У подальших дослідженнях планується розглянути алгоритми кооперативної навігації, урахування зміни погодних умов в процесі моделювання та інтеграцією у запропоновану систему методів машинного навчання для прогнозування ризиків у динамічних середовищах.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Debnath, D., Hawary, A.F., RaNmndan, M.I., Alvarez, F.V., Gonzalez F. QuickNav: An Effective Collision Avoidance and Path-Planning Algorithm for UAS. Drones. 2023, Vol.7(11), 678.
2. Debnath D., Vanegas F., Sandino J., Hawary A. F., Gonzalez F. A Review of UAV Path-Planning Algorithms and Obstacle Avoidance Methods for Remote Sensing Applications. Remote Sensing. 2024, Vol. 16(21), 4019.
3. Dovbysh I., Muraviov O., Momot A.; Bohdan H. Autonomous UAV navigation: technologies for orientation and localization. Applied questions of mathematical modeling. 2025, Vol. 8, 1, p. 57-64.
4. Bernevek I., Yaremko O. Investigation of unmanned aircraft autopiloting methods with real time route correction. ictet. 2024, Vol. 5, 2, pp.49-58.
5. Deng M., Yang Q., Peng Y. A Real-Time Path Planning Method for Urban Low-Altitude Logistics UAVs. Sensors (Basel). 2023, Vol. 23(17), 7472.

ІНТЕЛЕКТУАЛЬНА СИСТЕМА КАРТОГРАФУВАННЯ ВИБУХОНЕБЕЗПЕЧНИХ ПРЕДМЕТІВ НА ОСНОВІ ДАНИХ БПЛА ТА ГЛИБИННОГО НАВЧАННЯ

Інтенсивне мінування територій України від час повномасштабної війни призвело до формування великих площ потенційно небезпечних земель, які потребують обстеження та гуманітарного розмінування. За даними ДСНС, орієнтовно 156 тис. км² території країни є потенційно забрудненими вибухонебезпечними предметами (ВНП) [1]. Це ставить перед операторами протимінної діяльності задачу не лише оперативного виявлення ВНП, а й побудови точних, стандартизованих карт небезпечних зон, що відповідають вимогам міжнародних стандартів IMAS та можуть інтегруватися у національні інформаційні системи менеджменту протимінної діяльності.

Метою роботи є розробка інтелектуальної системи картографування вибухонебезпечних предметів на основі даних безпілотних літальних апаратів (БПЛА) та результатів глибинного навчання. Система поєднує автоматизоване виявлення ВНП за знімками з низьким перетином [2] за допомогою сервісів глибинного навчання, та подальшу просторову аналітику в середовищі ГІС. Запропонований підхід дозволяє перетворювати множину точкових детекцій глибинної моделі [2] на структуровані полігональні представлення небезпечних зон, придатні для використання у плануванні операцій розмінування та створенні формалізованих звітів.

Вхідні дані формуються під час послідовного UAV-обстеження ділянки. Спочатку виконується ручний політ-розвідка з метою оперативного огляду місцевості, уточнення меж ділянки інтересу та фіксації потенційних ВНП. Кожен кадр містить EXIF-координати, що забезпечує базову геоприв'язку. Далі для тієї ж території виконується зйомка з підвищеним перетином, яка слугує основою для побудови високоточної ортомозаїки в ArcGIS Pro [5].

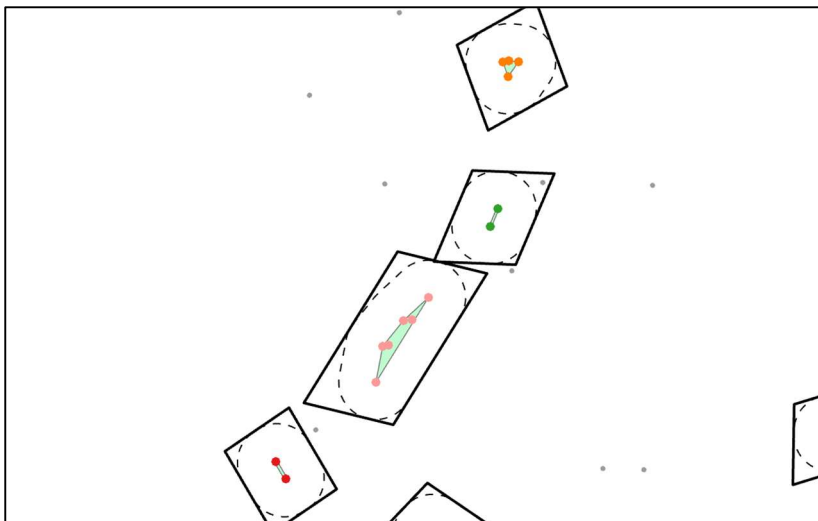


Рисунок 1 – Докази забруднення та етапи обробки в ArcGIS Pro (точки, кластери, опукла оболонка, буфер, фінальний прямокутник)

Для забезпечення коректності аналізу точковий шар попередньо очищується інструментом Definition Query в ArcGIS Pro, що дозволяє залишити лише точки видимих ВНП та релевантних доказів. Після цього застосовується просторове групування точок алгоритмом DBSCAN (Density-Based Spatial Clustering of Applications with Noise), реалізованим у вигляді інструменту Density-based Clustering. Кожен виділений кластер інтерпретується як окрема небезпечна ділянка; для нього обчислюються площа, кількість доказів та інші атрибутивні характеристики за допомогою інструменту Summarize Within, а на основі множини точок будується опукла оболонка (convex hull), що утворює базовий полігон небезпечної зони.

Для врахування запасу безпеки опуклі полігони буферизуються на задану відстань, у результаті чого формується згладжений контур небезпечної території. Наступним кроком є побудова мінімальних орієнтованих прямокутників, що охоплюють буферизовані полігони. Такі прямокутні контури спрощують опис меж (координати поворотних точок, азимуту, довжини сторін) та полегшують відтворення зон на місцевості під час маркування. Для кожного прямокутника автоматично обчислюються геометричні характеристики (MBG_Length, MBG_Width, MBG_Orientation), які можуть використовуватися як у

подальшій аналітиці (порівняння розмірів зон, аналіз напрямку мінних смуг), так і при генерації звітної документації.

Уся просторово-атрибутивна інформація зберігається у файловій геобазі FGDB, інтегрованій з ArcGIS Pro та ArcGIS Online, що забезпечує публікацію результатів у вигляді вебкарт та вебшарів. ArcGIS Survey123 використовуються для формування звітів за шаблонами. Таким чином реалізовано повний технологічний цикл – від збору польових даних із БПЛА та глибинних детекцій, аналітичної обробки, виділення небезпечних зон прямокутної форми та підготовки документів, сумісних із вимогами IMAS та національної системи IMSMA Core. Запропонована інтелектуальна система підвищує оперативність і відтворюваність картографування ВМП і може бути використана операторами гуманітарного розмінування в Україні як практичний інструмент підтримки прийняття рішень.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. 156 тисяч квадратних кілометрів території України наразі потенційно небезпечні та потребують обстеження й розмінування – ДСНС – Медіацентр Україна. URL: <https://mediacenter.org.ua/uk/156-tisyach-kvadratnih-kilometriv-teritoriyi-ukrayini-narazi-potentsijno-nebezpechni-ta-potrebuyut-obstezhennya-j-rozminuvannya-dsns/> (дата звернення: 18.10.2025).

2. Чупина В. Є. Створення інтерактивної мапи вибухонебезпечних предметів на основі даних з безпілотних апаратів : кваліфікаційна робота. 2024. 81 с.

3. Artificial intelligence. SafePro AI. URL: <https://safeproai.com/artificial-intelligence-landmine-detection-through-drones/> (дата звернення: 18.10.2025).

4. NPA Ukraine conducts pilot project with Safe Pro AI for automated processing of drone imagery. Norwegian People's Aid. URL: <https://www.npaid.org/mine-action-and-disarmament/news/npa-ukraine-conducts-pilot-project-with-safe-pro-ai-for-automated-processing-of-drone-imagery/> (дата звернення: 18.10.2025).

5. ArcGIS Pro. URL: <https://www.esri.com/en-us/arcgis/products/arcgis-pro/overview> (Last accessed: 18.10.2025).

ІНФОРМАЦІЙНА СИСТЕМА З ОЦІНКИ ГОТОВНОСТІ СТАРТАПУ ДО МАСШТАБУВАННЯ

Проблема прийняття управлінських рішень щодо переходу стартапу від стадії пошуку бізнес-моделі до стадії активного зростання є критичною в інноваційному менеджменті. Згідно зі звітом Startup Genome, передчасне масштабування (Premature Scaling) є причиною краху близько 74% технологічних проєктів [1]. Існуючий інструментарій характеризується фрагментарністю та переважним використанням лінійних чек-листів, які ігнорують специфіку бізнес-моделі та структурні дисбаланси. Метою дослідження є підвищення обґрунтованості управлінських та інвестиційних рішень щодо масштабування стартапів шляхом розроблення інформаційної системи з оцінки готовності стартапу до масштабування, яка інтегрує сучасні методи та алгоритми багатокритеріального аналізу (АНР, TOPSIS), а також алгоритми логічного виведення для автоматизованого виявлення ризикових патернів та структурних дисбалансів у розвитку стартапу.

Методологічну основу системи «VentureGrade» становить гібридна модель багатокритеріального оцінювання.

Для визначення вагових коефіцієнтів критеріїв реалізовано метод ієрархічної декомпозиції Сааті (АНР). Система забезпечує адаптивність оцінювання через два режими: ручне калібрування пріоритетів (матриці парних порівнянь) або використання «Експертних профілів» (Expert Profiles) — наборів ваг для типових моделей (B2B SaaS, Marketplace, DeepTech), сформованих на основі галузевих бенчмарків.

Для розрахунку інтегрального індексу готовності застосовано модифікований метод TOPSIS. На відміну від класичного підходу, у роботі впроваджено елементи некомпенсаторної логіки. Стандартні лінійні моделі дозволяють високим показникам за одним критерієм (наприклад, «Фінанси») нівелювати критичні провали в іншому (наприклад, «Product-Market Fit»). У розробленій системі імплементовано механізм «вето-порогов»: при падінні значень критичних метрик нижче допустимого рівня активується штрафний коефіцієнт, що суттєво знижує загальний рейтинг готовності, наближаючи математичну модель до реальної логіки венчурного інвестора.

Інтелектуальна складова системи посилена модулем аналізу ризиків, який базується на продукційній базі знань. Система в автоматичному режимі виконує діагностику структурних дисбалансів розвитку, застосовуючи набір евристичних правил логічного виведення. Наприклад, детекція патерну «Пастка Інженера» (Technology Push) формалізується через умову суттєвого перевищення показників технологічної зрілості над показниками ринкової валідації.

Архітектура системи реалізована за модульним принципом на мові Python (див. рис. 1). Інкапсуляція математичної логіки (AHP/TOPSIS) та продукційних правил виконана в окремих модулях бекенду. Взаємодія з користувачем здійснюється через кросплатформний інтерфейс Telegram.

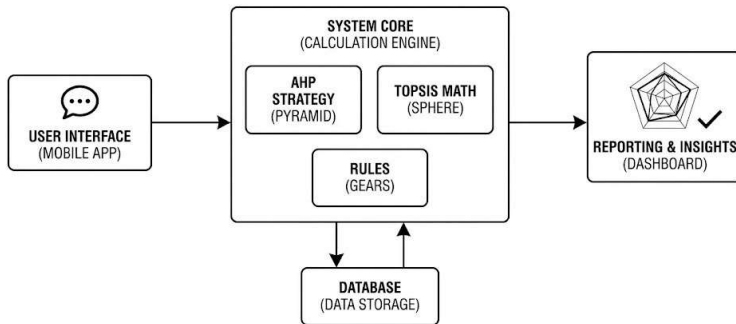


Рисунок 1 – Архітектура інформаційної системи оцінки стартапів

Для валідації розробленої моделі проведено ретроспективний аналіз (Back-testing) на історичних даних відомих компаній станом на момент прийняття ними рішення про масштабування.

Для перевірки адекватності розробленої моделі було проведено експериментальне тестування методом ретроспективного аналізу (Back-testing). Суть методу полягає у введенні в систему історичних даних відомих стартапів станом на момент прийняття ними рішення про масштабування та порівнянні прогнозу системи з реальними наслідками.

Як тестові сценарії було обрано два контрастні кейси:

1. Startup A (Quibi, 2020) — Кейс провалу. На момент запуску проєкт мав аномально високі фінансування (\$1.75 млрд), але низький рівень валідації ринкової потреби.

2. Startup B (Grammarly, Round A) — Кейс успіху. Проєкт характеризувався органічним зростанням, високим утриманням користувачів та збалансованими показниками.

Результати моделювання наведено в таблиці 1.

Таблиця 1 – Результати ретроспективного аналізу

Параметр	Startup A (Quibi)	Startup B (Grammarly)
Вхідні дані (шкала 1-10)	Фінанси: 10, Команда: 10, Ринок: 2	Фінанси: 6, Команда: 8, Ринок: 9
Профіль стратегії	"Media Platform" (Акцент на контент)	"SaaS / DeepTech" (Акцент на продукт)
TOPSIS Score (Результат)	0.54 (54%)	0.82 (82%)
Статус системи	NEEDS OPTIMIZATION	READY TO SCALE
Діагноз системи	«Пастка ресурсів»: Високий бюджет при відсутності PMF	Дисбалансів не виявлено. Рекомендовано масштабування

Система коректно ідентифікувала ризики для Startup A. Попри ідеальні фінансові показники, алгоритм TOPSIS, зважений через матриці АНР, суттєво знизив загальний рейтинг через критично низьку оцінку за критерієм «Ринок». Спрацювало продукційне правило модуля аналітики, що вказало на ризик передчасного масштабування. Для Startup B система показала високу готовність, що історично підтвердилося успішним зростанням компанії.

Таким чином, розроблена інформаційна система дозволяє проводити експрес-діагностику стартапів з використанням науково обґрунтованих методів. Поєднання АНР та TOPSIS забезпечує гнучкість налаштування під різні бізнес-моделі, а інтеграція продукційної експертної системи надає інтерпретовані рекомендації, що знижує ризики інвестиційних втрат.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Startup Genome. The Global Startup Ecosystem Report 2024. San Francisco : Startup Genome LLC, 2024. 350 p.
2. Behzadian M., Otaghsara S. K., Yazdani M., Ignatius J. A state-of-the-art survey of TOPSIS applications. Expert Systems with Applications. 2012. Vol. 39, No. 17. P. 13051–13069.
3. Saaty T. L. Decision Making for Leaders: The Analytic Hierarchy Process for Decisions in a Complex World. Pittsburgh : RWS Publications, 2008. 323 p.

UDC: 004.89

*Shvets Y.O., Malakhov E.V.
Odesa I.I.Mechnykov National University*

DYNAMIC RECONFIGURATION AND TASK REDISTRIBUTION TECHNOLOGY

The application of swarm complexes of unmanned aerial vehicles (UAVs) is already a generally accepted standard in many high-tech and critical spheres, including emergency relief, scanning of hard-to-access territories, and monitoring of dangerous zones where human involvement can be mortally dangerous. However, modern demands for mission speed and adaptability require a departure from homogeneous swarms, where the identity of the units limits the functional breadth of the system and leads to excessive duplication of effort [1]. Instead, the focus shifts to heterogeneous swarms, where drones have different functional characteristics, such as different speed, sensor resolution, or scanning width [2]. The main task of such systems is the optimization of territory scanning, which consists in ensuring maximum coverage with minimum time and resources. This, in turn, creates a need for the development of intelligent algorithms that can promptly change the configuration of the swarm and its individual units in accordance with environmental conditions and the nature of the operations being performed [3].

For achieving maximum efficiency and resilience to failures, dynamic reconfiguration of the heterogeneous swarm is critically important, as it will allow for the elimination of effort duplication, prevent a decrease in observation quality, and avoid an increase in mission time. The development and use of Artificial Intelligence (AI)-based reconfiguration technology is currently a promising approach.

Such an approach is aimed at not merely choosing the swarm's configuration at the beginning of the task, but also at ensuring its operational change during the execution of assigned tasks with the redistribution of tasks among the swarm nodes. Moreover, there should be a possibility for task redistribution in case of situations involving a change in the number of swarm nodes.

The goal of this technology is to ensure dynamic stability and adaptive response to changing conditions. The most important scenario for such reconfiguration arises when an individual drone completes its assigned route, but the global goal of territory coverage (for example, 95%) remains unachieved.

In this situation, the drone automatically transitions to a “waiting” state and initiates a request for a new task. This process is fully managed by a priority redistribution system based on a profitability analysis. For this purpose, this technology implements and maintains a pool of residual tasks, which is formed from territory segments whose coverage percentage remains below a defined critical threshold.

The selection of a new task from the pool is neither random nor sequential but constitutes a two-stage optimization process, consisting of a hard check for feasibility and a soft optimization based on cost. The first stage is a hard check for energy feasibility. Before a task is considered a candidate, the drone performs a careful calculation of the full energy required for its completion. This energy cost is calculated based on the energy consumption coefficients for transit and scanning modes, as well as the total length of the scanning path. The task is considered unfeasible and rejected if the necessary energy exceeds the drone's current energy reserve, taking into account a defined reserve factor.

The second stage is the priority metric, where from all physically feasible tasks, the one that minimizes the complex cost metric is chosen. This metric weighs four key factors: distance cost, which is a direct penalty for transit and prioritizes closer tasks; relative energy expenditure cost, which penalizes tasks that consume a large fraction of the remaining energy; mismatch cost, which penalizes drones for performing tasks that do not match their scanning width, thereby prioritizing the use of the heterogeneous swarm; and priority/urgency cost, which allows the system to weigh the importance of the least-scanned zones. The task selection is determined as the argument that minimizes this function.

To ensure coordinated work distribution and avoid duplication, the drone, upon selection, immediately employs a locking mechanism – it “locks” this task in the common pool and broadcasts this status to all other members of

the swarm. This guarantees effective and coordinated distribution of the residual work.

As a result, the proposed technology allows for the implementation of an autonomous priority task redistribution mechanism. It leads to the minimization of human intervention, an increase in swarm autonomy, and ensures its adaptability to new challenges in dynamic environments. Thanks to the locking and status broadcasting mechanism, coordinated distribution of residual work among all free agents is ensured.

Literature

1. Albrekht, Y., Pysarenko, A. Exploring the power of heterogeneous UAV swarms through reinforcement learning. *Technology Audit and Production Reserves*, 6 (2 (74)), 6–10, 2023. DOI: <https://doi.org/10.15587/2706-5448.2023.293063>

2. Kengyel, Daniela & Hamann, Heiko & Zahadat, Payam & Radspieler, Gerald & Wotawa, Franz & Schmickl, Thomas. Potential of Heterogeneity in Collective Behaviors: A Case Study on Heterogeneous Swarms. 2015. DOI: https://10.1007/978-3-319-25524-8_13.

3. Kozlov M.S., Shvets Y.O. Heterogeneous swarm complex multilevel reconfiguration based on Swarm Chemistry. *Applied Aspects of Information Technology*, 8:147-161, 2025. DOI: <https://doi.org/10.15276/aait.08.2025.10>

УДК 004.94

Шиян С.І.

*Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ОСОБЛИВОСТІ ПОБУДОВИ ІНТЕЛЕКТУАЛЬНОЇ СИСТЕМИ КОНТРОЛЮ ПАРАМЕТРІВ ПММ

У часи коли ціни на паливо коливаються, а фактичне споживання стає критичним показником, фахівцям важливо перевіряти рівень паливно-мастильних матеріалів (ПММ) в системах їх зберігання. Незалежно від того, чи це транспортний засіб, резервуар для зберігання, чи ізольований об'єкт, хибні та помилкові показники датчика або відсутність достовірної інформації можуть призвести до поломок, відмов, затримок або навіть прихованих витрат.

Системи моніторингу резервуарів (TMS) – це системи, які використовується для дистанційного моніторингу рівнів, стану та інших параметрів резервуарів для зберігання рідин ПММ або газу в режимі

реального часу. Ця прецизійна технологія забезпечує чітке уявлення про контрольоване споживання рідини/газу протягом року та виявляє будь-які закономірності використання. Це часто призводить до зниження експлуатаційних витрат на АЗС та в інших системах зберігання, покращення управління запасами, підвищення ефективності та усуває ризик простою. Завдяки сучасним системам контролю можливо проводити точні вимірювання рівня палива, автоматизувати моніторинг та виявляти будь-які аномалії: витік, крадіжку або несправність датчика.

При побудові інтелектуальної системи контролю ПММ необхідно враховувати наступні особливості і проблеми:

- Експлуатаційні витрати та неконтрольовані втрати. Втрати палива виникають з кількох джерел:
 - Витоки через знос або поломку бака
 - Ручне відкачування або автоматична крадіжка під час паркування.
 - Перевитрата палива через погане обслуговування двигуна або неефективне керування.
 - Людські помилки під час зчитування або введення вимірних рівнів.

Відсутність надійної системи контролю та моніторингу рівня ускладнює виявлення цих розбіжностей. Результат: витрати, які залишаються непоміченими, спожиті обсяги, які погано враховуються, та неадекватний контроль витрат на паливо. На основі дослідження транспортного та логістичного сектору в 2024 році на ПММ припадає до 35% логістичного бюджету професійного автопарку.

Окрім економічних аспектів, неефективне управління ПММ створює екологічну проблему. Невиявлені витоки, перевитрата через несправності двигуна або переповнення можуть призвести до серйозних пошкоджень, непотрібних викидів CO₂ та ризиків забруднення ґрунту (особливо у разі переповнення бака). Крім того, певні нормативні акти заохочують (або вимагають) від компаній вимірювати та зменшувати свій вуглецевий слід.

Це особливо стосується промислових секторів, на які поширюються такі стандарти, як:

- Екологічний менеджмент ISO 14001
- SEQE-EU (Європейська система вуглецевих квот)
- Місцеві податкові пільги для технологій моніторингу енергії

Автоматизуючи точне вимірювання об'єму палива, компанії можуть не тільки обмежити втрати, але й підвищити свої зобов'язання щодо охорони навколишнього середовища.

Сучасна система контролю та моніторингу палива складається на основі вимірювальних датчиків (ультразвукові, тискові, манометри, шина CAN, тощо). Ефективна система моніторингу рівня палива залежить, перш за все, від вірного вибору адаптованого датчика. Кожна технологія має свої переваги з точки зору точності, вартості, обсягу та умов монтажу. В таблиці 1 представлено особливості сучасних систем контролю і моніторингу ПММ.

Таблиця 1. Особливості сучасних систем контролю ПММ

Принцип роботи	Принцип вимірювання	Точність	Місце розташування
Ультразвук	Звукова хвиля, відбита рідиною	Достатня	Поза контактом (у верхній частині резервуара)
Датчики тиску	Різниця в тиску рідини	Дуже добра	Фіксований на дні резервуара
Радіохвилі	Радарна хвиля, відбита поверхнею рідини	Дуже добра	Безконтактний (верхня частина резервуара)
Радіометричний рівень	Зміна напруги пропорційна вимірюваному рівню (часто на основі змінного резистора, підключеного до поплавця)	Дуже добра	Кріпиться безпосередньо до контролюючого пристрою
Аналоговий рівень	Поплавок, з'єднаний з градуйованим стрижнем	Варіюється	Інтегровано в обладнання
Емкісний замір	Електрична ємність змінюється з рівнем рідини	Від доброї до самої точної	Занурення в рідину
CAN bus	Дані, отримані через електроніку автомобіля	Дуже висока	Без фізичних модифікацій

Підсумком є те, що інтелектуальна система контролю і моніторингу ПММ для промислових систем збереження ПММ повинна на системній методології об'єднувати різні технології та програмно-апаратні рішення та враховувати особливості експлуатації та сучасні екологічні вимоги. Подальші дослідження доцільно зосередити на розробці інтелектуальних систем обробки даних на основі технологій машинного навчання.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1.Reynolds, G., (2009) "The reduction of GHG emissions from shipping – A key challenge for the industry", World Maritime Technology Conference (WMRC), Mumbai, January 21-24.

2. Kim, A-R. and Seo, Y-J., (2019) "The reduction of SOx emissions in the shipping industry: The case of Korean companies", Marine Policy, Vol. 100, pp. 98-106.

3. DNV, (2010) "Assessment of measures to reduce future CO2 emission from shipping", Research and Innovation, Position paper No. 5, Det Norsk Veritas, Hovik.

4 .Yuan, J. and Nian, V., (2018) "Ship Energy Consumption Prediction with Gaussian Process Metamodel", Energy Procedia, Vol. 152, pp. 655-660

УДК 004.42

Яслик О. М.

*Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ІНТЕЛЕКТУАЛЬНА СИСТЕМА ПЕРВИННОЇ ДІАГНОСТИКИ COVID-19 НА ОСНОВІ NAIVE BAYES

Пандемія COVID-19, що триває у всьому світі, створює критичну потребу у швидких та точних методах діагностики. Раннє визначення можливого зараження сприяє своєчасному медичному втручання, зменшенню ризику масового поширення вірусу та оптимізації навантаження на медичні заклади. У цьому контексті застосування методів машинного навчання, зокрема алгоритму Naive Bayes, є ефективним підходом для створення систем класифікації стану пацієнта на основі ключових медичних ознак: температура тіла, рівень кашлю, наявність або відсутність нюху, головний біль, сатурація крові (SpO₂) та вік. Використання Naive Bayes забезпечує високу швидкість обчислень,

простоту інтеграції в прикладні програми та можливість роботи з обмеженими за обсягом наборами даних, що є критичною вимогою для медичних інформаційних систем.

Метою даної роботи було розробити програмний застосунок для діагностики COVID-19 з інтегрованим графічним інтерфейсом користувача, що забезпечує швидкий та інтуїтивний ввід даних пацієнта та миттєвий розрахунок ймовірності зараження. Для навчання моделі використовувався спеціально підготовлений датасет із числовими показниками стану пацієнтів з підтвердженим та відсутнім діагнозом COVID-19. Дані проходили багаторівневу попередню обробку: видалення пропущених значень, фільтрацію викидів, стандартизацію та нормалізацію ознак.

Система класифікації базується на Naive Bayes, що дозволяє обчислювати ймовірність належності нового пацієнта до категорії «COVID-позитивний» або «COVID-негативний». Для підвищення точності прогнозів застосовано тестування на виділеному наборі даних, оцінку показників точності та збалансованого F1-score, що дозволяє визначити ефективність моделі на практичних прикладах. Розроблений застосунок також включає алгоритми генерації текстового медичного звіту, що інформує користувача про рівень ризику, можливі ускладнення та рекомендації щодо подальших дій, таких як проведення ПЛР-тесту, самоспостереження або негайне звернення до лікаря.

Медичний аналіз COVID-19

Система Аналізу Ризику COVID-19

Температура (°C):	38.4
Рівень кашлю (0-3):	2
Втрата нюху (0-1):	1
Головний біль (0-3):	2
Вік:	11
Сатурація SpO ₂ (%):	89

Проаналізувати

=== МЕДИЧНИЙ АНАЛІЗ ПАЦІЄНТА ===

Температура: 38.4 °C
Рівень кашлю: 2/3
Втрата нюху: так
Головний біль: 2/3
Вік: 11 років
SpO₂: 89%

Ймовірність COVID: 1.000
Діагноз: COVID-ПОЗИТИВНИЙ

- ⚠ Низька сатурація! Ризик ускладнень.
- ⚠ Сильний кашель.
- ⚠ Втрата нюху – типовий симптом COVID.

📌 Рекомендація: Негайно звернутися до лікаря.

Рисунок 1 – Графічний інтерфейс застосунку для діагностики COVID-19 на основі алгоритму Naive Bayes

Використання графічного інтерфейсу на базі Tkinter забезпечує зручну взаємодію користувача з системою. Система дозволяє вводити медичні параметри пацієнта у числовому форматі та миттєво отримувати інтерпретовані результати, що робить її придатною для застосування як у медичних закладах, так і в умовах дистанційного моніторингу стану здоров'я. Тестування показало, що модель досягає високої точності на тестовому наборі даних, що підтверджує її практичну придатність для первинної діагностики COVID-19.

Таким чином, проведене дослідження демонструє ефективність інтеграції методів машинного навчання з інтерфейсними рішеннями для створення систем діагностики. Основний внесок роботи полягає у показі того, що навіть простий статистичний алгоритм, як Naive Bayes, може успішно використовуватись для побудови інформаційної системи класифікації ризику COVID-19, забезпечуючи високу швидкість обчислень, точність прогнозів та доступність для кінцевого користувача. Система має потенціал для подальшої адаптації під інші інфекційні хвороби та розширення функціоналу з урахуванням нових медичних показників.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Scikit-learn: Naive Bayes. URL: https://scikit-learn.org/stable/modules/naive_bayes.html (date of access: 29.11.2025)
2. Pandas. URL: <https://pandas.pydata.org/> (date of access: 29.11.2025).
3. Tkinter Documentation. URL: <https://docs.python.org/3/library/tkinter.html> (date of access: 29.11.2025).

УДК 621.398:004.71/.72

*Деятеріков Р.Г., Мельничук С.В., Воробець Г.І.
Чернівецький національний університет
імені Юрія Федьковича, м. Чернівці, Україна*

ОСОБЛИВОСТІ АНАЛІЗУ І СИНТЕЗУ АДАПТИВНОЇ ПОДІЄВО-ОРІЄНТОВАНОЇ SDN-МОДЕЛІ ПІДВИЩЕННЯ ПРОДУКТИВНОСТІ ТА ВІДМОВСТІЙКОСТІ ІОТ-СИСТЕМ

Системи Інтернету речей (IoT) наразі швидко розвиваються і охоплюють все новіші галузі застосування. Протокол MQTT є одним з ключових стандартів обміну повідомленнями між пристроями в IoT. Зі збільшенням кількості клієнтів і їх географічним розподіленням постає потреба в масштабуванні та підтримці управління пристроями [1, 2].

Традиційні механізми маршрутизації (DNS, Anycast, статичне балансування) не завжди здатні реагувати на зміну мережових умов або навантаження в реальному часі, що призводить до затримок у трафіках, перевантаження окремих вузлів і зниження відмовостійкості систем.

Останні наукові праці у сфері балансування навантаження в SDN-IoT системах зосереджуються переважно на мережевому рівні – оптимізації маршрутів, розподілі навантаження між контролерами або прогнозуванні трафіку. Такі рішення ефективно зменшують затримки на рівні OpenFlow і каналів зв'язку, проте ігнорують поведінку клієнтів MQTT та стан брокерів, де фактично відбувається обмін даними у більшості IoT-сценаріїв. Можливості протоколу MQTT 5.0, Server Redirection, відкривають шлях до створення адаптивної архітектури, де клієнт може автоматично перейти до оптимального брокера залежно від поточного стану мережі [3-5].

Пропонується концепція керування сеансами взаємодії між брокером та клієнтським пристроєм з використанням механізму Server Redirection. Це дозволяє на основі аналізу стану брокера та параметрів з'єднання між брокером і клієнтом реалізувати динамічне переспрямування окремого клієнта до менш завантаженого або оптимального з погляду мережевої топології брокера. Дана концепція:

- поєднує SDN-телеметрію та протокольні можливості MQTT 5 (*Server Redirection*), забезпечуючи адаптивне керування брокерами на рівні окремого клієнта;
- діє у допередсесійному (pre-session) режимі, що дозволяє приймати рішення ще до початку обміну даними — це мінімізує затримку підключення;
- реалізує розподілене балансування, де рішення приймаються децентралізовано через брокерну логіку, а не виключно контролером;
- забезпечує відмовостійкість і масштабованість, оскільки клієнти можуть динамічно переходити між брокерами без втрати сесій.

Загалом запропонована система переносить балансування з рівня мережі на рівень протоколу, роблячи його більш гнучким, швидкодійним і придатним для масштабних IoT-середовищ, де традиційні SDN-рішення вже не забезпечують необхідної адаптивності.

Модель процесу вимірювання мережових параметрів, їх аналізу та прийняття рішення щодо адаптивного налаштування мережі можна представити наступними процедурами:

1. Вимірювання затримки рівня протоколу (RTT_MQTT). Визначається як час між відправленням клієнтом повідомлення PUBLISH і отриманням підтвердження від брокера (PUBACK для QoS 1 або PUBCOMP для QoS 2):

$$RTT_MQTT = t_ack - t_publish$$

Цей показник відображає реальну затримку обробки повідомлення на рівні протоколу MQTT та включає час реакції брокера.

2. Мережева затримка (RTT_TCP). Може бути отримана з параметра $SRTT$ (Smoothed Round Trip Time) TCP-з'єднання брокера або визначена за інтервалом між повідомленнями $PINGREQ$ і $PINGRESP$:

$$RTT_TCP = t_pingresp - t_pingreq \parallel \text{inet:getstat(Socket, [rtt])}$$

Цей параметр характеризує чисту мережеву затримку без урахування часу обробки на рівні застосунку.

3. Ефективна затримка (інтегральний показник). Для прийняття рішень використовується інтегральне значення, що враховує обидві складові:

$$RTT_eff = \alpha \times RTT_MQTT + (1 - \alpha) \times RTT_TCP$$

де α — ваговий коефіцієнт (у межах від 0,7 до 0,8), який визначає пріоритет прикладного рівня в оцінці якості зв'язку.

4. Згладжування вимірювань. Отримане значення RTT_eff підлягає експоненційному згладжуванню (метод ЕМА), що дає змогу зменшити вплив випадкових коливань мережевих параметрів і підвищити стабільність оцінки.

5. Показник RTT_eff інтегрується у мультикритеріальну функцію вибору брокера як ключовий параметр оцінки якості мережевого шляху між клієнтом і брокером, впливаючи на рішення про перенаправлення клієнта.

Запропонована модель визначення затримки між клієнтом і брокером поєднує прикладний рівень MQTT та мережевий рівень TCP, що забезпечує більш точну оцінку якості з'єднання. На відміну від традиційних підходів, які враховують лише транспортну затримку, ця модель дозволяє оцінити повний шлях обробки повідомлення – від клієнта до брокера та у зворотному напрямку. Використання показників RTT_MQTT і RTT_TCP , згладжених методом ЕМА, забезпечує стабільність вимірювань навіть за умов шуму мережі. Дані про затримку можуть отримуватися безпосередньо від брокера через статистику сокета ($SRTT$) або через MQTT-повідомлення $PINGREQ/PINGRESP$, що робить метод гнучким і універсальним. Інтеграція цих параметрів у мультикритеріальну функцію вибору брокера дозволяє приймати обґрунтовані рішення про перенаправлення клієнтів.

Таким чином, запропонована концепція та модель її реалізації забезпечує реалістичне, вимірюване та протокольно узгоджене визначення затримки, необхідне для реалізації адаптивного балансування в SDN-керованих MQTT-системах.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Shvaika A. et al. A distributed architecture for MQTT messaging: the case of TBMQ. // Journal of Big Data. <https://journalofbigdata.springeropen.com/articles/10.1186/s40537-025-01271-x> <https://doi.org/10.1186/s40537-025-01271-x>
2. Банар А. Ю., Воробець Г. І. Алгоритми штучного інтелекту для оптимізації функціонування SDN: сучасні підходи та перспективи. *Зв'язок*. 2025. № 4. С. 11-18. DOI: 10.31673/2412-9070.2025.041241.
3. Jiménez, M.B., Fernández, D., Rivadeneira, J.E., Bellido, L. & Cárdenas, A.A. Survey of the Main Security Issues and Solutions for the SDN Architecture. *IEEE Access*, 2021, vol. 9, pp. 122016-122038. DOI: 10.1109/ACCESS.2021.3109564.
4. Ehsan Shahri, Paulo Pedreiras, and Luis Almeida. Response Time Analysis for RT-MQTT Protocol Grounded on SDN. In Fourth Workshop on Next Generation Real-Time Embedded Systems (NG-RES 2023). Open Access Series in Informatics (OASIs), Volume 108, pp. 5:1-5:15, Schloss Dagstuhl – Leibniz-Zentrum für Informatik (2023) <https://doi.org/10.4230/OASIs.NG-RES.2023.5>
5. Fatma Hmissi, Sofiane Ouni. SDN-DMQTT: SDN-Based Platform for Re-configurable MQTT Distributed Brokers Architecture. // Lecture Notes of the Institute for Computer Sciences, Social Informatics and Telecommunications Engineering ((LNICST, volume 594, pp. 393-411)) / International Conference on Mobile and Ubiquitous Systems: Computing, Networking, and Services. https://link.springer.com/chapter/10.1007/978-3-031-63992-0_26

МАШИННЕ НАВЧАННЯ ТА ШТУЧНИЙ ІНТЕЛЕКТ

УДК 004.85; 004.89

*Воїнова О.В., Гунченко Ю.О.
Одеський національний університет
імені І.І.Мечникова, м. Одеса, Україна*

МОДЕЛЮВАННЯ НЕЙРОМЕРЕЖЕВОЇ СИСТЕМИ РОЗПІЗНАВАННЯ ТА ПЕРЕКЛАДУ МОВИ ЖЕСТІВ

Мова жестів є повноцінним засобом комунікації для мільйонів людей з порушеннями слуху, однак її застосування в повсякденному житті часто обмежується недостатньою кількістю перекладачів та відсутністю технологій, здатних працювати у реальному часі. Це створює бар'єри у спілкуванні, доступі до освіти, отриманні медичних чи державних послуг.

Мова жестів використовує певні рухи рук, міміку та вираз обличчя для передавання алфавіту, цифр, слів та емоцій. Кожна країна має свою систему жестів з різними рухами рук, наприклад американська мова жестів (ASL) та індіанська мова жестів (ISL). Найпоширеніша мова жестів у світі — ASL.

З ростом використання комп'ютерів і смартфонів стало зручніше застосовувати їх для спілкування з глухими людьми. За даними Всесвітньої організації охорони здоров'я, глухі люди становлять близько 6% населення світу, тому необхідно впроваджувати технології для інтеграції людей з інвалідністю в спільноти. [1]. Потрібні переклади з жестової мови на мовлення та навпаки, щоб полегшити це спілкування.

З останніх досліджень досліджень перекладу жестів можна виділити наступні:

1. У роботі «Алгоритм лінійного дискримінантного аналізу (LDA)» індійську жестову мову розпізнавали в MATLAB, використовуючи LDA-класифікатор. Набір даних складався з десяти зображень для кожного з 26 жестів індійського алфавіту, а результати перетворювали у текст і голос [2].

2. «Згортоква нейронна мережа» — застосування CNN для класифікації жестів у відеопотоці (30–40 кадрів/с). Попередню обробку

виконували за допомогою OpenCV, а розпізнані жести озвучували через Google Audio. Набір даних включав жести літер, цифр і статичних слів [3].

3. У дослідженні «Глибоке навчання» (1000 відео індійських жестів) запропонували модель реального часу, де MobileNet використовували для класифікації 1000 дій [4].

4. «Мультиmodalний з урахуванням скелета» — представлена система, яка працює з трьома наборами: турецьким, китайським та американським [5].

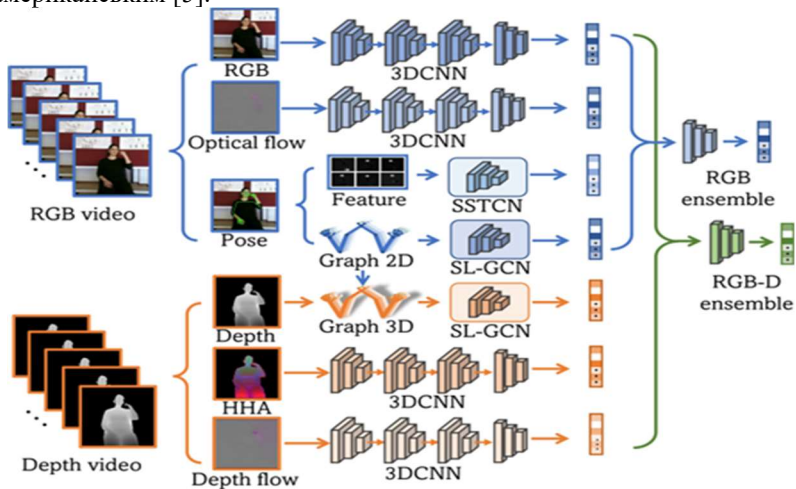


Рисунок 1 – Концепція системи розпізнавання жестової мови на основі скелету з глобальною ансамблевою моделлю (SAM-SLR-v2)

Оскільки представлені дослідження суттєво відрізняються за методами, обсягами даних та результатами, важливо порівняти їх між собою, щоб визначити сильні сторони кожного підходу та загальні тенденції розвитку систем розпізнавання жестової мови, що представлено у таблиці 1 [2].

Сучасні підходи до розпізнавання й перекладу жестової мови засновані на неймережах і забезпечують підвищену точність та швидкість. Найперспективніші CNN та мультиmodalні моделі працюють добре навіть за складних умов, але часто обмежені даними або обладнанням.

В роботі проведено аналіз існуючих систем перекладу жестової мови, показано, що подальший розвиток таких систем має бути спрямований на створення універсальних моделей, здатних працювати

з різними мовами жестів, підтримувати динамічні послідовності та забезпечувати роботу в режимі реального часу. Це дозволить значно розширити можливості комунікації людей із порушеннями слуху та сприятиме їхній інтеграції у цифрове середовище та суспільство загалом.

Таблиця 1 – Порівняння методів розпізнавання жестів

Підхід	Використаний набір даних	Мова жестів	Показники ефективності
Метод лінійного дискримінантного аналізу (LDA)	10 зображень для кожного з 26 жестів індійського алфавіту	Індійська	Заявлена висока точність
Згортова нейронна мережа	Комплекс зображень літер, цифр та статичних жестів	Американська	Точність становила 86%
Глибоке навчання з використанням MobileNet	1000 дій, по 30 відео (30 кадрів) на кожен жест	Індійська	Точність — 92,1%
Мультимодальна скелетна система	Турецький, китайський та американський набори даних	Турецька, китайська, американська	ТОП-1 точність: 99,76%, 99,95%, 84,94%

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Y. Shovkovyi . Automatic sign language translation system using neural network technologies and 3D animation. *Innovative technologies and scientific solutions for industries*. 2023. № 4(26). С. 108–121. URL: <https://doi.org/10.30837/itssi.2023.26.108> (дата звернення: 25.11.2025).

2. Najib F. M. Sign language interpretation using machine learning and artificial intelligence. *Neural computing and applications*. 2024. URL: <https://doi.org/10.1007/s00521-024-10395-9> (дата звернення: 25.11.2025).

3. Shubham Kr Mishra. Recognition of hand gestures and conversion of voice for betterment of deaf and mute people. *CrossRef listing of deleted dois*. 2019. Vol. 1. URL: https://doi.org/10.1007/978-3-8349-6966-8_5 (дата звернення: 25.11.2025).

4. Amrithaa I. S. Sandhiya A. Translating Indian sign language to text using deep learning. *International journal of innovative science and research technology*. 2022. Т. 7, № 5. С. 7154–7161.

5. Jiang S., Sun B. Sign language recognition via skeleton-aware multi-model ensemble. *Svf*. 2021. URL: <https://doi.org/10.48550/arXiv.2110.06161> (дата звернення: 25.11.2025).

УДК 004.415.2

Гожий В.О.

*Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ОСОБЛИВОСТІ ВПРОВАДЖЕННЯ ТЕХНОЛОГІЙ ШТУЧНОГО ІНТЕЛЕКТУ В МЕНЕДЖМЕНТІ ІТ ПРОЄКТІВ

Штучний інтелект суттєво змінив способи планування, виконання та реалізації ІТ проєктів [1]. Штучний інтелект (ШІ) та машинне навчання (МН) суттєво трансформували галузь управління проєктами та зробили значний внесок у підвищення ефективності, прийняття рішень та оптимізацію розподілу ресурсів. Незаперечно важлива роль управління проєктами у різних галузях, оскільки він відіграє вирішальну роль у забезпеченні структури та керівництва, необхідних для досягнення успішних результатів. В останні роки використання, застосування та інтеграція ШІ у галузі управління проєктами привернули увагу дослідників та науковців. ШІ може оптимізувати операції за рахунок автоматизації завдань, що повторюються, оптимізації розподілу ресурсів і поліпшення оцінки ризиків, що призводить до більш ефективного виконання проєктів [2].

Технології штучного інтелекту надають численні можливості в управлінні проєктами, підвищуючи ефективність, прийняття рішень та результати проєкту. Такі технології, як штучні нейронні мережі (ШНМ), нечітка логіка та генетичні алгоритми (ГА), надають сучасні інструменти для автоматизації рутинних завдань, удосконалюють прийняття складних рішень, оптимізувати планування та покращення ресурсів (ГА) [3, 4]. Тим не менш, успішне застосування ІІ в управлінні проєктами залежить від низки критичних факторів. Ключові фактори включають потребу у великих даних для навчання моделей ШІ, передові алгоритми, такі як ШНМ та моделювання Монте-Карло для оцінки ризиків та прогнозування продуктивності, а також сильну інституційну або державну підтримку для створення необхідної інфраструктури для проєктів. Наприклад, лабораторія ШІ та дорожня карта ШІ *Smart Dubai* пропонують структуру, яка сприяє інноваціям та підтримує впровадження ШІ. Більш того, інтеграція ІІ в інформаційні системи управління проєктами забезпечує безперебійний моніторинг

життєвого циклу проекту, дозволяючи проводити оновлення в режимі реального часу та покращувати процеси прийняття рішень [3, 4].

Незважаючи на переваги сучасних технологій ШІ, існує низка перешкод, що перешкоджають ефективному використанню ШІ в управлінні проектами. Неефективне впровадження ШІ може бути результатом недостатнього навчання та низької якості даних [4]. Крім того, висока вартість інструментів та інфраструктури ШІ створює фінансові проблеми, особливо для невеликих організацій. Більше того, деякі організації опираються змін та використовують технології в проектах, оскільки впровадження ІІ потребує значних коригувань у звичайному режимі роботи. Перешкоди пов'язані з етичними та нормативними проблемами, включаючи питання довіри, конфіденційності та управління, що наголошує на необхідності поступового та ретельно контролюваного процесу впровадження.

Технології штучного інтелекту здійснили революцію в різних галузях, дозволивши системам імітувати людський інтелект для вирішення проблем та прийняття рішень. Центральною концепцією в сучасному ШІ є машинне навчання (МН), технології і методи, що дозволяють системам виявляти закономірності та взаємозв'язки даних без явного програмування. МН може опрацьовувати великі набори даних і генерувати прогнози, які вважаються ключовими для успішного управління проектами, де управління складною інформацією створює значні проблеми [5, 6, 7]. Методи машинного навчання (МН) мають перспективу та можливості для забезпечення точності прогнозування, оптимізації розподілу ресурсів та просування процесів прийняття рішень. Наприклад, методи МН можуть використовувати історичні дані проекту для прогнозування термінів та вимог бюджету, тим самим покращуючи планування та дозволяючи приймати рішення на основі даних, які сприяють більш успішним результатам проекту. Такий підхід особливо корисний для вирішення таких завдань, як класифікація недоліків та оцінка ризиків, надаючи цінну інформацію про особливості та недоліки проекту. Понад те, будучи найважливішим компонентом революції великих даних, методи машинного навчання дозволяють системам аналізувати великі набори даних, виявляти закономірності і робити обґрунтовані прогнози, постійно адаптуючись і вдосконалюючись із часом.

Хоча інтерес до взаємодії методів ШІ та управління проектами зростає, аналіз літератури виявляє значну недостачу інформації. Хоча в багатьох дослідженнях вивчалися потенційні застосування ШІ в рамках управління проектами, жодне комплексне дослідження не проводило систематичного аналізу бар'єрів та факторів, пов'язаних з інтеграцією

ІІІ в управлінні проектами. Існуючі дослідження зазвичай зосереджені на окремих аспектах, таких як автоматизація, прийняття рішень та оцінка ризиків, не забезпечуючи цілісного розуміння проблем, з якими стикаються організації при впровадженні рішень на основі ІІІ. Цей недолік обмежує можливість розробки ефективних стратегій для подолання бар'єрів впровадження та використання всього потенціалу ІІІ в управлінні ІТ проектами [7, 8].

Ретельний аналіз необхідний поглиблення знань у цій галузі та підвищення ефективності впровадження ІІІ в ІТ фірмах та організаціях. Розуміння цих бар'єрів та факторів, що сприяють впровадженню ІІІ, дозволить заінтересованим сторонам розробляти цільові стратегії для зниження ризиків, використання можливостей та оптимізації методів управління ІТ проектами на основі методів штучного інтелекту.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Zhou, W.; Yan, Z.; Zhang, L. A comparative study of 11 non-linear regression models highlighting autoencoder, DBN, and SVR, enhanced by SHAP importance analysis in soybean branching prediction. *Sci. Rep.* 2024, *14*, 5905.

2. Taboada, I.; Daneshpajouh, A.; Toledo, N.; de Vass, T. Artificial Intelligence Enabled Project Management: A Systematic Literature Review. *Appl. Sci.* 2023, *13*, 5014.

3. Prifti, V. Optimizing project management using artificial intelligence. *Eur. J. Form. Sci. Eng.* 2022, *5*, 30–38.

4. Auth, G.; Jöhnk, J.; Wiecha, D.A. A Conceptual Framework for Applying Artificial Intelligence in Project Management. In Proceedings of the 2021 IEEE 23rd Conference on Business Informatics (CBI), Bolzano, Italy, 1–3 September 2021; Volume 1, pp. 161–170.

5. Polonevych, O.V.; Sribna, I.M.; Mykolaychuk, V.R.; Tkalenko, O.M.; Shkapa, V.V. Artificial Intelligence Applications for Project Management. *Connectivity* 2020, *146*, 054855.

6. Gil, J.; Torres, J.M.; González-Crespo, R. The application of artificial intelligence in project management research: A review. *Int. J. Interact. Multimed. Artif. Intell.* 2021, *6*, 54–66.

7. Jiang, F.; Jiang, Y.; Zhi, H.; Dong, Y.; Li, H.; Ma, S.; Wang, Y.; Dong, Q.; Shen, H.; Wang, Y. Artificial intelligence in healthcare: Past, present and future. *Stroke Vasc. Neurol.* 2017, *2*.

8. Ribeiro, J.; Lima, R.; Eckhardt, T.; Paiva, S. Robotic process automation and artificial intelligence in industry 4.0—a literature review. *Procedia Comput. Sci.* 2021, *181*, 51–58.

ІНТЕГРАЦІЯ НЕЙРОМЕРЕЖЕВИХ КОМПОНЕНТІВ У БАГАТОШАРОВІ АРХІТЕКТУРИ ТРАНСФОРМАЦІЇ ДАНИХ

Якість даних є критичним фактором успіху для систем бізнес-аналітики ВІ та штучного інтелекту AI/ML. Надійність аналітичних висновків безпосередньо залежить від точності, повноти та узгодженості інформації. Для управління зростаючими обсягами даних розроблені сучасні архітектурні парадигми, такі як Lakehouse. Цей підхід об'єднує гнучкість озер даних Data Lakes із надійністю та структурованістю сховищ даних Data Warehouses, усуваючи необхідність підтримки надлишкових систем та спрощуючи архітектуру безпеки.

Найпоширенішим патерном реалізації Lakehouse є «медальйонна архітектура», що передбачає три рівні обробки даних. Бронзовий шар (Bronze Layer) забезпечує зберігання «сирих» даних у вихідному форматі для гарантування можливості аудиту та відтворення історії. Наступним є срібний шар (Silver Layer), який містить очищені, дедупліковані та інтегровані дані, придатні для запитів, адже саме тут відбувається основна трансформація та валідація. Золотий шар (Gold Layer) складається з агрегованих вітрин даних, що оптимізовані безпосередньо для бізнес-аналітики.

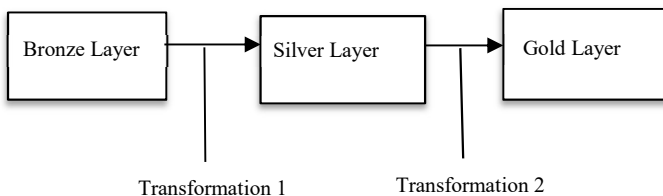


Рисунок 1 – Концептуальна схема медальйонної архітектури.

Постановка проблеми.

Традиційно переміщення даних між зазначеними шарами забезпечується процесами ETL (Extract, Transform, Load). Проте класичні ETL-скрипти мають суттєві недоліки: вони вимагають значних ручних зусиль для розробки, є «крихкими» до змін у схемі вхідних

даних та обмежені синтаксичною логікою обробки. Звичайний скрипт не здатен виявити семантичну аномалію або ефективно обробити неструктурований текст без жорстко заданих правил.

Мета та гіпотеза дослідження.

Метою роботи є дослідити можливість інтеграції великих мовних моделей (BMM) безпосередньо в ETL-конвеєри як внутрішнього компонента обробки, а не лише як кінцевого споживача даних. Гіпотеза полягає в тому, що нейронні мережі можуть замінити ручні правила на інтелектуальні агенти, здатні до семантичного розуміння контексту, що відкриває можливості для структурування PDF-документів, семантичної валідації та виявлення чутливої інформації у тексті.

Огляд існуючих рішень.

Аналіз літератури свідчить, що використання BMM в інженерії даних перебуває на експериментальній стадії, проте вже існують перспективні розробки:

- Clean Agent – агентний фреймворк, який автоматизує стандартизацію даних, поєднуючи можливості BMM з декларативними бібліотеками очищення даних;
- Cosoop – система, що використовує гібридний підхід, комбінуючи статистичне профілювання з семантичним розумінням BMM для генерації SQL-запитів очищення;
- RetClean – архітектура, яка застосовує підхід RAG (Retrieval-Augmented Generation) для відновлення пропущених значень шляхом пошуку релевантних прикладів у базах знань.

Індустрія також рухається в даному напрямку. Платформа Databricks впроваджує AI Functions, наприклад, `ai_query`, `ai_classify`, які дозволяють викликати генеративні моделі безпосередньо з SQL-запитів під час завантаження даних у срібний шар.

Виклики.

Попри перспективи, впровадження BMM в ETL стикається з трьома фундаментальними проблемами:

- Недетермінізм. ETL-процеси вимагають відтворюваності, тоді як BMM можуть генерувати різні відповіді на однакові вхідні дані;
- галюцинації, адже ризик спотворення даних є критичним для аналітичних систем;
- вартість та швидкість. Використання BMM для обробки кожного рядка великих масивів даних є економічно неефективним та повільним порівняно з традиційним SQL.

Напрямки подальших досліджень.

Для подолання зазначених викликів пропонується зосередити увагу на розробці гібридних архітектур. Ключовим завданням є створення

алгоритмів, які класифікують дані за складністю: прості структуровані записи обробляються ефективним SQL, тоді як складні неструктуровані випадки, наприклад, логи помилок у JSON, передаються на обробку ВММ. Перспективним є підхід «генерації коду», де ВММ не обробляє дані безпосередньо, а генерує оптимізований скрипт очищення для конкретного набору даних, що вирішує проблему швидкості та вартості. Важливою також є розробка шарів валідації для нівелювання недетермінізму та адаптація моделей до специфічного доменного контексту організації.

Висновки.

Використання ВММ в ETL є новим трендом, який доповнює традиційні методи. Майбутнє інженерії даних полягає у гібридних системах, здатних поєднувати надійність та швидкість детермінованих алгоритмів із семантичним інтелектом нейронних мереж.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Armbrust M., Ghodsi A., Xin R. et al. Lakehouse: A New Generation of Open Platforms for Data Lakehouses. CIDR 2021. 2021. URL: https://www.cidrdb.org/cidr2021/papers/cidr2021_paper17.pdf (дата звернення: 29.11.2025).
2. Qi D., Wang J. CleanAgent: Automating Data Standardization with LLM-based Agents. arXiv preprint arXiv:2403.08291. 2024. DOI: <https://doi.org/10.48550/arXiv.2403.08291> (дата звернення: 29.11.2025).
3. Zhang S., Huang Z., Wu E. Cocoon: Data Cleaning Using Large Language Models. arXiv preprint arXiv:2410.15547. 2024. DOI: <https://doi.org/10.48550/arXiv.2410.15547> (дата звернення: 29.11.2025).
4. Naeem Z. A., Ahmad M. S., Eltabakh M. et al. RetClean: Retrieval-Based Data Cleaning Using LLMs and Data Lakes. Proceedings of the VLDB Endowment. 2024. Vol. 17, No. 12. P. 4421–4424.
5. AI Functions [Electronic resource] // Databricks Documentation. URL: <https://docs.databricks.com/en/large-language-models/ai-functions.html> (дата звернення: 29.11.2025).

ДОСЛІДЖЕННЯ ЗГОРТКОВИХ НЕЙРОННИХ МЕРЕЖ ЗА ДОПОМОГОЮ МЕТОДІВ ПОЯСНЮВАНОГО ШТУЧНОГО ІНТЕЛЕКТУ

Розвиток технологій машинного навчання, зокрема згорткових нейронних мереж (CNN), у поєднанні зі зростанням обчислювальних потужностей та сучасними підходами до обробки великих масивів даних суттєво підвищили точність і ефективність систем розпізнавання. Загалом CNN набули суттєвого значення для аналізу зображень і відео, забезпечуючи високу здатність виділяти складні ознаки і працювати з різними структурами даних. Проте із зростанням складності або інтеграції таких моделей їх інтерпретованість суттєво зменшується — процес прийняття рішень може здаватися непрозорим навіть для розробників, а система перетворюється на «чорний ящик». Це ускладнює оцінку роботи моделі, її узагальнювальні властивості та безпосередньо впливає на подальший розвиток.

У зв'язку з цим виникає потреба у методах, здатних пояснити внутрішню логіку роботи подібних складних моделей. Методи пояснюваного (або інтерпретованого) штучного інтелекту (Explainable Artificial Intelligence, XAI) дозволяють зрозуміти, які ознаки CNN вважає найважливішими під час прийняття рішення, як змінюється поведінка мережі при модифікації вхідних даних і чи не присутні у моделі небажані упередження [1].

Інтегровані градієнти (англ. Integrated Gradients, IG) — це метод пояснюваного штучного інтелекту, який визначає внесок кожної вхідної ознаки у фінальний прогноз моделі [2]. Він порівнює градієнти виходу моделі відносно вхідних даних уздовж шляху між “базовим” вхідним значенням (наприклад, чорним зображенням) і реальним прикладом. Інтегруючи ці градієнти, IG надає стабільні й теоретично обґрунтовані оцінки важливості ознак.

Інтеграл інтегрованих градієнтів можна ефективно апроксимувати за допомогою підсумовування у точках, що знаходяться на досить малих інтервалах вздовж прямолінійного шляху від базової лінії x' до вхідного значення x , кінцевий вигляд для розрахунку наведено нижче (див. формулу 1).

$$\text{IntegratedGrads}_i^{\text{approx}}(x) ::= (x_i - x_i') \times \sum_{k=1}^m \frac{\partial F\left(x' + \frac{k}{m} \times (x - x')\right)}{\partial x_i} \times \frac{1}{m} \quad (1)$$

де x – вхідні дані;

x' – базова маска;

$\frac{\partial F(x)}{\partial x_i}$ – градієнт $F(x)$ вздовж i -го виміру;

m – число кроків у наближенні інтеграла Рімана.

XRAI — інший метод, який розширює ідею Integrated Gradients. Він використовує IG для оцінки важливості пікселів, а потім об'єднує їх у змістовні регіони, аналізуючи, які області зображення найбільше сприяють прогнозу. У результаті XRAI дає більш інтерпретовані та структуровані карти важливості, ніж методи, що працюють на рівні окремих пікселів. В свою чергу Grad-CAM створює теплову карту важливих регіонів зображення, використовуючи градієнти, що проходять через останні згорткові шари нейронної мережі. Метод показує, на які просторові області модель найбільше “спиралась” під час класифікації. Grad-CAM ідеально підходить для візуального пояснення CNN у задачах комп'ютерного зору, як і XRAI.

Для дослідження було побудовано малу модель класифікації, за основу взято попередньо навчену модель MobileNetV3. Тренування відбувалося на власному наборі даних, що містив зображення будівель, кожне з яких мало відповідну мітку: віціліла будівля, пошкоджена або зруйнована. Точність моделі на тестовому наборі склала 80%, що є достатнім, оскільки першочерговою ціллю дослідження є не отримання найкращої моделі, а інтерпретація її результатів.

У проєкті була застосована бібліотека, створена розробниками Google у межах ініціативи PAIR, що реалізовує як авторський алгоритм XRAI, так й інші [4]. На рис. 1 наведено приклад із застосуванням декількох алгоритмів для тестового зображення.



Рисунок 1 – Результати застосування різних методів XAI, наведені у послідовності: оригінальне зображення, теплова карта XRAI зі значущістю 30%, Vanilla IG і Blur IG

Як видно з рисунку, різні методи ХАІ формують різні типи сегментованих зображень, що відображають області найбільшої значущості для нейронної мережі під час класифікації. Зокрема, Vanilla IG створює більш фрагментовані теплові карти, які частково акцентують увагу на зелених зонах замість основних об'єктів — будівель. Натомість більш сучасні й модернізовані методи, такі як Blur IG і XRAI, точніше виділяють регіони, де розташовані зруйновані будівлі (відповідно до класу зображення), що свідчить про їх краще вміння до інтерпретації результатів моделі. Також варто підкреслити можливість регулювання порогу відсоткової значущості в XRAI, що дозволяє гнучко підбирати параметри візуалізації.

Отримані результати демонструють, що методи ХАІ надають цінні інструменти для інтерпретації моделей розпізнавання різної складності. Вони дозволяють простежити логіку прийняття рішень моделі відповідно до вхідних даних, глибше зрозуміти внутрішні процеси мережі та визначити шляхи її покращення або подальшого застосування.

Перспективним напрямом розвитку є використання ХАІ для моделей і систем, які працюють у реальному часі, а також у складніших задачах розпізнавання як семантична сегментація. Розвиток методів пояснюваного інтелекту відкриває шлях до більш надійних й інтерпретованих систем для практичного застосування у багатьох сферах як безпечні системи керування транспортними засобами, медичній діагностиці, військовій справі й інших критично важливих сферах.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. A comprehensive review of explainable artificial intelligence (XAI) in computer vision / Z. Cheng et al. *Sensors*. 2025. Vol. 25, no. 13. P. 4166. DOI: 10.3390/s25134166.
2. Sundararajan M., Taly A., Yan Q. Axiomatic attribution for deep networks: *Proceedings of the International Conference on Machine Learning (ICML 2017)*. PMLR, 2017. С. 3319–3328.
3. XRAI: better attributions through regions / A. Kapishnikov et al. *2019 IEEE/CVF international conference on computer vision (ICCV)*, Seoul, Korea (South), 27 October – 2 November, 2019. DOI: 10.1109/iccv.2019.00505.
4. Saliency Library. Github : вебсайт. URL: <https://github.com/PAIR-code/saliency> (дата звернення: 20.11.2025).

ІНТЕЛЕКТУАЛЬНА СИСТЕМА ПРОГНОЗУВАННЯ ВАРТОСТІ КРИПТОВАЛЮТ

Сучасний ринок криптовалют є одним із найдинамічніших сегментів світової фінансової системи, що характеризується високою волатильністю, відсутністю централізованого регулювання та залежністю від численних зовнішніх факторів – економічних, технологічних і соціальних. Його динамічний розвиток і висока волатильність зумовлюють потребу у застосуванні ефективних методів аналізу та прогнозування. Значні коливання вартості цифрових активів ускладнюють процес прийняття інвестиційних рішень, тому розробка інтелектуальних систем прогнозування набуває особливої актуальності.

Розроблена інтелектуальна система прогнозування вартості криптовалют має на меті підвищення точності передбачення ринкових змін на основі аналізу історичних даних та індикаторів технічного аналізу. Система поєднує методи машинного навчання, зокрема нейронні мережі типів LSTM, BiLSTM, GRU та гібридні архітектури CNN+LSTM, із класичними статистичними підходами, що забезпечує формування коротко- та середньострокових прогнозів зміни вартості криптовалют.

У ході дослідження здійснено детальний аналіз сучасних програмних засобів і методологій прогнозування цінних змін на криптовалютному ринку. Визначено основні групи інструментів, зокрема статистичні пакети Python з бібліотеками Pandas, NumPy, Matplotlib, PyTorch, а також спеціалізовані сервіси CoinGecko та CoinMarketCap, які забезпечують доступ до ринкових даних у режимі реального часу.

Попри широкий аналітичний потенціал цих інструментів, їхнє використання вимагає високого рівня технічної підготовки, налаштування середовищ і моделей, що ускладнює їх застосування звичайними користувачами. Тому розроблено інтелектуальну систему з інтуїтивно зрозумілим інтерфейсом, яка автоматизує всі етапи аналізу та прогнозування тенденцій ринку криптовалют (див. рис. 1).

Архітектура системи включає такі основні модулі:

– модуль збору даних – отримання котирувань криптовалют через API сервісів CoinMarketCap, CoinGecko та Yahoo Finance;

- модуль попередньої обробки – очищення часових рядів, нормалізація даних і усунення пропусків;
- аналітичний модуль – розрахунок ключових індикаторів технічного аналізу (MA, EMA, RSI, MACD, Bollinger Bands), а також криптовалютних індексів (BDI, AMCI, TCMC, CFGI);
- модуль прогнозування – побудова моделей глибинного навчання (LSTM, BiLSTM, GRU, CNN+LSTM) для передбачення майбутніх цін;
- візуалізаційний інтерфейс – створення динамічних графіків і таблиць, які відображають історичні та прогнозні значення цін у зручному форматі.

Для реалізації проєкту використано мову програмування Python та фреймворк Streamlit, що забезпечують гнучкість, масштабованість та інтерактивність системи. Для візуалізації застосовано Plotly Express, а для комунікації з API – бібліотеку Requests.

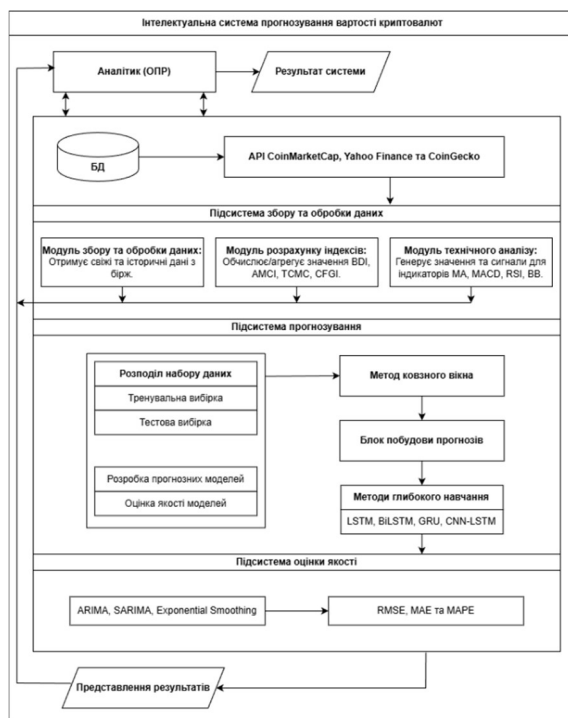


Рисунок 1 – Схема архітектури інтелектуальної системи прогнозування вартості криптовалют

У процесі тестування система була перевірена на прикладі кількох популярних криптовалют, для яких здійснювалося прогнозування вартості на основі історичних даних за певний період часу. Отримані результати підтвердили здатність моделі відображати основні ринкові тенденції та зміни у динаміці цін (див. табл. 1).

Таблиця 1 – Порівняння метрик оцінки якості прогнозування за певний проміжок часу

Криптовалюта Біткоїн (BTC)	Метрики якості прогнозів		
Hidden size: 50; Layers: 2; Learning rate: 0.001; Epochs: 200;			
30 днів / 60 точок / 4 точки кожні 6 годин	MAE	RMSE	MAPE
LSTM	1,343.24	1,027.95	1.01
GRU	858.33	621.46	0.61
BiLSTM	1,128.79	845.70	0.83
CNN+LSTM	1,281.24	933.89	0.92
ARIMA	4,178.13	3,275.29	3.24
SARIMA	4,125.21	3,212.26	3.18
Exp. Smoothing	3,551.41	2,703.59	2.68

Розроблена система реалізує повний цикл аналітики – від збору ринкових даних до формування прогнозів і аналітичних висновків у реальному часі. Її інтерфейс забезпечує можливість перегляду поточних цін, динаміки індексів та прогнозованих показників, що робить систему ефективним інструментом для аналітиків, трейдерів і дослідників ринку криптовалют. Завдяки інтерактивності, візуалізації результатів та автоматизації обчислень користувач може швидко оцінювати тенденції, виявляти потенційні ризики та приймати обґрунтовані інвестиційні рішення.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Hochreiter, S., Schmidhuber, J. (1997). Long Short-Term Memory. Neural Computation, URL: <https://doi.org/10.1162/neco.1997.9.8.1735>.
2. Chollet, F. (2021). Deep Learning with Python. 2nd ed. Manning Publications.
3. Bakar, N., Rosbi, R. (2023). Forecasting Cryptocurrency Market Using Deep Learning Models. Journal of Computational Finance, 18(4), 112–128.
4. CoinMarketCap API Documentation. [Online]. URL: <https://coinmarketcap.com/api/>.

*Дирда І. Ю., Сіденко Є. В.
Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ІНТЕЛЕКТУАЛЬНА СИСТЕМА КЛАСИФІКАЦІЇ РАКОВИХ КЛІТИН ЗА ГІСТОЛОГІЧНИМИ ЗОБРАЖЕННЯМИ НА ОСНОВІ МЕТОДІВ МАШИННОГО НАВЧАННЯ

Автоматизація діагностики онкологічних захворювань залишається однією з найважливіших задач сучасної медичної інформатики, оскільки раннє виявлення злоякісних новоутворень безпосередньо впливає на ефективність лікування пацієнтів. Зростання обсягів медичних даних вимагає впровадження інтелектуальних систем, здатних автоматично аналізувати гістологічні зображення з високою точністю.

Розробка інтелектуальної системи класифікації ракових клітин здійснюється з метою підвищення точності та швидкості діагностики, зниження навантаження на лікарів-патологів та мінімізації впливу людського фактора. Для досягнення поставленої мети використано такі методи машинного навчання, як згорткові нейронні мережі (Convolutional Neural Networks, CNN) - ResNet50, DenseNet121 та EfficientNet-B0. Застосування цих алгоритмів у поєднанні з трансферним навчанням дозволило провести багатфакторний аналіз морфологічних ознак клітин, що підвищило точність класифікації.

Аналіз аналогічних досліджень показав ефективність різних підходів до автоматизації медичної діагностики. Наприклад, F. A. Spanhol et al. [1] запропонували набір даних BreaKHis та продемонстрували застосування CNN для класифікації гістологічних зображень з точністю понад 90%. Автори K. He et al. [2] розробили архітектуру ResNet, яка вирішила проблему зникаючих градієнтів у глибоких мережах. При цьому, A. Esteva et al. [4] досягли рівня діагностики дерматологів у класифікації раку шкіри з використанням глибоких нейронних мереж. Це підтверджує ефективність комбінування сучасних архітектур CNN та методів трансферного навчання для підвищення надійності систем медичної діагностики.

Для навчання та оцінювання ефективності розробленої інтелектуальної системи класифікації ракових клітин було використано відкритий набір даних BreaKHis [6], який охоплює гістологічні зображення тканини молочної залози, зняті при різних рівнях

збільшення (40×, 100×, 200×, 400×). Набір містить 7 909 зображень, збалансованих за класами (доброякісні та злоякісні). Додатково використано методи аугментації даних для збільшення різноманітності навчальної вибірки. Сукупне використання цих джерел забезпечило репрезентативність, збалансованість і відтворюваність результатів, а також підвищило точність і узагальнювальну здатність моделей класифікації у системі.

Архітектура розробленої інтелектуальної системи класифікації ракових клітин побудована на основі модульного підходу, що забезпечує її гнучкість, масштабованість та зручність у подальшій інтеграції з медичними інформаційними системами. Система включає інтерфейс користувача, реалізований засобами Streamlit (Python); серверну частину на PyTorch, яка обробляє запити користувача та виконує класифікацію зображень; модулі попередньої обробки зображень (масштабування, нормалізація кольору, аугментація); блоки машинного навчання, що використовують алгоритми CNN (ResNet50, DenseNet121, EfficientNet-B0); модуль оцінки якості класифікації, а також сховище даних для моделей і результатів. Така структура забезпечує повний цикл обробки даних - від введення зображення користувачем до формування діагностичного висновку і збереження результатів аналізу (див. Рис. 1).

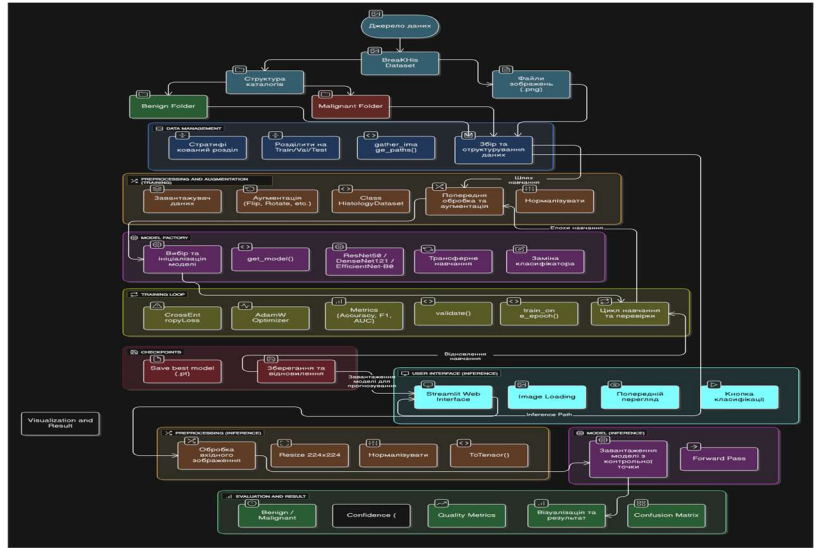


Рисунок 1 – Схема архітектури інтелектуальної системи класифікації ракових клітин

У процесі навчання моделей машинного навчання для класифікації гістологічних зображень було оцінено ефективність різних алгоритмів. Моделі для обробки зображень (ResNet50, DenseNet121, EfficientNet-B0) показали точність понад 93%, причому EfficientNet-B0 продемонструвала найкращу узагальнювальну здатність з точністю 95.8%, а DenseNet121 - найвищий рівень класифікації для складних морфологічних структур. Моделі для аналізу різних рівнів збільшення досягли точності до 96%, підтверджуючи ефективність глибокого навчання при роботі з гістологічними зображеннями. Додаткове навчання з аугментацією даних підвищило узагальнювальну здатність моделей і адаптивність до різних типів гістологічних препаратів. Система класифікації, побудована на наборі даних BreaKHis, досягла 95.8% загальної точності, успішно розпізнаючи як доброякісні, так і злоякісні категорії клітин.

Таким чином, реалізована інтелектуальна система класифікації ракових клітин показала високу ефективність і надійність у розпізнаванні патологічних змін, а поєднання алгоритмів глибокого навчання (ResNet50, DenseNet121, EfficientNet-B0) із попередньою обробкою даних, модулем оцінки якості та веб-інтерфейсом забезпечило точну класифікацію, адаптивність до різних форматів гістологічних зображень та зручну взаємодію з користувачем.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Spanhol, F. A., Oliveira, L. S., Petitjean, C., Heutte, L. A. (2016). A Dataset for Breast Cancer Histopathological Image Classification. *IEEE Transactions on Biomedical Engineering, 63(7), 1455-1462*. URL: <https://doi.org/10.1109/TBME.2015.2496264>.
2. He, K., Zhang, X., Ren, S., Sun, J. (2016). Deep Residual Learning for Image Recognition. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 770-778*. URL: <https://doi.org/10.1109/CVPR.2016.90>.
3. Tan, M., Le, Q. V. (2019). EfficientNet: Rethinking Model Scaling for Convolutional Neural Networks. *International Conference on Machine Learning (ICML), 6105-6114*. URL: <https://doi.org/10.48550/arXiv.1905.11946>.
4. Esteva, A., Kuprel, B., Novoa, R. A., Ko, J., Swetter, S. M., Blau, H. M., Thrun, S. (2017). Dermatologist-level classification of skin cancer with deep neural networks. *Nature, 542(7639), 115-118*. URL: <https://doi.org/10.1038/nature21056>.
5. Litjens, G., Kooi, T., Bejnordi, B. E., Setio, A. A. A., Ciompi, F., Ghafoorian, M., van der Laak, J. A., van Ginneken, B., Sánchez, C. I. (2017).

A survey on deep learning in medical image analysis. *Medical Image Analysis, 42, 60-88*. URL: <https://doi.org/10.1016/j.media.2017.07.005>.

6. BreaKHis: Breast Cancer Histopathological Database. Laboratory of Computer Vision (LVC), Federal University of Paraná. URL: <https://web.inf.ufpr.br/vri/databases/breast-cancer-histopathological-database-breakhis/>.

УДК 004.93

***Кінер Світлана Юрійвна**
Одеський національний університет
імені І. І. Мечникова, м. Одеса, Україна*

МЕТОДИ ВИЯВЛЕННЯ ТА РОЗПІЗНАВАННЯ ОБЛИЧ: СУЧАСНИЙ СТАН, ВИКЛИКИ ТА ПЕРСПЕКТИВИ

Упродовж останніх років технології роботи з обличчям залишаються дуже динамічними напрямками комп'ютерного зору. Вони дедалі частіше інтегруються в побутові пристрої, інтерфейси автентифікації, освітні платформи, системи відеоспостереження та інтелектуальні сервіси. Незважаючи на популярність, побудова надійної системи розпізнавання облич залишається складною задачею, що потребує глибокого аналізу методів, їх переваг та недоліків, а також розуміння реальних обмежень під час роботи з даними.

Розвиток методів виявлення облич історично починався з простих алгоритмів, у яких використовувалися ручні ознаки. Алгоритм Віюлі–Джонса зі своїми каскадами Хаара став першим практичним рішенням, яке могло працювати в реальному часі. Його перевагою була швидкість, однак конструкція слабких класифікаторів робить метод нестійким до зміни освітлення та поворотів голови. Підхід HOG+SVM забезпечив більшу точність за рахунок використання градієнтних ознак, однак ці моделі також потребують якісного зображення й демонструють обмеження в складних сценах. Поява глибинних нейронних мереж змінила сам підхід до задачі. Моделі MTCNN [3], RetinaFace [2], BlazeFace [3] та подібні їм здатні не лише знаходити обличчя, а й одночасно визначати ключові точки та якість зображення. Такі архітектури орієнтовані на реальні умови, включно з різним освітленням, ракурсами та перекриттями. Вони є поточним стандартом для промислових систем, однак їхня потужність залежить від доступності GPU, що не завжди можливо у навчальних або локальних проєктах.

У задачі розпізнавання облич спостерігається подібна еволюція – від класичних статистичних методів до глибинних моделей. PCA (Eigenfaces) та LDA (Fisherfaces) базуються на лінійних перетвореннях та зберігають швидкість обробки навіть на слабких пристроях, однак їхня ефективність падає у випадках зміни освітлення чи міміки. LBPН залишається одним із найстабільніших класичних алгоритмів, оскільки використовує локальні текстурні ознаки обличчя, що робить його стійким до частини варіацій, але не до поворотів голови чи складних сцен. Глибинні моделі, такі як FaceNet, VGGFace та ArcFace, сфокусувалися на іншій ідеї – не класифікувати обличчя безпосередньо, а навчати модель будувати ембеддинги – вектори ознак, між якими можна вимірювати відстань. Найпоширенішим сучасним рішенням є ArcFace [1], яке формує високостійкі простори ознак завдяки спеціальній функції втрат.

Попри розвиток технологій, галузь стикається з низкою обмежень, що актуальні як для промислових, так і для навчальних рішень. Зовнішні фактори – нерівномірне освітлення, тіні, низька роздільність, сильні повороти голови – залишаються складними для моделей будь-якого рівня. Іншою важливою проблемою є недостатня кількість якісних даних. Багато моделей демонструють хороші результати на великих, добре збалансованих наборах, але суттєво втрачають точність на малих персональних вибірках, що характерно для навчальних систем або закритих корпоративних застосувань. Проблема упередженості (bias) у моделях також набула значної уваги. Нерівномірне представлення різних демографічних груп у навчальних наборах може спричиняти некоректні прогнози або зниження точності для окремих категорій користувачів. Крім того, сучасні системи залишаються вразливими до атак, зокрема до використання друкованих фотографій, маскування та підроблених зображень, створених за допомогою генеративних моделей.

Сучасні дослідження демонструють кілька напрямів прогресу. Lightweight-моделі (MobileFaceNet [4], ShuffleFaceNet) дозволяють переносити розпізнавання на мобільні пристрої без значної втрати точності. Self-supervised learning відкриває можливість навчати моделі на немаркованих даних, що особливо важливо у випадках, коли розмітка є дорогою або недоступною. Few-shot learning та metric learning дозволяють досягати високих результатів при мінімальній кількості зразків [5] – це робить методи більш придатними для невеликих баз користувачів. Окремим напрямом розвитку є multimodal biometric systems – поєднання даних обличчя, голосу, поведінкових ознак або термографії, що підвищує надійність системи та зменшує кількість

помилки. Важливим трендом є *privacy-preserving ML: federated learning* та *differential privacy*, які дозволяють будувати системи без передачі необроблених зображень користувачів на сторонні сервери.

Проведений огляд показує, що жоден із наявних підходів не є універсальним: класичні методи залишаються доступними та швидкими, але обмежені в умовах змінного освітлення та ракурсів, тоді як глибинні моделі демонструють високу точність, однак потребують значних ресурсів і великих масивів даних. У практичних задачах, де обсяг зразків є невеликим, а обчислювальні можливості – обмеженими, постає потреба у підході, який поєднував би сильні сторони обох напрямів. Саме на основі цього аналізу формується концепція власного методу розпізнавання обличчя, що має бути достатньо легким для навчання на малих даних і водночас достатньо стійким до типових варіацій зображень. Оптимальним рішенням у таких умовах є гібридний підхід, який використовує класичні дескриптори для опису локальних особливостей обличчя, але організований за принципом ознакового простору, характерного для сучасних глибинних моделей. Запропонована структура методу може включати такі етапи:

1. Вирівнювання обличчя за ключовими точками. Це зменшує вплив поворотів та масштабу й створює єдину базу для подальшого аналізу.

2. Виділення локальних дескрипторів (LBP, HOG або їх комбінація). Вони забезпечують стабільність до освітлення та дозволяють отримати інформативні локальні ознаки без використання глибинних мереж.

3. Формування ознакового вектора. Гістограми та локальні характеристики об'єднуються у компактний вектор – спрощене представлення *embedding*-простору.

4. Нормалізація та, за потреби, зменшення розмірності. Це підвищує узагальнювальну здатність моделі та зменшує вплив шумів.

5. Класифікація у просторі ознак. Для невеликих наборів даних доцільно застосовувати SVM, kNN або інші методи метричного навчання.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Deng J., Guo J., Xue N., Zafeiriou S. ArcFace: Additive Angular Margin Loss for Deep Face Recognition. *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2019, pp. 4690–4699. DOI: 10.48550/arXiv.1801.07698.

2. Wang X., Zhang Z., Yu Y., et al. RetinaFace: Single-Shot Multi-Level Face Localisation in the Wild. CVPR, 2020, pp. 5203–5212. DOI: 10.1109/CVPR42600.2020.00525.

3. Zhang K., Zhang Z., Li Z., Qiao Y. Joint Face Detection and Alignment Using Multitask Cascaded Convolutional Networks. IEEE Signal Processing Letters, 2019, vol. 26, no. 10, pp. 1571–1575. DOI: DOI: 10.48550/arXiv.1604.02878.

4. He Z., Zhang Z., Sun Z. MobileFaceNet: A Compact Deep Neural Network for Face Verification. IEEE International Conference on Biometrics (ICB), 2019. DOI: 10.1109/ICB45273.2019.8987377.

5. Ji Z., Chen H., Cao Z. Few-Shot Face Recognition with Multi-center Metric Learning. Pattern Recognition, 2022, vol. 125. DOI: 10.1016/j.patcog.2021.108551.

УДК 004.82

Ковалів О. П., Сіденко Є. В, Кондратенко Г. В.

*Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ОПТИМІЗАЦІЯ ПРОГНОЗУВАННЯ ІНФЕКЦІЙНИХ ЗАХВОРЮВАНЬ ЗА ДОПОМОГОЮ НЕЙРОННИХ МЕРЕЖ З ВИКОРИСТАННЯМ ЕКЗОГЕННИХ ЗМІННИХ

Оптимізація прогнозування має важливу роль у сфері нейронних мереж, оскільки вона безпосередньо впливає на точність, стабільність та узагальнюваність моделей для прогнозування. Завдяки вдосконаленню параметрів моделі, архітектур та стратегій навчання, оптимізація прогнозування зменшує поширення помилок та покращує здатність моделі фіксувати складні, нелінійні закономірності в даних. Крім того, оптимізація прогнозування дозволяє нейронним мережам ефективніше адаптуватися до динамічних середовищ, покращуючи їхню застосовність у чутливих до часу областях, таких як фінанси, кліматичне моделювання чи розповсюдження інфекційних захворювань.

Однією з передових моделей у сфері прогнозуванні часових рядів є мережа N-BEATS [1]. Ця модель характеризується своєю здатністю самостійно вивчати часові структури в даних, знаходячи тренди та сезонність в часових рядах.

В ході дослідження базовою моделлю в прогнозуванні часових рядів було використано модель NaïveSeasonal. Цю модель стала основою для

порівняння та оцінки продуктивності основної моделі - N-BEATS. Для цього аналізу було використано дані зі щомісячних звітів з національних баз даних охорони здоров'я, а саме Центру громадського здоров'я МОЗ України. В звітах наведено абсолютні цифри та інтенсивні показники на 100 000 населення.

Набір даних було розділено на навчальний та тестовий набори для прогнозування з горизонтом прогнозування 6 місяців. Візуалізацію частини набору даних зображено нижче (рисунок 1). Крім того, значення в наборі даних були нормалізовані для врахування коливання кількості населення України.

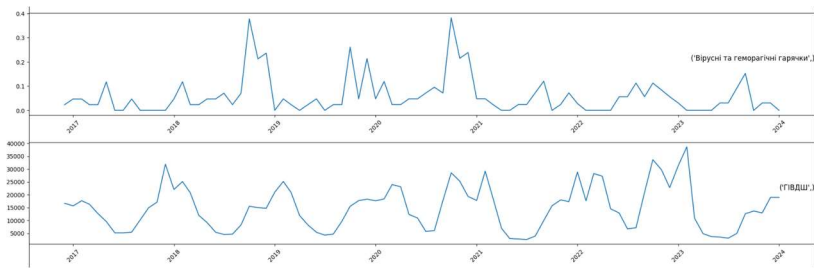


Рисунок 1 – Візуалізація частини набору даних

Модель NaiveSeasonal використовує сезонний наївний підхід, враховуючи тривалість сезону 24 години, помножену на 7 днів (один тиждень). Так як більшість часових рядів не мають очевидної сезонності, модель показала незадовільні результати. Візуалізацію порівняння між прогнозованими та істинними значеннями зображено нижче (рисунок 2). Середня абсолютна помилка (MAE) становить приблизно 108.24%.

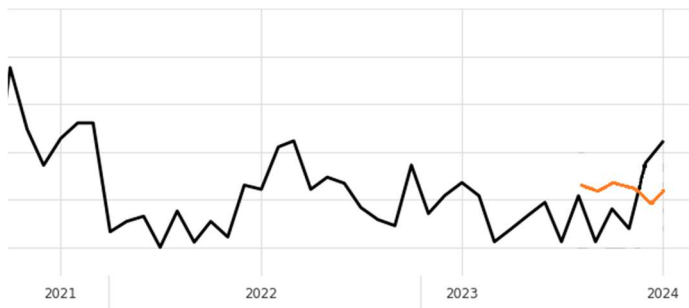


Рисунок 2 – Візуалізація прогнозування моделі NaiveSeasonal

Модель N-BEATS показала кращі результати, завдяки здатності моделі виявляти сезонність та тренди. Візуалізацію порівняння між прогнозованими та істинними значеннями зображено нижче (Рисунок 3). Середня абсолютна помилка (MAE) становить приблизно 63.57%.

Модель N-BEATS пропонує гнучкість для включення екзогенних змінних. Екзогенні змінні в часових рядах – це зовнішні фактори, які впливають на цільову змінну, але не залежать від неї [2]. Ці змінні мають бути відомі або передбачені заздалегідь і можуть значно покращити точність прогнозу часового ряду, враховуючи зовнішні події. Прикладами є економічні показники, погода, свята або рекламні заходи.

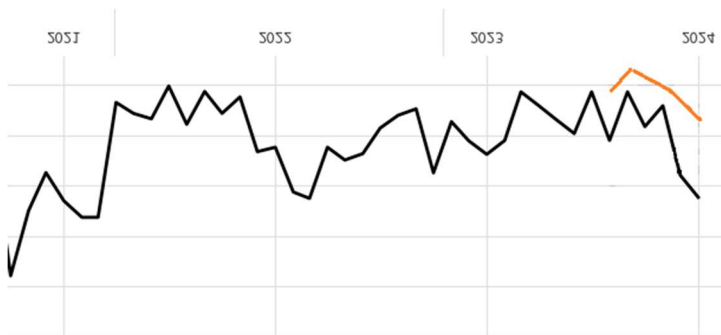


Рисунок 3 – Візуалізація прогнозування моделі N-BEATS.

Для аналізу було передано такі змінні, як місяць року та коефіцієнт надзвичайної ситуації. Значення коефіцієнту надзвичайної ситуації були підібрані з урахуванням екологічної та політичної ситуації в Україні. Так, в період з 2017 по 2020 роки коефіцієнт надзвичайної ситуації коливався в значеннях від 0.15 до 0.3, в період з 2020 по 2022 роки - в значеннях від 0.5 до 0.65 та в період з 2022 по 2025 роки - в значеннях від 0.8 до 1. Ці екзогенні змінні мали на меті покращити її прогностичні можливості, фіксуючи додаткові закономірності та кореляції, присутні в даних. Візуалізацію порівняння між прогнозованими та істинними значеннями зображено нижче (Рисунок 4). Середня абсолютна помилка (MAE) становить приблизно 37.94%.

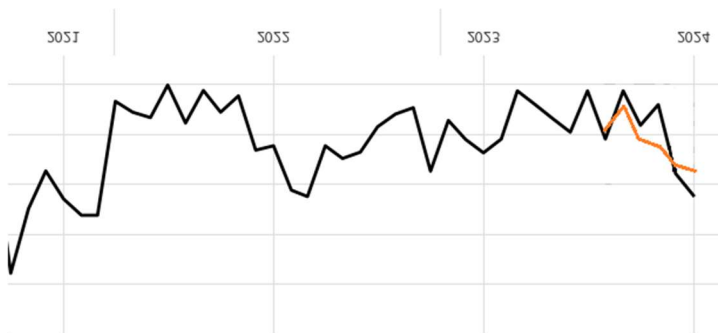


Рисунок 4 – Візуалізація прогнозування моделі N-BEATS з екзогенними змінними.

В результаті даної роботи було проведено аналіз впливу екзогенних змінних на результат прогнозування часових рядів за допомогою їх інтеграції у модель N-BEATS. Включивши такі змінні, як місяць року та коефіцієнт надзвичайної ситуації, було зафіксовано додаткові часові закономірності та розширено прогностичні можливості моделі. Модель N-BEATS показала багатообіцяючі результати з MAE 37.94%, що вказує на потенційні переваги використання екзогенної інформації для прогнозування поширення інфекційних захворювань, позиціонуючи її як найкращу модель у цьому дослідженні.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Boris N. Oreshkin, Dmitri Carpow, Nicolas Chapados, Yoshua Bengio N-BEATS: Neural basis expansion analysis for interpretable time series forecasting. URL: <https://arxiv.org/abs/1905.10437>
2. Yuxuan Wang, Haixu Wu, Jiaxiang Dong, Yong Liu, Yunzhong Qiu, Haoran Zhang, Jianmin Wang, Mingsheng Long TimeXer: Empowering Transformers for Time Series Forecasting with Exogenous Variables. URL: <https://arxiv.org/html/2402.19072v1>

ВИКОРИСТАННЯ НАВЧАННЯ З ПІДКРІПЛЕННЯМ ДЛЯ НАВІГАЦІЇ ДРОНІВ

Автономна навігація безпілотних літальних апаратів (БПЛА) у складних, динамічних та частково спостережуваних середовищах є однією з ключових проблем сучасної робототехніки та штучного інтелекту. Актуальність дослідження обумовлена стрімким зростанням потреб у безпілотних системах, здатних виконувати завдання без постійного зв'язку з оператором, особливо в умовах дії засобів радіоелектронної боротьби та протиповітряної оборони. Метою даної роботи є розробка програмної моделі середовища функціонування дронів, яку можна використовувати в межах стандартизованих алгоритмів навчання з підкріпленням для вирішення проблем навігації та виконання завдань дронами з використанням навчання з підкріпленням.

В якості основної платформи для моделювання та навчання агентів використано ігровий рушій Unity з інтегрованим пакетом ML-Agents Toolkit, що дозволяє створювати фізично коректні симуляції та отримувати візуальні дані високої точності. Для забезпечення інтерфейсу взаємодії між середовищем та Unity з метою тестування алгоритмів навчання застосовано бібліотеку Gymnasium. На відміну від класичних підходів, де навчання відбувається на фіксованих картах, що часто призводить до перенавчання агентів під конкретну топологію, у роботі запропоновано підхід, що базується на безперервній зміні конфігурації оточення [1]. Для формування тренувальних полігонів нами було реалізовано алгоритми процедурної генерації перешкод, що дозволяє створювати унікальні сценарії для кожного епізоду навчання. Зокрема, для імітації дрібних хаотичних уламків та перешкод застосовано алгоритми генерації білого шуму, для створення плавних ландшафтних структур та рельєфу використано Сімплес шум, а моделювання складних структурованих перешкод, як печери або коридори будівель реалізовано за допомогою Cellular Automata та DFS (Depth-First-Search).

Експериментальне дослідження проводилося у середовищах зростаючої складності, починаючи від спрощеного дискретного 2D Grid World для відпрацювання базової логіки, переходячи до 3D Grid World

і завершуючи повноцінним фізичним середовищем Continuous Physics World, де дрон керується силами тяги та обертальними моментами. Сенсорна система агента базується на наборі віртуальних променевих сенсорів (Raycast sensors), які імітують LIDAR, що надають інформацію про відстань до перешкод. Враховуючи, що агент отримує лише локальну інформацію про середовище, задача навігації формалізується як частково спостережуваний марковський процес прийняття рішень (POMDP). Для компенсації обмеженої видимості в архітектуру нейронної мережі інтегровано рекурентні модулі довгої короткострокової пам'яті (LSTM), що дозволяють агенту утримувати контекст попередніх станів [2].

Основним алгоритмом навчання обрано PPO (Proximal Policy Optimization), який демонструє кращу стабільність навчання та вищу ефективність у складних сценаріях з багатьма перешкодами. Проте для задач, що вимагають тонкого контролю швидкості та орієнтації, алгоритм SAC (Soft Actor-Critic) показує кращі результати. Для прискорення процесу навчання та уникнення потрапляння у локальні оптимуми застосовано методику навчання за розкладом (Curriculum Learning) [3], за якої складність середовища та щільність перешкод зростають поступово по мірі покращення навичок агента. Додатково впроваджено механізм навчання на основі допитливості (Curiosity-Driven Learning) [4] з використанням модуля внутрішньої винагороди, що стимулює агента досліджувати нові стани середовища та ефективно вирішувати навігаційні задачі у лабіринтах складної конфігурації.

Результати моделювання підтверджують, що інтеграція процедурної генерації з LSTM-модулями дозволяє отримати навігаційну політику, здатну ефективно працювати у нових середовищах. Система демонструє високу точність уникнення колізій та досягнення цільових координат, що відкриває перспективи її застосування у реальних пошуково-рятувальних операціях. Подальші дослідження будуть зосереджені на проблематиці перенесення навчених моделей (Sim-to-Real) на фізичні платформи дронів.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Josh Tobin, Rachel Fong, Alex Ray. Domain Randomization for Transferring Deep Neural Networks from Simulation to the Real World [Електронний ресурс] – Режим доступу до ресурсу: <https://arxiv.org/abs/1703.06907>.
2. Victoria Hodge, Richard Hawkins. Deep reinforcement learning for drone navigation using sensor data [Електронний ресурс] – Режим доступу до ресурсу:

https://www.researchgate.net/publication/342348844_Deep_reinforcement_learning_for_drone_navigation_using_sensor_data.

3. Iskandar, A., & Kovács, B. (2024). Investigating the Impact of Curriculum Learning on Reinforcement Learning for Improved Navigational Capabilities in Mobile Robots. *Inteligencia Artificial*, 27(73), 163–176. [Електронний ресурс] – Режим доступу до ресурсу: <https://doi.org/10.4114/intartif.vol27iss73pp163-176>.

4. Deepak Pathak, Pulkit Agrawal, Alexei A. Efros and Trevor Darrell. Curiosity-driven Exploration by Self-supervised Prediction. In *ICML 2017* [Електронний ресурс] – Режим доступу до ресурсу: <https://pathak22.github.io/noreward-rl/resources/icml17.pdf>.

Рябов Д.М., Пенко В.Г.

*Одеський Національний університет
імені І.І. Мечникова, м. Одеса, Україна*

МЕТОД ПІДВИЩЕННЯ ЕФЕКТИВНОСТІ НАВЧАННЯ З ПІДКРІПЛЕННЯМ ЗАВДЯКИ КЛАСТЕРИЗАЦІЇ ПРОСТОРУ СТАНІВ

Анотація

У роботі пропонується новий метод дослідження простору станів у задачах навчання з підкріпленням [1], заснований на кластеризації та імовірнісній оцінці перспективності кластерів. Запропонований підхід групує спостереження в кластери та динамічно розподіляє увагу агента залежно від їх інформаційної цінності. Перспективність кластера визначається на основі середньої накопиченої винагороди в кластері. Експерименти показали, що запропонований метод здатний перевершувати або зрівнюватися за ефективністю з популярними стратегіями дослідження [2].

Ключові слова: навчання з підкріпленням; кластеризація; дослідження станів, агент, стратегія

1. Опис алгоритму

1.1. Представлення станів

Традиційно кожен стан середовища будемо позначати як s_i і представляти у формі вектора числових ознак $\{\phi_i\} \in \mathbb{R}^d$.

Спочатку в ході експериментів ми обираємо певні стани середовища, і отриману множину зібраних станів на поточному етапі

$\mathcal{D} = \{s_1, s_2, \dots, s_N\}$, де N – кількість станів.

1.2. Кластеризація

На множині станів $\mathcal{D}$ виконується кластеризація за допомогою алгоритму К-середніх (або іншого відповідного методу), в результаті чого формується набір кластерів:

$$\mathcal{C} = \{C_1, C_2, \dots, C_K\}.$$

1.3. Оцінка перспективності кластерів

Для кожного кластера C_k обчислюється показник якості q_k , що відображає його «перспективність» для подальшого дослідження.

$$q_k = V(C_k) = f(\{s_i \in C_k\}),$$

де $V(C_k)$ - традиційно використовувана в навчанні з підкріпленням функція цінності стану.

1.4. Призначення вірогідності розширення кластерів

На основі оцінок q_k формуються ймовірності p_k розширення відповідних кластерів. Для зручності використовується нормалізація:

$$p_k = \frac{\exp(\alpha q_k)}{\sum_{j=1}^K \exp(\alpha q_j)}$$

де параметр $\alpha > 0$ регулює ступінь «фокусування» на найбільш перспективних кластерах.

1.5. Вибір кластера та розширення

На кожному кроці або епізоді агент вибирає кластер для пріоритетного дослідження, випадково відповідно до розподілу $\{p_k\}_{k=1}^K$.

Потім всередині обраного кластера проводиться внутрішній семплінг станів за рахунок виконання дій, спрямованих на отримання нових даних, що розширюють цей кластер.

Після чого знову проводимо послідовне розбиття станів на кластери, оцінку їх перспективності, розширення найбільш пріоритетних кластерів і подальшу переоцінку цих показників у міру накопичення досвіду.

1.6. Періодична переоцінка та адаптація

Після такого розширення і накопичення нових станів процес переоцінки якості кластерів і оновлення ймовірностей p_k повторюється.

1.7. Інтеграція з RL-агентом

Запропонований механізм може бути реалізований як додатковий модуль для базового RL-алгоритму (наприклад, DQN, PPO або A3C), що впливає на стратегію вибору станів або дій.

2. Обчислювальний експеримент

В рамках дослідження проведено порівняльний аналіз п'яти стратегій дослідження в завданнях навчання з підкріпленням: випадкове дослідження, Count-based Exploration, Модельно-орієнтований, Prioritized Experience Replay, а також кластерно-адаптивний метод з імовірнісним розширенням кластерів. Тестування проводилося в трьох типових середовищах з пакету Gymnasium: CartPole-v1, MountainCar-v0 і FrozenLake-v1.

Результати показують, що ефективність стратегії залежить від властивостей середовища.

Кластерно-адаптивний метод продемонстрував конкурентоспроможні або кращі результати у всіх трьох середовищах, завдяки здатності переоцінювати пріоритети дослідження на основі спостережуваних нагород у кластерах станів. Він особливо ефективний у завданнях з вираженою структурою і необхідністю балансувати між дослідженням і експлуатацією.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Sutton, R. S., & Barto, A. G. (2018). Reinforcement Learning: An Introduction (2nd ed.). MIT Press.
2. Ladosz, P., Weng, L., Kim, M., & Oh, H. (2022). Exploration in Deep Reinforcement Learning: A Survey. arXiv:2205.00824.

УДК 004.42

*Трунов О.М., Савчук А.О.,
Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ПОРІВНЯЛЬНИЙ АНАЛІЗ НЕЙРОМЕРЕЖЕВИХ МЕТОДІВ СЕГМЕНТАЦІЇ ЗОБРАЖЕНЬ ІНТЕРФЕЙСІВ АСК ДЛЯ АВТОМАТИЗОВАНОГО КОГНІТИВНОГО АНАЛІЗУ

Проведено порівняльний аналіз сучасних нейромережевих методів сегментації зображень інтерфейсів автоматизованих систем керування (АСК), зокрема U-Net, Mask R-CNN, DeepLabv3+, Segment Anything Model (SAM) та Vision Transformers (SegFormer). Розглянуто їхню архітектуру, алгоритми піксельної та об'єктної сегментації, вимоги до навчання та здатність працювати зі складними UI-

структурами. Запропоновано експериментальну схему оцінювання, що включає метрики IoU, Dice, Precision/Recall та когнітивні показники: щільність елементів, візуальне навантаження та відхилення від еталонної структури інтерфейсу. Показано, що моделі Transformer-типу забезпечують найвищу генералізацію, SAM — найстійкішу роботу без навчання, а U-Net і Mask R-CNN — кращий компроміс для локалізованих елементів контролю. Визначено підходи до формування еталонних карт сегментації для систем автоматизованого когнітивного аналізу та описано застосування статистичних тестів (*t*-тест Велча, розмір ефекту Коена) для порівняння альтернативних UI-рішень.

Ключові слова: сегментація; машинний зір; нейронні мережі; інтерфейси АСК; когнітивний аналіз; SAM; U-Net; Mask R-CNN.

ВСТУП

Ефективність роботи оператора автоматизованих систем керування (АСК) критично залежить від ергономіки інтерфейсу та швидкості зчитування інформації. Застосування методів автоматизованого когнітивного аналізу дозволяє об'єктивно оцінити ці параметри, не залучаючи дорогих експертів. Проте, надійність такої оцінки напряму залежить від того, наскільки точно алгоритми комп'ютерного зору (Computer Vision) здатні розпізнати та виокремити елементи керування, індикатори й мнемосхеми на екрані.

МЕТА ДОСЛІДЖЕННЯ

Мета дослідження — проведення порівняльного аналізу сучасних нейромережових методів сегментації та визначення оптимального підходу для формування семантичних карт інтерфейсів, які надалі використовуються для розрахунку когнітивних метрик.

ОГЛЯД ІСНУЮЧИХ МОДЕЛЕЙ

У дослідженні ми розглянули як класичні згорткові мережі (CNN), так і новітні архітектури на базі трансформерів (ViT), які на сьогодні є стандартом у задачах обробки зображень.

Таблиця 1 – Характеристика досліджуваних нейромережових моделей

№	Модель	Сутність архітектури	Основні метрики	Переваги	Недоліки
1	U-Net	Енкодер-декодер зі skip-connections	IoU, Dice	Чудова точність визначення меж, добре працює на малих вибірках	Важко враховує глобальний контекст всього екрану
2	Mask R-CNN	Двоетапна Instance Segmentation	mAP, Mask IoU	Дозволяє розділити окремі екземпляри кнопок, що стоять поруч	Повільна робота, високі вимоги до обчислювальних ресурсів
3	DeepLab v3+	Atrous Spatial Pyramid Pooling (ASPP)	mIoU	Ефективно захоплює об'єкти різного масштабу	Може втрачати дрібні деталі на стиках UI-елементів
4	SegFormer	Ієрархічний Vision Transformer	mIoU, Accuracy	Стійкість до шумів, врахування глобального контексту	Потребує значного обсягу даних для навчання
5	SAM	Promptable Segmentation (Zero-shot)	Prompt IoU	Універсальність, працює без донавчання	Надмірна ресурсоемність, схильність до надлишкової деталізації

Таблиця 2 – Порівняння ефективності підходів

Підхід	Тип архітектури	Точність (IoU)	Швидкість (FPS)	Здатність до генералізації
U-Net	CNN (FCN)	Висока	Середня	Низька
Mask R-CNN	CNN + ROI Align	Дуже висока	Низька	Середня
DeepLab v3+	CNN (Dilated)	Середня	Середня	Середня
SegFormer	Transformer	Висока	Висока	Висока
SAM	ViT-based	Змінна	Низька	Дуже висока

МЕТОДИКА ПОРІВНЯННЯ

Для коректного порівняння методів ми використали ідентичні набори даних інтерфейсів АСК. Оцінка проводилася за такими критеріями:

– **Об'єктивні метрики:** Intersection over Union (IoU), Dice Score та Boundary F1 (точність країв).

– **Технічні показники:** час інференсу та вимоги до GPU.

– **Інтегральна оцінка:** оскільки суто технічні метрики не завжди відображають придатність маски для аналізу, було запропоновано індекс якості Q:

$$Q = \sum w_i q_i$$

де q_i — це нормовані значення точності та відхилення від еталонної структури. Для перевірки статистичної значущості відмінностей між моделями застосовувався t-тест Велча разом із розрахунком розміру ефекту d Коена.

ЕКСПЕРИМЕНТАЛЬНА ЧАСТИНА

Експерименти проводилися на власному датасеті, що складається з 300 скріншотів реальних SCADA-систем.

Таблиця 3 – Приклад результатів оцінювання (фрагмент)

ID	Модель	IoU	Dice	Час (мс)	C_dev (відхилення)	Q (індекс)
1	U-Net	0.8 9	0.91	45	0.05	0.88
2	U-Net	0.8 7	0.90	48	0.07	0.85
3	SAM	0.9 2	0.94	120	0.02	0.91
4	SAM	0.9 1	0.93	115	0.03	0.90
5	Mask R-CNN	0.9 0	0.92	150	0.04	0.87
6	SegFormer	0.8 8	0.89	30	0.06	0.86

РЕЗУЛЬТАТИ

Отримані дані, підтверджені t-тестом Велча ($p < 0,05$), дозволяють стверджувати, що вибір архітектури суттєво залежить від задачі:

- U-Net залишається надійним вибором для чітко стандартизованих інтерфейсів, забезпечуючи високу точність меж,
- SAM демонструє вражаючу здатність працювати з новими, незнайомими інтерфейсами (Zero-shot), що робить його ідеальним інструментом для автоматичної розмітки датасетів, хоча його використання в реальному часі обмежене ресурсами,
- SegFormer завдяки механізму Self-Attention показує найкращий баланс між швидкістю роботи та якістю сегментації;

ВИСНОВКИ

Запропонований підхід, що поєднує нейромережеву сегментацію зі статистичною верифікацією результатів, дозволяє інженерам обґрунтовано обирати архітектуру під конкретний клас систем АСК. Використання інтегрального індексу Q дає змогу врахувати не лише піксельну точність, а й придатність отриманих даних для подальшого когнітивного аналізу.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Ronneberger O., Fischer P., Brox T. U-Net: Convolutional Networks for Biomedical Image Segmentation. MICCAI, 2015.
2. He K., Gkioxari G., Dollár P., Girshick R. Mask R-CNN. Proceedings of the IEEE International Conference on Computer Vision (ICCV), 2017.
3. Goodfellow I., Bengio Y., Courville A. Deep Learning. MIT Press, 2016.
4. Zhang Y., Li K., Wang W. A Review on Image Segmentation Techniques Using Deep Learning. IEEE Access, 2020.

УДК 004.42

*Салютін М. О., Кондратенко Ю. П., Сіденко Є. В.
Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

МЕТОДИ UNDER- ТА OVERSAMPLING У ЗАДАЧІ РОЗПІЗНАВАННЯ ВІЙСЬКОВИХ ОБ'ЄКТІВ

У сучасних системах комп'ютерного зору задача розпізнавання військових об'єктів залишається однією з найскладніших через надзвичайно незбалансовану природу реальних наборів даних. У більшості сценаріїв військові об'єкти займають лише малу частину зображень порівняно з фоновим середовищем (лісами, полями, пустелями тощо), що призводить до значного дисбалансу класів і суттєвого погіршення якості детекції рідкісних цілей. Незбалансованість та маленькі розміри об'єктів які треба розпізнавати особливо критична в оборонних застосуваннях, де пропуск навіть одного стратегічно важливого об'єкта може мати неприпустимі наслідки. Проблема посилюється дефіцитом якісно підготовлених відкритих наборів даних для військової тематики. Існуючі набори даних або містять недостатню кількість прикладів окремих класів техніки, або мають значну кількість некоректно розмічених зображень, або взагалі відсутні через режим секретності. Це суттєво обмежує результати навчання глибоких нейронних мереж сучасної архітектури, які потребують ретельної попередньої обробки даних для кращого результату і обробки всіх можливих сценаріїв [1].

Традиційні підходи до обробки зображень, такі як аугментація поворотами, змінами яскравості чи масштабуванням, хоча й покращують узагальнення, не вирішують проблему розподілу класів з малою кількістю об'єктів у порівнянні із домінуючими. У таких умовах

методи undersampling [2] більшості класу та oversampling [3] меншості набувають особливої актуальності, оскільки дозволяють безпосередньо впливати на розподіл прикладів у тренувальній вибірці та забезпечувати стабільніше навчання детекторів об'єктів навіть за умов обмеженої кількості якісних даних.

Для оцінки ефективності методів балансування даних у задачі розпізнавання військових об'єктів було використано відкритий набір даних Military Decision-Making Dataset [4], що містить 16 809 анотованих зображень у форматі YOLO. Об'єкти розподілено за одинадцятьма класами військової техніки: танк (TANK), бойова машина піхоти (IFV), бронетранспортер (APC), інженерна машина (EV), штурмовий гелікоптер (AH), транспортний гелікоптер (TH), штурмовий літак (AAP), транспортний літак (TA), зенітна установка (AA), буксирувана артилерія (TART) та самохідна артилерія (SPART). Зображення отримані з реальних сцен бойових дій, техніка представлена в широкому діапазоні ракурсів, масштабів, умов освітлення та ступеня маскування, що забезпечує високу варіативність тренувальної вибірки та сприяє формуванню стійких до реальних умов моделей детекції.

Навчання проводилося на базі архітектури YOLOv11 з використанням двох конфігурацій малої моделі (YOLOv11-s), що відрізнялися лише стратегією балансування даних. Усі експерименти виконувалися з роздільною здатністю зображень 1024×1024 пікселі, оптимізатором AdamW, загальною кількістю епох 100 та механізмом ранньої зупинки (patience = 20).

Перша модель була навчена без додаткового балансування, тоді як друга включала комплексне застосування методів under-/oversampling з подальшим контролем якості розмітки, що й стало предметом порівняльного аналізу. При підготовці даних виявлено 16 809 анотованих зображень, проте після автоматичної перевірки розмітки вилучено 471 зображення з некоректними координатами bounding box (вихід за межі [0;1] та від'ємні значення). Остаточний чистий набір даних склав 16 338 зображень. Подальший аналіз розподілу об'єктів по класах підтвердив значний дисбаланс: класи APC і TANK містили понад 6000 екземплярів кожен, тоді як класи IFV, EV, TH, AAP та AA мали менше 1500 об'єктів, що створювало високий ризик погіршення детекції рідкісних цілей.

Для усунення дисбалансу застосовано комбінацію методів undersampling і oversampling. На етапі undersampling кількість об'єктів домінуючих класів (TANK, APC, SPART) штучно обмежена верхньою межею 3000 екземплярів на клас шляхом вибіркового виключення

зображень, що містили надлишкові приклади цих категорій. На етапі oversampling для малопредставлених класів (IFV, EV, TH, AAP, AA) виконано цілеспрямоване генерування нових зображень з використанням бібліотеки Albumentations до досягнення тієї ж межі 3000 об'єктів на клас. Застосовувалися лише детекційно-безпечні трансформації: горизонтальне віддзеркалення ($p=0.5$), випадкова зміна яскравості та контрасту ($p=0.4$), зміна відтінку та насиченості ($p=0.3$), розмиття руху ($p=0.2$) та поворот до $\pm 10^\circ$ ($p=0.5$). Координати bounding box автоматично перераховувалися для кожної аугментації з збереженням нормалізованого формату YOLO.

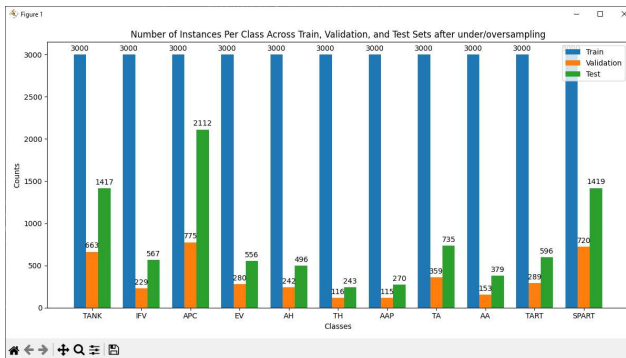


Рисунок 1 – Збалансований набір даних

В результаті сформовано збалансовану тренувальну вибірку, в якій усі одинадцять класів мають приблизно однакову кількість об'єктів (3000 на кожен клас).

Базова модель, навчена на оригінальному незбалансованому наборі даних, продемонструвала високі загальні показники завдяки домінуванню поширених класів, проте при детальному аналізі по класах виявлено суттєве падіння точності на рідкісних категоріях (IFV, EV, TH, AAP, AA). Модель, навчена на збалансованій вибірці з використанням комбінованих методів under- та oversampling, досягла кращих узагальнених метрик за більш жорстким критерієм $mAP@0.5:0.95$ та показала стабільно високу точність по всіх одинадцяти класах без значного просідання на малопредставлених об'єктах.

Таблиця 1 – Результати навчання обраних моделей

Тип моделі	mAP@0.5	mAP@0.5:0.95	Precision	Recall	F1-Score
Незбалансована модель	0.990	0.857	0.9846	0.9735	0.979
Збалансована модель	0.993	0.878	0.9892	0.9801	0.985

Виходячи із таблиці 1 можна зробити висновок, що балансування класів за допомогою комбінації undersampling домінуючих і oversampling рідкісних класів дозволяє суттєво підвищити якість детекції військових об'єктів у реальних умовах експлуатації, особливо для стратегічно важливих малопоширених типів техніки, і є необхідним етапом підготовки даних при використанні сучасних моделей комп'ютерного зору у задачах з природно незбалансованим розподілом класів.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Zhang, C., Zhou, W., Zhao, X. et al. A machine learning framework for missing and imbalanced data in marketing analytics. J Market Anal (2025). <https://doi.org/10.1057/s41270-025-00432-4>
2. Cheruvathoor, L.G., Thomas, J. (2025). Unbalanced Dataset Preprocessing Using Hybrid Combination Algorithm for Arrhythmia Detection. In: Kriksciuniene, D., Bajaj, A., Abraham, A. (eds) Bio-Inspired Computing. IBICA 2023. Lecture Notes in Networks and Systems, vol 1228. Springer, Cham. https://doi.org/10.1007/978-3-031-78937-3_10
3. Shaikh, M.Z., Jatoi, S., Baro, E.N. et al. FaultSeg: A Dataset for Train Wheel Defect Detection. Sci Data 12, 309 (2025). <https://doi.org/10.1038/s41597-025-04557-0>
4. Military Decision-Making Dataset. <https://www.kaggle.com/datasets/nzigulic/military-equipment>, 2024. Accessed: 22-Jul-2025.

ІНТЕЛЕКТУАЛЬНА СИСТЕМА КЛАСИФІКАЦІЇ ПАЛЬНОГО НА ОСНОВІ МЕТОДІВ МАШИННОГО НАВЧАННЯ

Класифікація типів пального є важливою задачею в нафтопереробній промисловості, автомобільному секторі та авіації, оскільки точне визначення характеристик пального забезпечує якість продукції, безпеку використання та відповідність екологічним стандартам. Традиційні методи лабораторного аналізу є трудомісткими, вимагають спеціалізованого обладнання та займають значний час. Впровадження інтелектуальних систем на основі машинного навчання дозволяє автоматизувати процес класифікації та підвищити швидкість прийняття рішень.

Розробка інтелектуальної системи класифікації пального здійснюється з метою автоматизації процесу визначення типу пального на основі його фізико-хімічних характеристик, підвищення точності класифікації та зменшення часу на проведення аналізу. Для досягнення поставленої мети використано алгоритм Random Forest (RF), який показав високу ефективність при роботі з багатовимірними числовими даними та забезпечив можливість оцінки важливості кожної ознаки.

Аналіз сучасних підходів до класифікації пального показав, що машинне навчання активно застосовується в нафтохімічній галузі. Дослідники використовують різні алгоритми, включаючи Support Vector Machine (SVM), Artificial Neural Networks (ANN) та ансамблеві методи для прогнозування властивостей пального та його класифікації. Random Forest демонструє переваги завдяки стійкості до перенавчання, здатності обробляти нелінійні залежності та можливості визначення найважливіших ознак для класифікації.

Для навчання та оцінювання ефективності розробленої системи було створено збалансований датасет обсягом 5000 зразків, який включає чотири типи пального: дизельне паливо (1000 зразків), бензин (2500 зразків), гас/авіаційне паливо (1000 зразків) та біопаливо (500 зразків). Датасет базується на реальних стандартах ASTM (American Society for Testing and Materials) та ISO (International Organization for Standardization) і містить такі фізико-хімічні характеристики: октанове число (або цетанове для дизелю), щільність (г/см^3), в'язкість ($\text{мПа}\cdot\text{с}$), температуру спалаху ($^{\circ}\text{C}$) та вміст сірки (ppm). Використання

нормального та логнормального розподілів при генерації даних забезпечило реалістичність зразків, а додавання випадкового шуму імітувало похибки реальних вимірювань.

Архітектура розробленої інтелектуальної системи класифікації пального побудована на основі модульного підходу, що забезпечує гнучкість та можливість подальшого розширення функціоналу. Система складається з наступних компонентів: веб-інтерфейс користувача, реалізований засобами HTML, CSS та JavaScript, що забезпечує інтуїтивне введення параметрів пального та візуалізацію результатів; серверна частина на Flask (Python), яка обробляє HTTP-запити та координує роботу всіх модулів; модуль попередньої обробки даних, що виконує нормалізацію вхідних параметрів за допомогою StandardScaler; модуль машинного навчання на основі Random Forest з 100 деревами рішень; модуль візуалізації результатів, який відображає ймовірності для кожного класу у вигляді інтерактивних графіків; сховище моделей у форматі pickle для швидкого завантаження навченої моделі та scaler'a. Така архітектура забезпечує повний цикл обробки даних – від введення характеристик користувачем до отримання результату класифікації з відповідними ймовірностями (див. рис. 1).

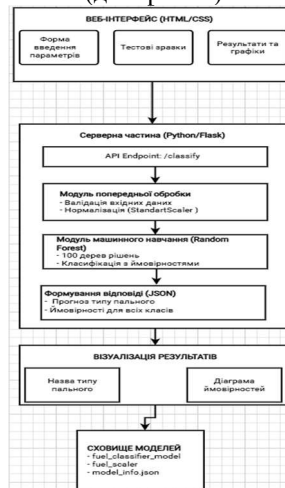


Рисунок 1 – Схема архітектури інтелектуальної системи класифікації пального

У процесі навчання моделі Random Forest на підготовленому датасеті було досягнуто високої ефективності класифікації. Дані були розділені на навчальну (80%, 4000 зразків) та тестову (20%, 1000 зразків) вибірки зі збереженням пропорцій класів (stratified sampling).

Перед навчанням всі числові ознаки були нормалізовані за допомогою StandardScaler для приведення до однакового масштабу. Модель Random Forest з параметрами $n_estimators=100$, $max_depth=15$, $min_samples_split=5$, $min_samples_leaf=2$ показала точність класифікації 97.2% на тестовій вибірці. Аналіз важливості ознак (feature importance) виявив, що найбільший вплив на класифікацію мають октанове число (35%), температура спалаху (28%) та щільність (22%), тоді як в'язкість (10%) та вміст сірки (5%) мають менший, але все ж значущий внесок. Матриця плутанини показала високу точність розпізнавання для всіх класів: дизель – 98%, бензин – 97%, гас – 96%, біопаливо – 95%. Час навчання моделі склав 2.3 секунди, а час класифікації одного зразка – менше 0.01 секунди, що підтверджує придатність системи для практичного застосування в режимі реального часу.

Веб-інтерфейс системи забезпечує зручну взаємодію з користувачем та наочну візуалізацію результатів. Форма введення дозволяє вводити п'ять параметрів пального з відповідними підказками про допустимі діапазони значень. Реалізовано набір тестових зразків для швидкої перевірки роботи системи: Дизель, Бензин, Гас та Біопаливо. Результати класифікації відображаються у вигляді назви типу пального українською мовою та інтерактивної діаграми з ймовірностями для кожного класу, відсортованими за спаданням. Інформаційна панель показує метрики моделі: загальну точність, розмір датасету, кількість класів та кількість аналізованих параметрів. Адаптивний дизайн забезпечує коректне відображення на різних пристроях, включаючи мобільні телефони та планшети.

Таким чином, реалізована інтелектуальна система класифікації пального на основі Random Forest показала високу ефективність (97.2% точності) у розпізнаванні чотирьох типів пального за їх фізико-хімічними характеристиками. Поєднання алгоритму машинного навчання з попередньою обробкою даних, веб-інтерфейсом та модулем візуалізації забезпечило точну класифікацію, швидкість та зручну взаємодію з користувачем. Система може бути використана в нафтопереробній промисловості, на нафтобазах, АЗС та в лабораторіях контролю якості для автоматизації процесу ідентифікації типів пального.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. ASTM D1655 - Standard Specification for Aviation Turbine Fuels. *ASTM International*. URL: <https://www.astm.org/d1655-24.html>.
2. ISO 3405:2019 - Petroleum and related products - Determination of distillation characteristics at atmospheric pressure. *International Organization for Standardization*. URL: <https://www.iso.org/standard/72795.html>.
3. Breiman, L. (2001). Random Forests. *Machine Learning*, 45(1), 5-32. URL: <https://doi.org/10.1023/A:1010933404324>.
4. Pedregosa, F., et al. (2011). Scikit-learn: Machine Learning in Python. *Journal of Machine Learning Research*, 12, 2825-2830. URL: <https://jmlr.org/papers/v12/pedregosa11a.html>.
5. Fuel Properties Dataset. *Kaggle*. URL: <https://www.kaggle.com/datasets/fuel-properties> (приклад типового джерела датасетів).
6. Flask Web Development Framework. *Official Documentation*. URL: <https://flask.palletsprojects.com/>.

УДК 004.85

Стрюк О.С.,

*Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

РОЗГОРТАННЯ ТА ІНЖЕНЕРНА ОПТИМІЗАЦІЯ ГЕНЕРАТИВНИХ ЗМАГАЛЬНИХ МЕРЕЖ НА КОРДОННИХ СИСТЕМАХ

У роботі представлено практичний підхід до розгортання та оптимізації генеративних змагальних мереж (ГЗМ) на кордонних системах (англ. edge systems) в умовах апаратно-параметричних обмежень, з використанням Raspberry Pi 5, як тестового стенда. Дослідження вирішує ключові виклики, притаманні edge-пристроєм: обмежену обчислювальну потужність, брак пам'яті та вимоги до енергоефективності [1, 2, 3].

Завданням дослідження було розгорнути та оптимізувати ГЗМ на Raspberry Pi 5 для генерації графічних даних в реальному часі. Основними викликами були обмежені ресурси процесора та графічного модуля Raspberry Pi 5, а також його енергоспоживання [3, 4].

Для подолання цих обмежень було реалізовано комплексний набір методів інженерної та програмної оптимізації як на рівні моделі, так і на рівні коду. Ключові стратегії включали [5]:

1. Квантування та прунінг (стиснення) моделі. Для запуску ГЗМ на Raspberry Pi 5 було проведено попереднє стиснення моделей. Квантування, метод зниження точності числових параметрів моделі (наприклад, перетворення 32-бітних чисел у 8-бітні), допомогло економити пам'ять. Прунінг, метод «обрізання» (видалення) менш важливих ваг у моделі, дозволив зменшити її розмір. Це дозволило зменшити вимоги до пам'яті та обчислювальних потужностей без значної втрати якості генерованих зображень.

2. Для проектування та навчання моделі було використано фреймворк PyTorch, що дозволив оптимізувати інференцію моделі, скоротити час навчання та зменшити затримку при генерації зображень.

3. Також було використано інструменти профілювання на базі PyTorch Profiler, а також розглянуто перспективи подальшої оптимізації шляхом використання TFLite, ONNX, TensorRT, моделювання апаратного прискорення та застосування методів дистиляції знань.

Під час експерименту було перевірено здатність платформи Raspberry Pi 5 забезпечити стабільну роботу ГЗМ при генерації зразків рукописних цифр. Використовуючи оптимізовану архітектуру, Raspberry Pi 5 забезпечив генерацію зразків у реальному часі. ГЗМ було успішно навчено безпосередньо на Raspberry Pi 5 після додаткового налаштування моделі, і важливо, що це було зроблено без використання спеціалізованих апаратних прискорювачів (як-от Google Coral) [5].

Під час проведення експериментів одним з головних інженерних викликів, окрім обчислювальних обмежень, став перегрів пристрою при тривалій роботі. Для забезпечення стабільної роботи системи було впроваджено додаткові заходи з активного охолодження: зовнішній корпус з додатковим кулером (з температурним контролем та PWM-керуванням) та алюмінієвий радіатор для покращеного відведення тепла.

Проведені експерименти підтвердили можливість ефективного впровадження ГЗМ на кордонних пристроях, аналогічних Raspberry Pi 5, при відповідній інженерній оптимізації. Досвід розгортання продемонстрував важливість розробки спеціалізованих алгоритмів оптимізації, орієнтованих на обмежені обчислювальні можливості.

Результати дослідження свідчать, що інженерні підходи до оптимізації ГЗМ здатні значно розширити область їх застосування в IoT, мобільній робототехніці та інших кордонних пристроях реального часу [3, 4].

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Goodfellow I. et al. Generative adversarial nets. *Advances in Neural Information Processing Systems* 27. Curran Associates, Inc., 2014. P. 2672–2680. URL: <http://papers.nips.cc/paper/5423-generative-adversarial-nets.pdf>.
2. Kingma D. P., Rezende D. J., Mohamed S., Welling M. Semi-Supervised Learning with Deep Generative Models. *Proceedings of the International Conference on Machine Learning*. 2014. P. 3581–3589. URL: <https://arxiv.org/abs/1406.5298>.
3. Стрюк О. С., Кондратенко Ю. П. Методи прикладного застосування генеративних змагальних мереж при обробці графічних даних. *Штучний інтелект*. 2023. № 3. С. 154–161. DOI: <https://doi.org/10.15407/jai2023.03.154>.
4. Striuk O., Kondratenko Y. Implementation of Generative Adversarial Networks in Mobile Applications for Image Data Enhancement. *Journal of Mobile Multimedia*. 2023. Vol. 19, no. 03. P. 823–838. DOI: <https://doi.org/10.13052/jmmm1550-4646.1938>.
5. Striuk O., Kondratenko Y. Optimization Strategy for Generative Adversarial Networks Design. *International Journal of Computing*. 2023. Vol. 22, no. 3. P. 292–301. DOI: <https://doi.org/10.47839/ijc.22.3.3223>.

УДК 004.89:004.912

Хіньов Д. О., Кулаковська І. В.

*Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ІНТЕЛЕКТУАЛЬНА СИСТЕМА КЛАСИФІКАЦІЇ МУЗИЧНИХ ЖАНРІВ НА ОСНОВІ СПЕКТРОГРАМ ІЗ ВИКОРИСТАННЯМ ЗГОРТКОВИХ НЕЙРОННИХ МЕРЕЖ

Стрімке зростання обсягів аудіоконтенту у цифровому просторі зумовлює потребу в автоматизованих системах обробки та аналізу музичних композицій. Сучасні музичні сервіси на кшталт Spotify, YouTube Music та Apple Music щодня отримують десятки тисяч нових треків, що унеможливорює ручне визначення жанрової приналежності. Жанр є ключовою ознакою для побудови рекомендаційних систем, персоналізації плейлистів, аналізу музичних тенденцій та впорядкування великих аудіобаз. Традиційні підходи до класифікації музики виявляються малоефективними через складність сигналу,

накладання інструментальних партій та нечіткі межі між жанрами, що зумовлює потребу у застосуванні глибинних моделей.

Метою роботи є створення інтелектуальної системи класифікації музичних жанрів на основі спектрограм аудіосигналів, здатної забезпечувати високу точність розпізнавання та інтегруватися у сучасні цифрові платформи. Для дослідження використано датасет GTZAN із 1000 аудіофрагментів десяти жанрів. Попередня обробка включала нормалізацію сигналів, сегментацію аудіо, аугментацію (time-stretching, pitch-shifting, шумові інжекції) та побудову мел-спектрограм і MFCC-ознак, що забезпечило високий рівень інформативності для подальшого аналізу.

Загальну архітектуру інтелектуальної системи класифікації музичних жанрів представлено на рисунку 1 (див. рис. 1). Вона включає послідовні етапи попередньої обробки аудіосигналу, його спектрального перетворення, аналізу згортковою нейронною мережею та відображення прогнозу користувачеві. Така структурна організація дає змогу реалізувати повний цикл обробки — від перетворення звукового сигналу у спектрограму до визначення музичного жанру на основі глибинних ознак.

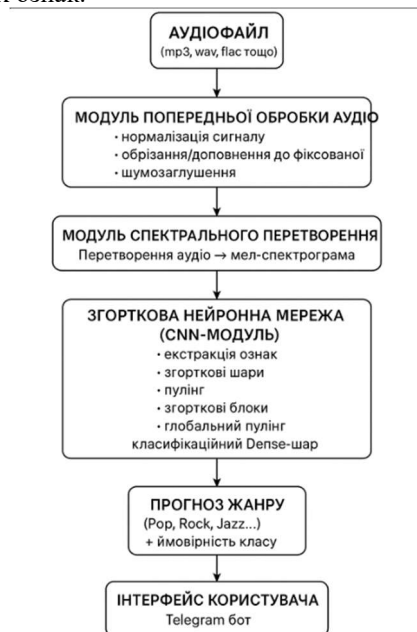


Рисунок 1 – Схема роботи інтелектуальної системи класифікації музичних жанрів

Запропонована система базується на згорткових нейронних мережах (Convolutional Neural Networks, CNN), здатних автоматично вилучати частотно-часові патерни зі спектрограм без необхідності ручного формування ознак. Було розроблено та протестовано кілька архітектур CNN із різними параметрами глибини, фільтрів та регуляризації, що дозволило досягти точності класифікації на рівні 85–90%. Такий результат підтверджує ефективність спектрограмного підходу та глибокого навчання для задач музичної аналітики.

Система реалізує повний технологічний цикл — від аналізу аудіосигналу до отримання прогнозу користувачем. Для забезпечення доступності розроблено Telegram-бот, що дає змогу завантажувати аудіофайли, автоматично формувати їх спектрограми та повертати користувачу визначений жанр у режимі реального часу. Це усуває потребу у встановленні спеціалізованого ПЗ та робить систему готовою до інтеграції у сервіси автоматичної каталогізації, пошуку музичних композицій та формування персоналізованих рекомендацій.

Отримані результати демонструють практичну цінність застосування глибокого навчання у сфері музичної класифікації. Подальші напрями розвитку включають розширення датасету, впровадження гібридних CNN–CRNN архітектур та дослідження можливостей трансформерних моделей для покращення аналізу часових патернів. Інтелектуальна система, розроблена в межах роботи, може стати основою для сучасних музичних аналітичних платформ і сервісів.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Tzanetakis G., Cook P. Musical genre classification of audio signals // IEEE Transactions on Speech and Audio Processing. — 2002. — Vol. 10, No. 5. — P. 293–302.
2. McFee B. et al. librosa: Audio and music signal analysis in Python // Proceedings of the 14th Python in Science Conference (SciPy). — 2015. — P. 18–25.
3. Choi K. et al. Convolutional recurrent neural networks for music classification // 2017 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP). — 2017. — P. 841–845.
4. GTZAN Dataset. URL: <https://www.tensorflow.org/datasets/catalog/gtzan> (дата звернення: 20.11.2025).
5. Humphrey E., Bello J. A tutorial on deep learning for music information retrieval // The Journal of the Acoustical Society of America. — 2017. — Vol. 142, No. 5. — P. 34–40.

6. Zhang X. et al. CNN-based music genre classification using mel-spectrograms // Journal of Physics: Conference Series. — 2021. — Vol. 1874. — P. 012001.

УДК 004.056.5

Цимбал А. Ю., Калініна І. О.

*Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ІНТЕЛЕКТУАЛЬНА СИСТЕМА ІДЕНТИФІКАЦІЇ ФІШИНГУ

Фішингові атаки залишаються одним із найпоширеніших видів кіберзагроз у сучасному цифровому середовищі, оскільки вони спрямовані на обман користувачів з метою отримання конфіденційної інформації. Зростання кількості таких атак вимагає впровадження інтелектуальних систем, здатних автоматично розпізнавати фішингові повідомлення та посилання (Uniform Resource Locator, URL).

Розробка інтелектуальної системи ідентифікації фішингових повідомлень та URL здійснюється з метою підвищення рівня кібербезпеки, своєчасного попередження користувачів про потенційні загрози та зниження ризику несанкціонованого доступу до конфіденційної інформації. Для досягнення поставленої мети використано такі методи машинного навчання (Machine Learning, ML), як Naive Bayes (NB), Logistic Regression (LR), Support Vector Machine (SVM), Random Forest (RF) та Gradient Boosting (GB). Застосування цих алгоритмів дозволило провести багатофакторний аналіз текстових та структурних ознак, що підвищило точність класифікації.

Аналіз аналогічних досліджень показав ефективність різних підходів до виявлення фішингових атак. Наприклад, Н. Saleem *et al.* [1] поєднали систему правил із ML (SVM, RF) для вебфішингу та інтегрували результати в розширення мережі Phish Net. Автори М. Hassnain *et al.* [2] застосували штучний інтелект на основі візуальної схожості для виявлення zero-day атак із точністю до 84%. При цьому, С. V. Swetha *et al.* [3] розробили мультимодальну архітектуру для фішингових листів із трансформерами та крос-увагою, досягнувши 99% точності. Це підтверджує ефективність комбінування правил, ML і трансформерів для підвищення надійності систем кіберзахисту.

Для навчання та оцінювання ефективності розробленої інтелектуальної системи виявлення фішингових повідомлень і URL було використано кілька відкритих та власних наборів даних. Основним

є «Phishing Dataset» [4] з платформи Hugging Face, який охоплює тексти, електронні листи, SMS (Short Message Service), URL-адреси й HTML-коди (HyperText Markup Language), збалансовані за класами. Додатково застосовано «Phishing Email Dataset» [5] і «Phishing URL Classifier Dataset Cleaned» [6] із платформи Kaggle, що містять структуровані вибірки електронних листів і URL із числовими ознаками. Сукупне використання цих джерел забезпечило репрезентативність, збалансованість і відтворюваність результатів, а також підвищило точність і узагальнювальну здатність моделей класифікації у системі.

Архітектура розробленої інтелектуальної системи ідентифікації фішингових повідомлень і URL побудована на основі модульного підходу, що забезпечує її гнучкість, масштабованість та зручність у подальшій інтеграції з іншими аналітичними сервісами. Система включає інтерфейс користувача, реалізований засобами HTML, CSS (Cascading Style Sheets) та JavaScript; серверну частину на Flask (Python), яка обробляє запити користувача та виконує класифікацію повідомлень і URL-адрес; модулі попередньої обробки тексту та URL (очищення, визначення мови, переклад, векторизація); блоки ML, що використовують алгоритми NB, LR, SVM, RF і GB; модуль рекомендацій і зворотного зв'язку, а також сховище даних для моделей і результатів. Така структура забезпечує повний цикл обробки даних – від введення повідомлення користувачем до формування рекомендацій і збереження результатів аналізу (див. рис. 1).

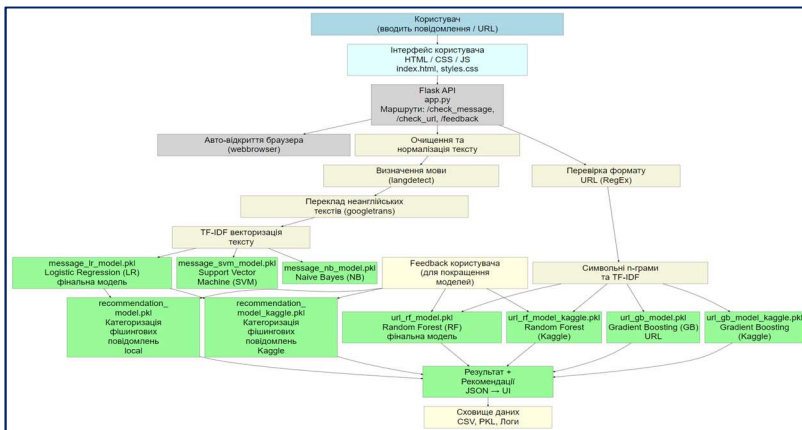


Рисунок 1 – Схема архітектури інтелектуальної системи ідентифікації фішингових повідомлень і URL

У процесі навчання моделей ML для класифікації фішингових повідомлень і URL-адрес було оцінено ефективність різних алгоритмів. Моделі для обробки текстових повідомлень (NB, LR, SVM) показали точність понад 90%, причому LR продемонструвала найкращу узагальнювальну здатність, а SVM – найвищий рівень класифікації для складних мовних структур. Моделі для аналізу URL (RF та GB) досягли точності до 97%, підтверджуючи ефективність ансамблевих методів при роботі з великими та незбалансованими наборами даних. Додаткове навчання на Kaggle-датасетах підвищило узагальнювальну здатність моделей і адаптивність до різних типів фішингових шаблонів. Рекомендаційна модель, побудована на змішаному наборі SMS, Email та користувацьких повідомлень, досягла 91% загальної точності, успішно розпізнаючи як фішингові, так і безпечні категорії.

Таким чином, реалізована інтелектуальна система ідентифікації фішингових повідомлень і URL показала високу ефективність і надійність у розпізнаванні різних кіберзагроз, а поєднання алгоритмів ML (NB, LR, SVM, RF, GB) із попередньою обробкою даних, модулем рекомендацій і веб-інтерфейсом забезпечило точну класифікацію, адаптивність до різних форматів контенту та зручну взаємодію з користувачем.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Saleem, H., Fuzail, M., Naeem, A., Aslam, N., Faiz, A. (2025). Model for the Detection of Web Phishing Attacks by using Machine Learning Algorithms. *Kashf Journal of Multidisciplinary Research*, 2(6), 223-239. URL: <https://doi.org/10.71146/kjmr501>.
2. Hassnain, M., Qureshi, I. A., Haider, A. (2025). Detection and Identification of Novel Attacks in Phishing using AI Algorithms. *International Journal of Computer Science and Mobile Computing*, 14(3), 20-27. URL: <https://doi.org/10.47760/ijcsmc.2025.v14i03.003>.
3. Swetha, C. V., Shaji, S. (2025). Applying Textual Fusion and Metadata in a Multimodal Transformer-Based Approach to Detect Phishing. *Third International Conference on Recent Trends in Computer Science and Data Analytics*. URL: https://www.researchgate.net/publication/390040540_Applying_Textual_Fusion_and_Metadata_in_a_Multimodal_Transformer-Based_Approach_to_Detect_Phishing.
4. Phishing dataset. *Hugging Face*. URL: <https://huggingface.co/datasets/ealvaradob/phishing-dataset>.
5. Phishing Email Dataset. *Kaggle*. URL: https://www.kaggle.com/datasets/naserabdullahalam/phishing-email-dataset?select=phishing_email.csv.

6. Phishing URL Classifier Dataset Cleaned. *Kaggle*. URL: <https://www.kaggle.com/datasets/adityachaudhary1306/phishing-url-classifier-dataset-cleaned>.

УДК 004.8

Чернова О. Ю., Антоненко О. С.
Одеський національний університет
І. І. Мечникова, м. Одеса, Україна

ОСОБЛИВОСТІ ПРЕДСТАВЛЕННЯ ЧАСОВИХ РЯДІВ ДЛЯ МОДЕЛІ МАШИННОГО НАВЧАННЯ В ПРОГНОЗУВАННІ ПОГОДИ

При побудові власної бюджетної IoT метеостанції є доречним створити модель для прогнозування певних метеорологічних параметрів погоди. Даний крок може включати наступне: збір даних з власної метеостанції протягом великого проміжку часу (для того, щоб навчена модель мала змогу більш точно та тривало спрогнозувати майбутні події) або тренування на існуючому наборі даних або об'єднання попередніх етапів (зазвичай це відбувається на початкових етапах, де спочатку використовуються дані із зовнішніх джерел, а далі з плином часу збираються власні дані з IoT метеостанції). При цьому невід'ємною частиною є фіксація дати та її подальше збереження, наприклад, в формат вигляду рік-місяць-день-час, який далі важливо правильно подати у якості вхідних даних. Певною проблемою є те, що такі дані є циклічними (повторюються через деякі проміжки часу) та важливо правильно представити їх природу.

Метою даного дослідження є експериментальна перевірка того, як вигляд вхідних даних впливає на точність побудованої моделі машинного навчання (зокрема, на k -найближчих сусідів для регресії). Зібрано дані за 5 років з міста Одеси за період з 1 квітня по 31 жовтня з безкоштовного джерела Метеопост [1] в файл формату csv. Дата об'єднана у вид 2021-04-01 00:00:00. Щоб наочно побачити та впевнитися, наскільки сильно на навчання обраної моделі впливає подача правильно конвертованих вхідних даних, а також самі точки даних, проведено експерименти, тому що подавати вхідні значення дати з прикладу "як є" взагалі неприпустимо. Дані є циклічними – місяці, дні та час. Зазвичай, популярним методом кодування таких значень є синусоїдальне та косинусоїдальне перетворення, які використовує автор дослідження [2], де зокрема порівнюється представлення вхідних

даних (часу) як порядкові значення (від 0 до 23) та як трансформовані за допомогою пари $\sin$ і $\cos$. Формули перетворення можна записати наступним чином [3]:

$$x_{\sin} = \sin\left(\frac{2\pi a}{\max(a)}\right), \quad x_{\cos} = \cos\left(\frac{2\pi a}{\max(a)}\right), \quad (1)$$

де $\max(a)$ – це максимально можливе значення a , для часу – 24, для дня з місяця вирішено взяти середнє значення 30.5, для дня з року – 365.

Проведено порівняння різних підходів представлення вхідних даних для алгоритму машинного навчання k -найближчих сусідів для регресії; цільові функції – температура повітря, атмосферний тиск (тиск моря) та вологість, варіанти дати:

- №1 – необроблені дані мають вигляд – рік, місяць, день і час (цілі значення: 0, 3, 6, ..., 21);

- №2 – дані трансформовано за допомогою синусоїдального та косинусоїдального перетворення: день (як день з місяця) та час;

- №3 – дані перетворено за допомогою синусоїдального та косинусоїдального перетворення: день (як день з року) та час;

- №4 – дані з урахуванням автокореляції (яку в даному випадку потрібно врахувати, адже перевірено, що всі параметри погоди мають високий коефіцієнт кореляції та графіки показують певну систематичну закономірність; саме наявні сезонні зміни, котрі відбуваються з певною періодичністю, можуть створювати автокореляцію [4]), тобто поєднання лагових ознак цільової функції [5] та конвертованих даних ($\sin$ і $\cos$ дня з року та часу).

Для вибору оптимальної кількості значень k написано функцію, яка із заданої кількості вибраних k (50) розраховує для кожного випадку середньоквадратичну похибку та з цього обирає оптимальне значення k . Дані розділено на 3 вибірки: навчальну, валідаційну та тестову. Результати тестової вибірки для температури повітря наведено в табл. 1, для атмосферного тиску (тиску моря) та вологості в табл. 2 та табл. 3 відповідно. Для кожного випадку в таблицях вказано кількість використаних оптимальних k , а також обчислено середньоквадратичну похибку (MSE) та коефіцієнт детермінації (R^2).

Таблиця 1 – Температура повітря

	№1	№2	№3	№4
k	12	16	50	8
MSE	21.553	38.548	9.749	2.269
R ²	0.439	-0.004	0.746	0.940

Таблиця 2 – Атмосферний тиск (тиск моря)

	№1	№2	№3	№4
k	34	49	50	18
MSE	19.169	19.664	19.344	0.311
R ²	0.025	0.000	0.017	0.984

Таблиця 3 – Вологість

	№1	№2	№3	№4
k	22	49	50	6
MSE	0.027	0.028	0.022	0.008
R ²	0.082	0.038	0.257	0.732

Отже, найбільш коректним та точним варіантом є №4, де в якості вхідних значень використовуються 2 минулі значення цільової функції (лагові ознаки) та перетворення $\sin$ та $\cos$ для дня з року та часу. Також варіант №3 непогано спрацював для прогнозування температури повітря, але в інших випадках показує погані результати. Щоб гарно навчити модель, краще мати дуже великий об'єм даних (також доречно їх проаналізувати на аномальні значення, які спричинені випадковістю чи програмними похибками) та правильно трансформувати вхідні значення для моделі машинного навчання.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Архів метеоданих. URL: <https://meteopost.com/weather/archive/> (дата звернення: 27.11.2025).
2. Mahajan T., Singh G., Bruns G. (2021, March 3). An experimental assessment of treatments for cyclical data [Conference session]. 2021 Computer Science Conference for CSU Undergraduates, Virtual.
3. Feature Engineering – Cyclical Variables. URL: <https://www.thedataschool.com.au/ryan-edwards/feature-engineering-cyclical-variables/> (дата звернення: 30.11.2025).
4. Пістунов І. М., Приходченко О. Ю. Економетрика. З розрахунками на Excel: навч. посіб. Дніпро: НТУ «ДП», 2024. 221 с. Режим доступу: <http://pistunovi.inf.ua/EkMe.pdf>

5. Lag Features. URL: <https://codesignal.com/learn/courses/preparing-financial-data-for-machine-learning/lessons/creating-lag-features-for-time-series-prediction> (дата звернення: 30.11.2025).

УДК 004.932:004.032.26

Шеремет К. О., Сіденко Є. В.

*Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ГЕНЕРАЦІЯ СИНТЕТИЧНИХ РЕНТГЕНІВСЬКИХ ЗНІМКІВ НА ОСНОВІ ОПИСІВ ІЗ ВИКОРИСТАННЯМ НЕЙРОННИХ МЕРЕЖ

Генерація синтетичних медичних зображень, зокрема рентгенівських знімків грудної клітки, є одним із ключових напрямів розвитку медичних інформаційних технологій. Зростання потреби в якісних і різноманітних даних, обмеження конфіденційності, нестача розмічених знімків і юридичні вимоги щодо персональних даних стимулюють застосування генеративних моделей для створення достовірних рентгенівських зображень, придатних для навчання та тестування медичних AI-систем.

Розробка системи генерації синтетичних рентгенівських зображень на основі текстових описів покликана підвищити доступність медичних даних, знизити ризики, пов'язані з використанням чутливої інформації, та створити безпечні умови для тренування діагностичних моделей. У роботі застосовано підхід до донавчання сучасної дифузійної архітектури Stable Diffusion, яка зарекомендувала себе високою ефективністю в задачах text-to-image та добре адаптується до медичних сценаріїв. Оскільки результат залежить від якості вихідних даних, одним із ключових етапів дослідження стало очищення датасету: усунення анонімізованих і некоректних фрагментів, нормалізація медичної термінології та узгодження текстових описів із відповідними зображеннями.

Огляд сучасних досліджень демонструє високу ефективність дифузійних моделей у генерації синтетичних медичних зображень. Автори Yuanfeng Ji *et al.* [1] показали, що латентні дифузійні моделі, зокрема Stable Diffusion, є високоефективними для генерації chest X-ray. Вони підкреслюють важливість великих текстово-візуальних датасетів та умовного керування генерацією для підвищення якості синтетики. Дослідники Tobias Weber *et al.* [2] запропонували каскадний дифузійний підхід для генерації chest X-ray високої роздільності.

Показано, що багаторівнева дифузія з текстовим керуванням суттєво підвищує реалістичність і стабільність синтетичних даних. В свою чергу, Abdullah al Nomaan Nafi *et al.* [3] довели, що дифузійні моделі дозволяють формувати статистично значущі синтетичні медичні датасети навіть у сферах з браком реальних знімків. Таким чином, сучасні наукові результати підтверджують доцільність вибору дифузійних моделей як базової технології для генерації синтетичних рентгенівських знімків.

Для навчання та підготовки даних було використано відкритий медичний набір Chest X-rays від Indiana University [4], доступний на платформі Kaggle, який містить рентгенівські знімки та два типи текстових описів – *findings* та *impression*. Цей датасет було обрано завдяки наявності структурованих клінічних звітів, достатньому обсягу матеріалу та можливості відновлення коректного текстово-візуального відповідника, що є критично важливим для формування якісного та репрезентативного навчального набору. Для реалізації самої системи генерації було використано офіційну реалізацію моделі Stable Diffusion XL [5] з платформи Hugging Face, яка забезпечує відтворюваність експериментів і створює стандартизоване середовище для подальшого донавчання та інтеграції моделі.

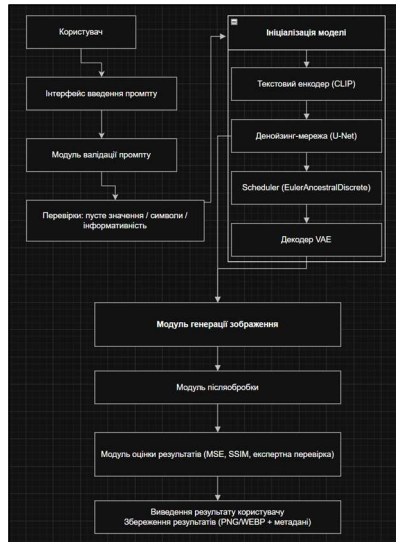


Рисунок 1 – Схема структури інтелектуальної системи генерації синтетичних рентгенівських знімків

Архітектура розробленої системи генерації синтетичних рентгенівських зображень на основі текстових описів побудована за модульним принципом, що забезпечує гнучкість, масштабованість і можливість подальшого удосконалення. Система включає веб-інтерфейс Streamlit для введення тексту та керування параметрами; модуль попередньої обробки тексту; модуль ініціалізації глибинної моделі Stable Diffusion XL; блок генерації зображень, що реалізує повний дифузійний цикл; модуль післяобробки та додавання метаданих; компонент оцінки якості результатів; а також підсистему збереження згенерованих зображень. Така архітектура забезпечує повний цикл обробки – від введення опису до отримання, оцінювання та збереження синтетичного рентгенівського зображення (див. рис. 1).

У процесі побудови та тестування системи було оцінено якість роботи моделі Stable Diffusion XL у поєднанні з підготовленими текстовими описами та спеціалізованими параметрами генерації. На основі тестового набору реальних chest X-ray знімків встановлено, що модель здатна стабільно відтворювати ключові анатомічні структури грудної клітки – легеневі поля, тінь серця та діафрагму. Під час експериментів досягнуто високої візуальної відповідності між синтетичними та реальними знімками та низького рівня артефактів за оптимальних параметрів дифузії. Використання модулів очистки та стандартизації текстових описів підвищило повторюваність результатів і зменшило кількість некоректних генерацій.

Таким чином, запропонована система синтетичної генерації медичних зображень на основі Stable Diffusion XL забезпечує високу якість відтворення рентгенівських структур, стабільність роботи й гнучкість налаштувань, що робить її ефективною основою для подальшого донавчання та тестування моделей у задачах медичної візуалізації.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Yuanfeng Ji et al. (2025). A Generative Foundation Model for Chest Radiography. *arXiv:2509.03903*. URL: <https://doi.org/10.48550/arXiv.2509.03903>.
2. Tobias Weber, Michael Ingrisch, Bernd Bischl, David Rügamer. (2023). Cascaded Latent Diffusion Models for High-Resolution Chest X-ray Synthesis. *Advances in Knowledge Discovery and Data Mining: 27th Pacific-Asia Conference on Knowledge Discovery and Data Mining, PAKDD, May 25–28, Proceedings, Part III, 180-191*. URL: <https://doi.org/10.48550/arXiv.2303.11224>.

3. Abdullah al Nomaan Nafi *et al.* (2024). Diffusion-Based Approaches in Medical Image Generation and Analysis. *arXiv:2412.16860*. URL: <https://doi.org/10.48550/arXiv.2412.16860>.
4. Chest X-rays (Indiana University). *Kaggle*. URL: <https://www.kaggle.com/datasets/raddar/chest-xrays-indiana-university>.
5. Stabilityai/stable-diffusion-xl-base-1.0. *Hugging Face*. URL: <https://huggingface.co/stabilityai/stable-diffusion-xl-base-1.0>.

ІНТЕЛЕКТУАЛЬНІ КОМП'ЮТЕРНО-ІНТЕГРОВАНІ ТЕХНОЛОГІЇ

УДК 004.42

Василенко В.І., Пенко В.Г.

*Одеський національний університет
імені І.І.Мечнікова, м. Одеса, Україна*

ПІДХІД ДО ЕФЕКТИВНОГО ВИКОРИСТАННЯ РЕСУРСІВ ЗА ДОПОМОГОЮ НАВЧАННЯ З ПІДКРІПЛЕННЯМ

Сучасні соціально-економічні та технічні системи функціонують в умовах постійного дефіциту й конкуренції за ресурси - матеріальні, енергетичні, інформаційні, фінансові та людські. Зростання масштабів задач логістики, виробництва, телекомунікацій, енергетики й сервісних систем вимагає прийняття рішень у реальному часі за умов невизначеності попиту, випадкових збоїв і змін зовнішнього середовища. Традиційні підходи надають цінний інструментарій, але їхня ефективність стрімко знижується при зростанні розмірності задачі, складності динаміки та вимог до адаптивності. У цих умовах перспективним є використання методів навчання з підкріпленням (Reinforcement Learning, RL), які розглядають керування ресурсами як послідовний процес прийняття рішень шляхом навчання на досвіді взаємодії агента із середовищем [1].

У роботі запропоновано технологічний підхід до розв'язання ресурсно-орієнтованих задач, що складається з таких етапів:

1. Формалізація задачі як марковського процесу прийняття рішень:

- стан описується вектором рівнів ресурсів, чергами заявок, параметрами попиту й зовнішнього середовища;
- простір дій включає рішення щодо замовлення, розподілу, пріоритетизації або обмеження використання ресурсів;
- функція винагороди моделює сукупні витрати, штрафи за дефіцит, простої та перевищення місткості.

2. Конструювання симуляційного середовища на основі інтерфейсу Gum [2], що реалізує цикл «спостереження – дія – винагорода - новий стан» і дозволяє варіювати сценарії попиту, структуру обмежень та політики початкового керування.

3. Вибір і налаштування алгоритмів RL. Для неперервних або складних просторів станів використано алгоритм Proximal Policy

Optimization (PPO). Для дискретних дій із великою розмірністю станів застосовано Deep Q-Network (DQN).

4. Інтеграція обмежень. Жорсткі технічні та ресурсні ліміти враховуються шляхом проєкції дій у допустиму множину, тоді як «м'які» вимоги (рівень сервісу, згладжування керувань) вбудовуються в функцію винагороди у вигляді штрафних термів.

5. Оцінювання якості політики. Для аналізу ефективності використовуються показники середньої сумарної винагороди, частоти дефіцитів, середньої довжини черг, завантаження ресурсів та стійкості результатів до змін параметрів середовища.

Запропонована технологія демонструє, що навчання з підкріпленням є перспективним інструментом для розв'язання задач раціонального використання ресурсів у динамічних стохастичних системах.

Подальші дослідження доцільно спрямувати на розроблення мультиагентних схем для великих децентралізованих систем, вивчення стійкості політик до змін режимів роботи та дослідження питань пояснюваності рішень RL-агентів у критично важливих застосуваннях.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Sutton R. S., Barto A. G. Reinforcement Learning: An Introduction. 2nd ed. Cambridge: MIT Press, 2018. 548 p.

2. Gymnasium Documentation: An API Standard for Reinforcement Learning Environments [Електронний ресурс]. Farama Foundation, 2024. URL: <https://gymnasium.farama.org> (дата доступу: 20.11.2025)

Кондрашов Д. А.,

*Криворізький національний університет
м. Кривий Ріг, Україна*

Музика І. О.,

*к. т. н, доцент, Криворізький національний університет
м. Кривий Ріг, Україна*

ІНТЕЛЕКТУАЛЬНЕ КЕРУВАННЯ ПРОЦЕСОМ ПРЕЦИЗІЙНОГО СОРТУВАННЯ РУД НА ОСНОВІ МАШИННОГО НАВЧАННЯ

Розглянуто актуальність застосування технологій машинного навчання для інтелектуального керування процесом прецизійного сортування руд. Проаналізовано існуючі системи сенсорного

сортування та виділено перспективні напрямки розвитку інтелектуальних систем керування в умовах нестационарності рудного потоку.

Сучасна гірничо-збагачувальна промисловість потребує підвищення ефективності переробки мінеральної сировини на всіх етапах технологічного процесу. Прецизійне (сенсорне) сортування руд є одним із ключових напрямків автоматизації збагачувального виробництва, що дозволяє здійснювати селективне відбракування некондиційних шматків руди ще на ранніх стадіях переробки. Це забезпечує зниження навантаження на подальші стадії збагачення, зменшення енергетичних витрат та підвищення якості кінцевого продукту.

Однак існуючі системи автоматичного сортування стикаються з низкою проблем, пов'язаних із нестационарністю рудного потоку, виникненням нестандартних ситуацій (накладання проєкцій декількох шматків руди, підсипка неідентифікованих різновидів сировини, зміна гранулометричного складу) та необхідністю адаптації до мінливих умов експлуатації. Традиційні методи керування, засновані на жорстких алгоритмах та порогових значеннях, не завжди здатні забезпечити належну точність класифікації в умовах високої варіативності вхідних параметрів.

Застосування технологій машинного навчання та нейронних мереж відкриває нові можливості для створення інтелектуальних систем керування, здатних автоматично розпізнавати складні закономірності у даних, адаптуватися до змінних умов та приймати рішення в режимі реального часу з високою точністю.

ОГЛЯД ІСНУЮЧИХ СИСТЕМ ТА ПІДХОДІВ

Сучасні системи сенсорного сортування руд базуються на різних фізичних принципах виявлення відмінностей між корисними мінералами та породою. Найбільш поширеними є такі типи сенсорів:

Оптичні системи використовують аналіз відбитого або пропущеного світла у різних діапазонах спектру (видимий, ближній інфрачервоний, рентгенівський). Компанії Tomra, Steinert, Mogensen пропонують комерційні рішення на основі RGB-камер, гіперспектральних сенсорів та рентген-флуоресцентного аналізу. Основні переваги полягають у високій швидкості обробки та можливості виявлення мінералогічних особливостей. Недоліком є залежність від чистоти поверхні матеріалу та умов освітлення.

Магнітні сенсори визначають магнітну сприйнятливість шматків руди шляхом вимірювання зміни індуктивності резонансного контуру. Перевагами таких систем є нечутливість до забруднення поверхні та можливість роботи у широкому діапазоні гранулометрії. Проте існують

обмеження щодо швидкодії (2-4 мс на вимірювання), взаємодії котушок та впливу форми і орієнтації об'єктів на результати вимірювання. Аналіз експериментальних даних показує, що для зразків з магнітною сприйнятливістю менше 1×10^{-3} SI спостерігається сильна кореляція відгуку сенсора з об'ємом та площею об'єкта, а не з його магнітними властивостями, що потребує додаткових методів компенсації.

Рентген-трансмсійні системи (XRT) аналізують ослаблення рентгеновського випромінювання при проходженні через матеріал, що дозволяє визначити щільність та атомний склад. Системи компаній XRT Sorter Technologies та інших широко застосовуються для сортування алмазовмісних кімберлітів, золотовмісних руд та кольорових металів.

Лазерне сканування (LIDAR) створює тривимірну модель поверхні шматка руди, що дозволяє оцінити його геометричні параметри, виявити тріщини та неоднорідності структури. Інтеграція 3D-даних з іншими сенсорними даними підвищує точність класифікації.

Аналіз літературних джерел показує, що більшість сучасних систем використовують традиційні статистичні методи класифікації (лінійний дискримінантний аналіз, метод опорних векторів) або прості порогові алгоритми. Застосування глибоких нейронних мереж (згорткових мереж для обробки зображень, рекурентних мереж для аналізу часових рядів) у задачах сортування руд поки що обмежене, хоча такі підходи демонструють значний потенціал у суміжних галузях автоматизації технологічних процесів.

ПЕРСПЕКТИВНІ НАПРЯМКИ ДОСЛІДЖЕННЯ

На основі аналізу стану питання можна виділити наступні перспективні напрямки розвитку інтелектуальних систем керування процесом сортування руд:

1. Розробка гібридних архітектур нейронних мереж для обробки мультимодальних даних (оптичні зображення, магнітні характеристики, 3D-геометрія). Доцільним є дослідження можливостей архітектур типу CNN для просторових даних, LSTM для послідовностей вимірювань та механізмів уваги (attention) для виділення найбільш інформативних ознак.

2. Методи виявлення аномалій та нестандартних ситуацій на основі автоенкодерів, генеративно-змагальних мереж (GAN) або методів novelty detection. Це дозволить системі автоматично ідентифікувати випадки накладання об'єктів, появи невідомих типів сировини або несправностей обладнання без попереднього навчання на таких ситуаціях.

3. Адаптивне навчання та transfer learning для швидкої адаптації моделей до зміни характеристик рудного потоку без необхідності

повного перенавчання системи. Актуальними є дослідження методів few-shot learning та meta-learning для роботи з обмеженими обсягами даних про нові типи руд.

4. Компенсація впливу дестабілізуючих факторів (температурний дрейф характеристик сенсорів, взаємодія котушок, вплив форми та орієнтації об'єктів, шум у динамічних режимах роботи конвеєра) шляхом розробки спеціалізованих моделей корекції на основі машинного навчання. Особливої уваги потребує проблема залежності відгуку магнітних сенсорів від геометричних параметрів об'єктів.

5. Оптимізація архітектури нейронних мереж з урахуванням обмежень реального часу. Дослідження методів квантизації, pruning та knowledge distillation для забезпечення роботи складних моделей на промислових обчислювальних платформах з обмеженими ресурсами при збереженні частоти опитування сенсорів 365 замірів на секунду.

6. Розробка критеріїв багатокритеріальної оптимізації для балансування між точністю класифікації, швидкістю системи, витратами на помилки першого та другого роду, енергоефективністю процесу.

ВИСНОВКИ

Створення інтелектуальних систем керування процесом прецизійного сортування руд на основі машинного навчання є актуальним науково-практичним завданням, що потребує комплексного підходу до розробки математичних моделей, методів обробки даних та алгоритмів прийняття рішень. Інтеграція сучасних технологій глибокого навчання з існуючими апаратно-програмними рішеннями у галузі сенсорного сортування може забезпечити значне підвищення ефективності збагачувального виробництва та точності відбраковки некондиційних шматків руди в умовах нестаціонарності рудного потоку.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Peukert D., Xu C., Dowd P. I. A Review of Sensor-Based Sorting in Mineral Processing: The Potential Benefits of Sensor Fusion // Minerals 2022, 12(11), 1364. URL: <https://doi.org/10.3390/min12111364> (дата звернення 26.11.2025).

2. Wang Q., Zhang K., Zou G. et al. Advances in deep learning-based ore particle size detection: A review of methods, challenges, and trends // MetaResource vol. 2, no. 2, pp. 83–104, 2025. DOI: 10.23919/METAR.2025.000006. URL: <https://ieeexplore.ieee.org/stamp/stamp.jsp?arnumber=11126021> (дата звернення 28.11.2025).

МАТЕМАТИЧНА МОДЕЛЬ РОЙОВОГО ІНТЕЛЕКТУ ДЛЯ КАРТОГРАФУВАННЯ ТЕРИТОРІЇ ЗА ДОПОМОГОЮ БПЛА

На сучасному етапі технологічного розвитку мультиагентні системи (МАС) набувають широкого застосування в багатьох наукових і технічних сферах, стаючи визначальним чинником прогресу.

Особливої важливості ці системи набувають в умовах збройних конфліктів, де гостро відчувається потреба в автономних системах із високорозвиненим інтелектуальним управлінням. Це підкреслює актуальність подальшого розвитку мультиагентних технологій[1].

Вже сьогодні ці технології активно використовуються в безпілотних літальних апаратах (БПЛА). Їхнє застосування охоплює не лише військову справу, але й такі галузі, як агропромисловий комплекс, пошуково-рятувальні операції, будівництво та екологічний моніторинг.

Для повного розуміння цього матеріалу радимо переглянути термінологію, описану в попередніх публікаціях авторів [2, 3].

Формалізовані моделі поведінки рою беруть свій початок у біологічних дослідженнях колективної поведінки тварин 19–20 століть. На початку 20-го століття з'явилися перші прості математичні моделі, як-от відома динаміка популяцій Лотки-Вольтерри.

Справжній активний розвиток обчислювальних моделей розпочався у другій половині 20-го століття, коли поява комп'ютерів уможливила симуляцію складних та великомасштабних систем. Сучасні моделі рою є результатом інтеграції спостережень за природними явищами, математичного аналізу та нових технологій.

Існують два принципово різні підходи для моделювання й аналізу динаміки рою: просторовий (spatial) та непросторовий (nonspatial) [4].

У непросторових підходах динаміка рою описується на рівні популяції через частотні розподіли груп різних розмірів, при цьому передбачається, що групи об'єднуються або розділяються залежно від внутрішньої динаміки, умов навколишнього середовища та взаємодії з іншими групами. Недоліком цього підходу є необхідність введення кількох «штучних» припущень щодо злиття та поділу груп, що негативно впливає на точність моделювання.

Просторовий підхід, навпаки, явно чи неявно враховує середовище і поділяється на дві окремі концепції: індивідуально-орієнтовану або лагранжеву (Lagrangian) та континуальну або ейлерову (Eulerian) [4]. У науковій літературі ці концепції можуть мати інші назви:

індивідуально-орієнтовані моделі іноді називають мікроскопічними, а континуальні — макроскопічними [5].

На основі детального аналізу наукових джерел [3-9], які стосуються мультиагентних систем, що використовуються, зокрема, в безпілотних літальних апаратах (БПЛА), для подальшого дослідження була обрана просторова Лагранжева модель, запропонована Т.І. Zohdi [10].

Ця модель є переважною завдяки її відносній простоті та водночас значному дослідницькому потенціалу, оскільки вона забезпечує точне відстеження траєкторії кожного агента (БПЛА) у тривимірному просторі, надає ефективний інструментарій для моделювання складних взаємозв'язків між агентами та базується на класичних фізичних принципах руху, що гарантує надійну та зрозумілу математичну основу.

Мета дослідження полягає у моделюванні процесу колективного картографування території за допомогою рою автономних агентів (дронів), як показано на Рисунку 1 . При цьому головним критерієм оптимізації є мінімізація сумарного часу, необхідного для обстеження всієї території.

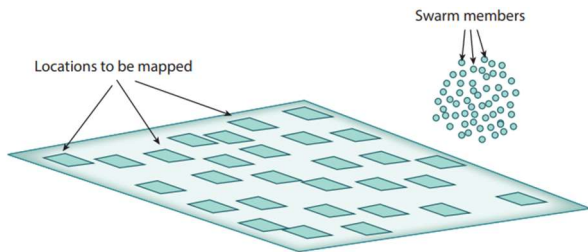


Рисунок 1 - Типова схема картографування території за допомогою рою
(Locations to be mapped — Локації, цілі для картографування;
Swarm members — Члени, агенти рою). [10]

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Military Drone Market Size, Share & Russia-Ukraine War Impact Analysis. Forecast, 2024-2032
URL:<https://www.fortunebusinessinsights.com/military-drone-market-102181> (дата звернення 23.12.2024)
2. Купін А.І., Косей М.П. Аналіз алгоритмів ройового інтелекту. «Системні технології» - Том 3 №152 (2024) – С. 69-80
3. Купін А.І., Косей М.П. Аналіз алгоритмів ройового інтелекту. International scientific and technical conference. Information Technologies in Metallurgy and Machine building – ITMM 2024 – С. 460-465

4. Gazi, Veysel & Passino, Kevin. (2003). Stability analysis of swarms. Automatic Control, IEEE Transactions on. 48. 692 - 697. 10.1109/TAC.2003.809765.
5. Lerman, Kristina & Galstyan, Aram & Hogg, Tad. (2004). Mathematical Analysis of Multi-Agent Systems.
6. Esteban Restrepo. Coordination control of autonomous robotic multi-agent systems under constraints. Automatic. Université Paris-Saclay, 2021. English. <NNT : 2021UPASG093>. <tel-03537341>
7. Fina, L.; Smith, D.S., Jr.; Carnahan, J.; Sevil, H.E. Entropy-Based Distributed Behavior Modeling for Multi-Agent UAVs. Drones 2022, 6, 164. <https://doi.org/10.3390/drones6070164>
8. Tchappi, Igor & NDAMLABIN MBOULA, Jean Etienne & Najjar, Amro & Mualla, Yazan & Galland, Stéphane. (2022). A Decentralized Multilevel Agent Based Explainable Model for Fleet Management of Remote Drones.
9. Martin Rosalie, Grégoire Danoy, Serge Chaumette, Pascal Bouvry, Chaos-enhanced mobility models for multilevel swarms of UAVs, Swarm and Evolutionary Computation, Volume 41, 2018, Pages 36-48, ISSN 2210-6502, <https://doi.org/10.1016/j.swevo.2018.01.002>.
10. Zohdi, T. I. (2018). Multiple UAVs for Mapping: A Review of Basic Modeling, Simulation, and Applications. Annual Review of Environment and Resources, 43(1), 523–543. DOI: 10.1146/annurev-environ-102017-025912.

УДК 004.8:004.056.53:004.9

Моторнюк С. О.

*Харківський національний економічний університет
ім. Семена Кузнеця, м. Харків, Україна*

ПРОБЛЕМАТИКА ДОСТОВІРНОСТІ ТА ПОВНОТИ ВІДКРИТИХ ДАНИХ: ВИКЛИКИ ТА НАПРЯМИ ІНТЕЛЕКТУАЛЬНОЇ ОБРОБКИ

Проблема достовірності та повноти відкритих даних є критичною для цифрового врядування та наукових досліджень, адже від якості оприлюднених наборів залежить точність аналітики та обґрунтованість управлінських рішень. Попри значну кількість робіт, досі бракує узгоджених підходів до вимірювання коректності, повноти й актуальності даних [1]. Особливої складності додає поєднання внутрішньої якості (точність, цілісність, актуальність) та зовнішньої

(придатність для конкретних завдань і повторного використання) [2]. Зростання обсягів і різноманітності джерел ускладнює автоматизований контроль, що знижує довіру до відкритих даних. Перспективним є застосування інтелектуальних методів, зокрема багатовимірного аудиту та поєднання великих мовних моделей із знаннями графами для автоматизованої перевірки й збагачення метаданих [3].

Останні публікації засвідчують інтенсивний рух від розпорошених трактувань «якості даних» до її системної операціоналізації для відкритих екосистем, однак консенсусу щодо складу вимірів та їхнього зв'язку з прикладними сценаріями досі бракує. Узагальнювальна праця Carvalho та колег [1] показує масштаб проблеми: попри десятиліття досліджень, відсутня узгоджена структура вимірів, яку можна безпосередньо покласти в основу прикладних фреймворків і міждоменної порівнюваності; автори ідентифікують широкий спектр характеристик і пропонують їх переорганізацію, але не знімають питання відповідності цих характеристик специфіці відкритих даних та їхніх метаданих. Для власне відкритих даних потужною виявилася запропонована автором González-Vidal дихотомія «внутрішньої» та «зовнішньої» якості: перша фокусується на репрезентативності, повноті, узгодженості й своєчасності, друга — на придатності наборів до конкретних завдань аналітики й навчання моделей; запропоновані метрики і прототипи автоматизації дають практичні інструменти для портальних репозитаріїв, однак лишається відкритим питання масштабування цих перевірок, їхньої доменної адаптації та інтеграції з життєвим циклом публікації даних [2]. На рівні інструментарію у деяких наукових дослідженнях помітний тренд до уніфікації багатомодального аудиту: нові бібліотеки поєднують показники для табличних, часових і візуальних даних, вводять таксономії «вхідної» та «синтетичної» валідації, а також інстанс-рівневі статистики для локалізації «пробов» якості, як це показано у дослідженні Façoso та співавторів; проте їх застосування у відкритих даних стикається з бідністю метаданих, різноманітністю ліцензій/форматів та відсутністю узвичаєних протоколів вбудовування перевірок у платформи публікації [3]. Паралельно науковцями формується напрям інтелектуальної обробки для підвищення змістовної повноти, як це показано у праці Benjira та співавторів: поєднання великих мовних моделей із знаннями графами дає змогу автоматизувати семантичне зіставлення відкритих джерел із формулами показників сталого розвитку, істотно покращуючи точність і повноту мепінгу; водночас залишаються невирішеними питання стабільності таких конвеєрів у багатомовних

середовищах, контроль помилкових відповідей і відтворюваність при оновленні даних та схем [4].

Отже, нерозв'язаними частинами загальної проблеми лишаються: відсутність доменно-агностичної, але операційно застосовної узгодженої системи вимірів для відкритих даних; розрив між метриками внутрішньої якості та показниками придатності до конкретних політик/моделей; нестача стандартних, масштабованих механізмів інтеграції перевірок у портальні й harvesting-процеси; обмежена мультимодальна підтримка та слабка «спостережуваність» якості на рівні окремих екземплярів; а також потреба в надійних LLM+KG-методах семантичного збагачення, що гарантовано зберігають відтворюваність і підтримують динамічну актуалізацію наборів [5].

Метою дослідження є підвищення достовірності та повноти відкритих даних шляхом формалізації вимірів якості, автоматизованого аудиту та інтелектуальної обробки. Для досягнення мети передбачено: уточнити ключові показники внутрішньої й зовнішньої якості; розробити методи автоматизованої перевірки та виявлення аномалій; створити підхід до семантичного збагачення й зіставлення даних; інтегрувати ці механізми в процеси публікації та оновлення наборів; оцінити ефективність запропонованих рішень на прикладі реальних відкритих даних.

Основний матеріал дослідження зосереджений на поєднанні системного аналізу вимірів якості відкритих даних та розробленні інтелектуальних механізмів їх автоматизованого контролю й збагачення. На першому етапі проведено формалізацію категорій «внутрішньої» та «зовнішньої» якості даних, що охоплюють точність, цілісність, актуальність, а також придатність для повторного використання та аналітики. Такий підхід дозволив структурувати чинники, що впливають на достовірність і повноту, та визначити критичні точки, де потрібні автоматизовані перевірки.

Далі запропоновано методика мультимодального аудиту, здатну оцінювати якість табличних, часових, текстових та візуальних даних. Методика поєднує статистичні та семантичні перевірки, виявлення пропусків та аномалій, а також перевірку зв'язності між об'єктами й метаданими. Розроблена модель інтегрується у процеси публікації та агрегації даних, забезпечуючи безперервний моніторинг якості та зворотний зв'язок авторам наборів.

Окрему увагу приділено використанню великих мовних моделей (LLM) у поєднанні зі знанневими графами для семантичного збагачення й зіставлення даних з онтологіями та індикаторами публічної політики.

Такий підхід дає змогу автоматизувати валідацію та підвищити повноту метаданих навіть для гетерогенних і багатомовних джерел.

Для перевірки результатів проведено експериментальні дослідження на реальних відкритих наборах даних. Порівняно базовий підхід (ручна перевірка та мінімальні автоматизовані тести) з розробленою інтелектуальною системою. Результати показали зростання показників якості за ключовими параметрами — насамперед за повнотою та достовірністю — а також зниження витрат часу на аудит (табл. 1).

Таблиця 1. Порівняння базового та запропонованого підходів до перевірки якості відкритих даних

Підхід	Повнота (%)	Достовірність (%)	Час перевірки (відн.)
Базовий (ручні перевірки)	68	72	1.0
Запропонований (інтелектуальний аудит)	89	91	0.35

Проведене дослідження підтвердило, що якість відкритих даних визначається не лише технічними аспектами їх публікації, а насамперед чіткою формалізацією вимірів достовірності та повноти, а також можливістю автоматизованого контролю на всіх етапах життєвого циклу даних. Запропоновано узгоджену систему показників внутрішньої й зовнішньої якості, методику мультимодального аудиту, здатну виявляти пропуски, аномалії та семантичні невідповідності, а також інтегрований підхід до використання великих мовних моделей та знанневих графів для збагачення метаданих і підвищення корисності наборів для повторного використання. Експериментальна перевірка засвідчила істотне зростання показників повноти та достовірності за одночасного скорочення часу на аудит, що підтверджує практичну ефективність запропонованих рішень.

Подальші дослідження доцільно зосередити на удосконаленні механізмів адаптації аудиту до різних доменів і форматів даних, розробленні масштабованих стандартів інтеграції перевірок у портальні платформи, а також на поглибленні підходів до семантичного зіставлення із залученням LLM у багатомовних середовищах. Перспективним є створення відкритих тестових корпусів і бенчмарків для порівняння методів оцінки якості, а також розроблення

інтерактивних інструментів, що дозволять розширити застосування запропонованих підходів у публічному управлінні та наукових дослідженнях.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Data Quality: revisiting dimensions towards new framework development [Електронний ресурс] / André M. Carvalho [та ін.] // Procedia Computer Science. – 2025. – Т. 253. – С. 247–256. – Режим доступу: <https://doi.org/10.1016/j.procs.2025.01.088> (дата звернення: 01.10.2025). – Назва з екрана.

2. González-Vidal A. Intrinsic and extrinsic quality of data for open data repositories [Електронний ресурс] / Aurora González-Vidal, Alfonso P. Ramallo-González, Antonio F. Skarmeta // ICT Express. – 2022. – Режим доступу: <https://doi.org/10.1016/j.ict.2022.06.001> (дата звернення: 01.10.2025). – Назва з екрана.

3. pyMDMA: Multimodal data metrics for auditing real and synthetic datasets [Електронний ресурс] / Ivo S. Façoco [та ін.] // SoftwareX. – 2025. – Т. 31. – С. 102256. – Режим доступу: <https://doi.org/10.1016/j.softx.2025.102256> (дата звернення: 01.10.2025). – Назва з екрана.

4. Automated mapping between SDG indicators and open data: An LLM-augmented knowledge graph approach [Електронний ресурс] / Wissal Benjira [та ін.] // Data & Knowledge Engineering. – 2025. – С. 102405. – Режим доступу: <https://doi.org/10.1016/j.datak.2024.102405> (дата звернення: 01.10.2025). – Назва з екрана.

5. CKAN code architecture – CKAN 2.9.11 documentation [Електронний ресурс] // Overview – CKAN 2.11.3 documentation. – Режим доступу: <https://docs.ckan.org/en/2.9/contributing/architecture.html>. – Назва з екрана.

КОНЦЕПЦІЯ ЦИФРОВОГО ДВІЙНИКА ДЛЯ СИСТЕМ ТЕХНІЧНОГО ОБСЛУГОВУВАННЯ

Останні декілька десятиліть супроводжувались стрімким розвитком телекомунікаційних технологій, засобів збору та обробки даних, а також проривними досягненнями у сфері штучного інтелекту та машинного навчання. Цей технологічний прогрес став каталізатором для четвертої промислової революції (Industry 4.0), яка полягає в масовому впровадженні кіберфізичних систем у виробництво та створенні «розумних» підприємств, де фізичні та цифрові процеси глибоко інтегровані. У цьому новому технологічному ландшафті, поряд із вже звичними інструментами – Інтернетом речей (IoT), хмарними обчисленнями, доповненою реальністю та іншими засобами та технологіями – особливої ваги набуває концепція цифрового двійника (Digital Twin, DT).

Ідея цифрового двійника полягає у тому, щоб створити цифрове відображення, яке було б еквівалентне деякому реальному фізичному об'єкту або процесу, та поєднати це відображення з цим об'єктом так, щоб вони могли обмінюватись даними в режимі реального часу [1].

Ключовою відмінністю цифрового двійника від звичайної імітаційної комп'ютерної моделі є постійний двонапрямлений обмін даними між віртуальною та реальною системами, що дозволяє динамічно оновлювати цифрову модель так, щоб вона в точності відповідала поточному стану реального об'єкта. Окрім даних про фізичний об'єкт-двійник, віртуальна модель також отримує дані і про навколишнє середовище, завдяки чому підвищується точність з якою цифровий двійник відображає реальну фізичну систему [2].

Так як цифровий двійник з високою точністю відповідає реальному об'єкту, то його можна застосовувати для задач прогнозування, аналізу та оптимізації, не змінюючи при цьому об'єкт-оригінал. Станом на сьогодні, цифрові двійники застосовуються у авіа-космічній, медичній галузях, у військовій справі, в машинобудуванні, урбаністиці, робототехніці, енергетиці [3].

В даній роботі концепція цифрових двійників розглядається в контексті систем технічного обслуговування. Інтеграція технології цифрових двійників у процеси організації технічного обслуговування

дозволяє трансформувати традиційні підходи до підтримання працездатності обладнання та здійснити перехід до парадигми прогностичного обслуговування (Predictive Maintenance).

Прогностичне технічне обслуговування полягає у передбаченні залишкового ресурсу об'єкта (Remaining useful life, RUL) на основі даних, що характеризують цей об'єкт. Класичним варіантом реалізації моделі прогностичного технічного обслуговування є використання машинного навчання на основі історичних даних про відмови обладнання та його параметрів. Важливою перевагою застосування цифрових двійників у цьому контексті є здатність працювати в умовах змінного навантаження та нестаціонарних робочих режимів. Якщо традиційні методи машинного навчання часто втрачають точність, коли обладнання працює в умовах, відмінних від тих, на яких навчалася модель, то цифровий двійник, завдяки інтеграції фізичних законів та поточної телеметрії, здатен адаптуватися до нових умов експлуатації та надавати коректні прогнози RUL навіть у непередбачуваних ситуаціях [4]. На рисунку 1 показано як може виглядати структура системи технічного обслуговування з використанням цифрового двійника.



Рисунок 1 – Структура системи PdM з цифровим двійником

Таким чином, застосування цифрових двійників у системах технічного обслуговування забезпечує суттєво більш надійне, точне та достовірне прогнозування залишкового ресурсу (RUL) та потенційних відмов, що є особливо критичним для складних систем із динамічними, змінними та нестаціонарними умовами експлуатації (авіаційні та газотурбінні двигуни, вітрові турбіни, кар'єрні самоскиди й

екскаватори, прокатні стани, доменні печі, морські судна тощо). На відміну від традиційних моделей машинного навчання, які втрачають точність при відхиленні від навчальної вибірки, цифровий двійник постійно синхронізується з реальним об'єктом через потік телеметрії, враховує актуальні навантаження, режими роботи, деградацію матеріалів та зовнішні впливи.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Review of digital twin about concepts, technologies, and industrial applications / M. Liu та ін. *Journal of Manufacturing Systems*. 2020. URL: <https://doi.org/10.1016/j.jmsy.2020.06.017>
2. Barricelli B. R., Casiraghi E., Fogli D. A Survey on Digital Twin: Definitions, Characteristics, Applications, and Design Implications. *IEEE Access*. 2019. Vol. 7. P. 167653–167671. URL: <https://doi.org/10.1109/access.2019.2953499>
3. Digital Twin: Data Exploration, Architecture, Implementation and Future / M. S. Dihan et al. *Heliyon*. 2024. P. e26503. URL: <https://doi.org/10.1016/j.heliyon.2024.e26503>
4. Overview of predictive maintenance based on digital twin technology / D. Zhong et al. *Heliyon*. 2023. P. e14534. URL: <https://doi.org/10.1016/j.heliyon.2023.e14534>

УДК 004.738.5:004.8

*Світнєв О. Ю., Волощук Л. О.
Одеський національний університет
ім. І.І.Мечникова, м. Одеса, Україна*

ПОРІВНЯЛЬНИЙ АНАЛІЗ МЕТАЕВРИСТИЧНИХ АЛГОРИТМІВ ДЛЯ ЗАДАЧІ РОЗМІЩЕННЯ СЕРВІСІВ У ТУМАННИХ ОБЧИСЛЮВАЛЬНИХ СИСТЕМАХ

Туманні обчислення розміщують ресурси ближче до джерел даних, але виникає складна задача: де саме розмістити кожен сервіс для оптимальної роботи системи. Уявіть розумне місто з тисячами IoT пристроїв та десятками fog вузлів різної потужності. Треба вирішити, який сервіс на якому вузлі запускати. Неправильне рішення може збільшити затримку в 2-3 рази [1]. Ця задача належить до класу NP-складних - для системи навіть середнього розміру неможливо перебрати всі варіанти за прийнятний час. Тому використовують метаевристичні алгоритми, що знаходять досить гарне рішення за розумний час. У цій

роботі порівнюються чотири основних підходи: генетичний алгоритм, алгоритм рою частинок, мурашині колонії та Firefly Algorithm.

Генетичний алгоритм працює як еволюція - створює популяцію варіантів розміщення, найкращі схрещує між собою, додає мутації. Через багато поколінь знаходить якісне рішення. Chen та колеги отримали покращення затримки на 15%, а багатоцільові версії показують до 25% кращого використання ресурсів [2]. Основна перевага - універсальність застосування до різних задач. Недоліки: багато обчислень, складне налаштування параметрів, потреба в модифікаціях для динамічних систем.

Алгоритм рою частинок імітує поведінку зграї птахів. Кожна частинка рухається в просторі варіантів на основі власного та колективного досвіду. Показує покращення на 10-15% для статичних конфігурацій [3]. Головна перевага - простота реалізації, всього кілька параметрів, швидка конвергенція. Критичний недолік - погана адаптація до змін навантаження, схильність до передчасної конвергенції в локальні оптимуми.

Алгоритм мурашиних колоній моделює пошук шляху через феромонові сліди. Мурахи залишають феромон на шляху, інші йдуть туди, де його більше. Феромон випаровується, дозволяючи забувати погані рішення. Найефективніший для оптимізації маршрутизації в fog мережах - покращення затримки на 18% [4]. Перевага - природний баланс між використанням знайдених рішень та пошуком нових. Недоліки: повільна конвергенція, багато параметрів, гірше працює для задач без графової структури.

Firefly Algorithm базується на притяганні світлячків до яскравішого світла. Яскравість відповідає якості рішення. Алгоритм автоматично ділить рій на підгрупи навколо різних оптимумів, уникаючи передчасної конвергенції. Особливо ефективний для енергоефективності - зменшення споживання на 22%, найкращий показник серед метаевристик [5]. Недоліки: багато параметрів, висока обчислювальна складність через порівняння всіх пар світлячків.

Порівнюючи чотири методи, виділяємо ключові моменти. За швидкістю лідує алгоритм рою частинок – він найшвидший, але ризикує застрягти в локальному оптимумі. Генетичний алгоритм (GA) повільніший, але надійніше досліджує простір. Мурашині колонії (ACO) найповільніші, проте добре балансують пошук та використання рішень. Firefly займає середню позицію і добре уникає передчасної конвергенції. За якістю рішення всі методи схожі після достатньої кількості ітерацій (різниця 5-10%). Критичніша різниця у стабільності: GA та Firefly дають стабільніші результати, PSO більш варіативний. За

енергоефективністю виділяється Firefly з 22% покращення, інші мають 10-15%.

Складність реалізації найнижча у PSO (простий код, мало параметрів). GA середньої складності (треба реалізувати оператори схрещування/мутації). АСО складніший через моделювання феромонових слідів та графової структури. Firefly також вимагає ретельної реалізації механізму притягання. Обчислювальна складність найвища у Firefly та GA через порівняння всіх пар/оцінку всієї популяції. PSO та АСО мають лінійну складність відносно розміру популяції, що важливо для великих систем.

На мою думку, вибір методу залежить від специфіки завдання. Для швидкого прототипування та систем з обмеженими ресурсами найкраще підходить алгоритм рою частинок - простий, швидкий, ідеальний для первинного налаштування fog системи. Якщо критична якість рішень і є час на налаштування, генетичний алгоритм буде кращим вибором. Особливо для складних багатокритеріальних задач, де треба балансувати затримку, енергію, вартість, надійність. Для систем з передбачуваним навантаженням та можливістю періодичної переоптимізації - це оптимальний варіант. Мурашині колонії мають вузьку, але важливу нішу - задачі з явною мережевою структурою. Для оптимізації маршрутизації в fog мережах, балансування навантаження між вузлами через динамічну зміну шляхів передачі даних - тут вони працюють краще за конкурентів. Природний механізм адаптації через випаровування феромону дає певну динамічність. Firefly Algorithm вартий уваги, коли енергоефективність - пріоритет номер один. Для автономних fog вузлів на батареях, edge пристроїв з обмеженим живленням, зелених дата-центрів - тут він показує найкращі результати. Також він добре працює для багатомодальних задач, де є кілька гарних областей рішень і треба їх всі знайти.

Спільне обмеження всіх цих методів — статичність. Вони оптимізують розміщення для конкретного моменту часу/навантаження. Коли ситуація змінюється, потрібна нова оптимізація. Для динамічних Fog-систем ці методи дають лише базову архітектуру, яку треба доповнювати механізмами адаптації в реальному часі.

Метаевристичні алгоритми залишаються основним практичним інструментом для розміщення сервісів у туманних системах. Кожен з розглянутих методів має свої переваги: PSO – швидкість та простота, GA – універсальність та якість, АСО – ефективність для мережесхемних задач, Firefly – енергоефективність. У реальних проектах часто ефективний гібридний підхід - використання різних алгоритмів на різних етапах або комбінування їх переваг. Проте в багатьох туманних

розподілених системах функціональне навантаження динамічно змінюється залежно від ситуації - час доби, день тижня, непередбачувані події. Усі розглянуті методи оптимізують розміщення для конкретного моменту часу, а при зміні умов потребують повторного запуску. Подальші дослідження мають зосередитися на розробці адаптивних варіантів цих алгоритмів, здатних динамічно підтримувати оптимальність розміщення без повної переоптимізації системи.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Ata A., Khan M. A., Abbas S., Ahmad G., Fatima A. Modelling smart road traffic congestion control system using machine learning techniques. Neural Network World. 2019. Vol. 29, no. 2. P. 99-110.

2. Chen L., et al. IoT microservice deployment in edge-cloud hybrid environment using reinforcement learning. IEEE Internet of Things Journal. 2021. Vol. 8, no. 16. P. 12610-12625.

3. Azimi S., Pahl C., Hosseini S. M. Performance management in clustered edge architectures using particle swarm optimization. Cluster Computing. 2021. Vol. 24. P. 2233-2257.

4. Guerrero C., Lera I. Multi-objective application placement in fog computing using graph neural network-based reinforcement learning. The Journal of Supercomputing. 2024. Vol. 80, no. 19. P. 27073-27094.

5. Pallewatta S., Kostakos V., Buyya R. QoS-aware placement of microservices-based IoT applications in Fog computing environments. Future Generation Computer Systems. 2022. Vol. 131. P. 121-136

УДК 681.5

Харламенко В.Ю., Тимохін Є.В.

*Криворізький національний університет
м. Кривий Ріг, Україна*

ОГЛЯД ПЕРСПЕКТИВНИХ МЕТОДІВ ПРОГНОЗНОГО КЕРУВАННЯ ВОДНИМИ РЕСУРСАМИ ЗБАГАЧУВАЛЬНОЇ ФАБРИКИ

Рациональне використання водних ресурсів у процесах збагачення корисних копалин є одним із ключових чинників підвищення енергоефективності, стабільності виробництва та екологічної безпеки. Сучасні фабрики зі збагачення залізородної сировини характеризуються складною гідравлічною інфраструктурою, значними коливаннями витрат води та високою залежністю від зовнішніх

факторів, включаючи змінні технологічні навантаження, погодні умови та якість руди. У таких умовах впровадження методів машинного навчання та систем штучного інтелекту стає стратегічним інструментом для оптимізації управління водою, підвищення точності прогнозування та зниження невизначеності гідравлічних моделей.

1. Моделювання та прогнозування водного споживання як основа для прогнозного управління.

Оперативна оцінка та прогнозування водоспоживання на окремих технологічних вузлах збагачувальної фабрики є визначальним елементом побудови систем прогнозного керування. Вода використовується на всіх стадіях збагачення – під час дроблення та подрібнення руди, у гідротранспорті, класифікації та флотації, а тому точне визначення миттєвого та погодинного попиту дозволяє стабілізувати роботу насосних станцій, зменшити гідравлічні втрати та підвищити ступінь повторного використання води.

Відповідно до сучасних підходів, для прогнозування погодинного водоспоживання можуть застосовуватися моделі на основі алгоритмів типу XGBoost, Random forest та штучних нейронних мереж, які інтегрують історичні дані, зовнішні фактори, часові патерни та технологічні індикатори. Комбінування багатовимірних вхідних змінних забезпечує значно вищу точність моделі порівняно з традиційними методами часового ряду, що було підтверджено у дослідженні, де XGBoost продемонстрував найвищий показник точності серед трьох протестованих алгоритмів [1]. Такий підхід може бути безпосередньо адаптований для потреб збагачувальних фабрик, де важливо враховувати як внутрішні технологічні цикли, так і зовнішні впливи, наприклад зміни у графіку подачі руди або особливості водного режиму хвостосховищ.

2. Підвищення надійності гідравлічних моделей через зменшення невизначеності.

Точність роботи будь-якої системи керування, що базується на моделюванні потоків води, залежить від рівня невизначеності вхідних даних. Невизначеність може виникати як через шум або аномалії в даних сенсорів, так і через структурні обмеження гідравлічної моделі. У гірничо-збагачувальних підприємствах ці фактори проявляються особливо помітно: датчики витрати та тиску працюють у складних умовах, транспортні траси операційно змінюються, а підземні та відкриті гідравлічні системи піддаються впливу середовища.

У дослідженні [1] запропонована комплексна структура роботи з невизначеністю, що включає попереднє очищення даних (кластеризація DBSCAN та метод ADR-CI), поділ вузлів на групи з подібними

характеристиками споживання та застосування методів асиміляції даних на основі фільтра Калмана. Такий підхід дозволяє значно зменшити як вимірювальну, так і модельну невизначеність та забезпечити стабільну роботу гідравлічної моделі в реальному часі. Для збагачувальних фабрик це означає можливість оперативного прогнозування тисків у трубопроводах, раннє виявлення нештатних ситуацій (наприклад, засмічення або надлишкового водовідбору) та підтримку оптимального водного балансу.

3. Прогнозне керування як засіб оптимізації водорозподілу

Прогнозне керування, кероване моделями машинного навчання, дозволяє будувати динамічні стратегії, здатні випереджати зміни технологічних навантажень і відповідно адаптувати режими роботи насосних станцій, аераційних систем або установок очищення.

У практичному застосуванні це підтверджено результатами впровадження самонавчальної feedforward-системи для очищення стічних вод, де точність прогнозування навантаження становила близько 88%, а споживання енергії на аерацію було знижено на $\approx 15\%$ у порівнянні з традиційним керуванням [2]. Такий ефект демонструє потенційні переваги й для збагачувальних фабрик, де водні системи значною мірою визначають енергоємність процесу. Використання прогнозного керування в операціях рециркуляції, стабілізації пульпи чи подачі технологічної води сприятиме зниженню витрат та підвищенню ефективності.

4. Інтелектуальна аналітика для операційного контролю

Окрім прогнозування та керування, штучний інтелект дозволяє автоматично виявляти аномалії у гідросистемах фабрики, що часто є маркерами прихованих дефектів, деградації обладнання або збоїв вимірювальної апаратури. Для цього можуть застосовуватись нейронні мережі на основі регресії квантилів, які формують діапазони очікуваних значень параметрів та сигналізують про відхилення від нормальної поведінки. У дослідженні [2] показано ефективність подібних моделей у виявленні аномалій процесу й інструментальних збоїв ще до того, як їх вплив стає критичним для роботи установки.

Для фабрик зі збагачення руди такі можливості мають особливе значення, оскільки гідромеханічні системи працюють у режимі постійних навантажень та зносу. Раннє виявлення відхилень у роботі насосів, деградації витратомірів або нерівномірного навантаження на гідроциклонні батареї дозволяє запобігати аваріям і підвищувати надійність всього виробничого циклу.

5. Перспективи впровадження на збагачувальних фабриках

Запропоновані підходи створюють передумови для формування інтегрованої інтелектуальної системи управління водними ресурсами на гірничо-збагачувальному підприємстві. Поєднання прогнозування, асиміляції даних, прогнозного контролю і автоматичного виявлення аномалій дозволяє забезпечити стабільну роботу гідросистем і насосного парку, підвищити рівень рециркуляції води і зменшити потребу у свіжій воді; одночасно створюються передумови для зменшення ризиків аварій, пов'язаних з гідравлічними перевантаженнями; стає можливим забезпечити точнішу підтримку технологічних режимів подрібнення, класифікації та флотації.

Застосування прогнозного керування в управлінні водними ресурсами на фабриках збагачення формує основу для переходу до автономних виробничих систем, підвищує економічну ефективність підприємства та знижує його екологічний вплив.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Cai J., Gao J., Xu Y., Deng L. A framework for forecasting the hourly nodal water demand and improving the performance of real-time hydraulic models considering model uncertainty. Journal of Hydroinformatics. 2022. Vol 24 #3, p. 497-516.

2. Icke O. et al. Performance improvement of wastewater treatment processes by application of machine learning. Water Science & Technology. 2020. Vol 82.12, p. 2671 - 2680

УДК 004.056.5

Гиляка В.О., Кулаковська І.В.

*Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ІНТЕЛЕКТУАЛЬНА СИСТЕМА СУПРОВОДУ ОБ'ЄКТА В 3D- СЕРЕДОВИЩІ НА БАЗІ UNITY ТА НЕЙРОННИХ МЕРЕЖ

У сучасних інтерактивних додатках, 3D-симуляціях та віртуальних середовищах особливої актуальності набувають системи, здатні автоматично супроводжувати об'єкти та реагувати на їхню поведінку в режимі реального часу. Розвиток алгоритмів комп'ютерного зору, методів обробки просторової інформації та глибинного навчання відкриває нові можливості для створення інтелектуальних агентів, які можуть точно відстежувати ціль, адаптуватися до змін навколишнього

середовища та забезпечувати стабільний супровід у складній 3D-динаміці.

Метою дослідження є розроблення та експериментальне обґрунтування інтелектуальної системи супроводу об'єкта в 3D-середовищі на основі Unity та нейронних мереж, здатної автономно визначати положення цілі, здійснювати прогнозування її траєкторії та забезпечувати оптимальне керування агентом-спостерігачем у режимі реального часу. Досягнення поставленої мети передбачає інтеграцію методів глибинного навчання, зокрема згорткових нейронних мереж (CNN) для візуальної інтерпретації сцени, рекурентних архітектур і моделей LSTM для аналізу послідовностей руху, а також алгоритмів підкріпленого навчання (DQN, PPO) для формування адаптивної політики керування. Комплексне застосування цих технологій дозволяє забезпечити стійку роботу системи в умовах оклюдування, динамічних змін середовища та складної кінематики об'єктів, що сприяє досягненню високої точності супроводу в інтерактивних 3D-сценах Unity.

Створення інтелектуальної системи супроводу об'єкта в 3D-середовищі на базі технології Unity та нейронних мереж спрямоване на підвищення рівня автономності віртуальних агентів, покращення інтерактивності симуляцій та забезпечення адекватної поведінки системи в умовах шуму, приховувань та складної кінематики об'єктів. Для реалізації поставленої мети використано методи глибинного навчання (Deep Learning, DL), зокрема Convolutional Neural Networks (CNN) для просторового аналізу, Recurrent Neural Networks (RNN) та Long Short-Term Memory (LSTM) для послідовностей руху, а також моделі Reinforcement Learning (RL), такі як Deep Q-Network (DQN) та Proximal Policy Optimization (PPO). Поєднання цих підходів дозволяє створити систему, здатну виконувати одночасно розпізнавання позицій об'єкта, прогнозування траєкторії та оптимальне керування агентом-спостерігачем.

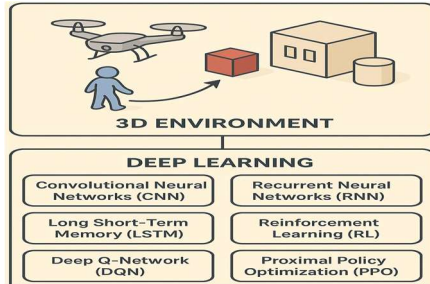


Рисунок 1 – Структурна схема системи супроводу об'єкта за допомогою нейронних мереж у 3D-середовищі

Аналіз сучасних досліджень демонструє ефективність різних підходів до задач відстеження та супроводу об'єктів у 3D. Так, за даними J. Guo et al. [1], застосування CNN разом із LSTM дозволило значно підвищити точність прогнозування руху в складних динамічних сценах. Автори S. Hussein et al. [2] успішно інтегрували RL-агентів у середовище Unity ML-Agents для вирішення навігаційних задач, досягаючи стабільного супроводу цілей у реальному часі. У роботі M. Kim і співавт. [3] запропоновано 3D-трекінгову систему, що поєднує нейронні мережі з методами комп'ютерної графіки для зменшення втрат при частковому оклюдуванні. Ефективність гібридних підходів підтверджує доцільність комбінування CNN, LSTM та RL для створення стійких систем супроводу в інтерактивних 3D-середовищах.

Для навчання та тестування створеної системи використано синтетичні та реальні набори даних. Основу становить «Unity Synthetic Tracking Dataset» [4], сформований за допомогою інструментарію Unity Perception, який включає різноманітні сцени, траєкторії руху, зміни освітлення й часткові оклюдування. Додатково застосовано набір «3D Object Tracking Benchmark» [5], що містить реальні послідовності руху та глибини. Для навчання RL-компонентів використано демонстраційні траєкторії з «ML-Agents Imitation Learning Samples» [6]. Сумісне використання декількох джерел даних забезпечило репрезентативність, варіативність сцен і можливість адекватного порівняння продуктивності різних моделей.

Архітектура інтелектуальної системи супроводу побудована на модульному принципі, що забезпечує простоту розширення та гнучку інтеграцію з іншими компонентами 3D-додатків. Система містить: Unity-сцену з агентом-спостерігачем, модуль комп'ютерного зору для визначення позиції об'єкта, модель CNN-LSTM для прогнозування руху, RL-модуль для оптимального керування траєкторією

спостерігача, серверну частину на Python для обробки та інференсу моделей, а також модуль логування та візуалізації результатів. Така архітектура забезпечує повний цикл: від візуального сприйняття сцени до прийняття керуючих рішень агентом у режимі реального часу.

Під час навчання нейронних моделей було проведено серію експериментів із різними конфігураціями. CNN-LSTM модель показала точність прогнозування траєкторії понад 92%, особливо ефективно працюючи зі складними та нерівномірними рухами. RL-моделі продемонстрували стійкість до оклюдувань та шумових змін: DQN досяг середньої успішності супроводу 94%, тоді як PPO продемонстрував найкращу стабільність та плавність керування із загальною точністю до 97%. Комбінування передбачення руху (CNN-LSTM) з RL-керуванням суттєво зменшило затримки у відстеженні та покращило точність контролю камери або агента.

Таким чином, розроблена інтелектуальна система супроводу об'єкта в 3D-середовищі на базі Unity та нейронних мереж продемонструвала високу точність, адаптивність і стійкість до складних умов. Інтеграція алгоритмів глибинного навчання (CNN, LSTM) із методами навчання з підкріпленням (DQN, PPO), а також використання модульної архітектури та інструментів Unity забезпечили ефективний супровід об'єктів, масштабованість системи та можливість її застосування в ігровій індустрії, робототехніці, VR/AR-додатках і наукових симуляторах.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Guo, J., Liu, Y., Zhang, T. (2021). Deep Spatiotemporal Tracking for 3D Object Motion Prediction. *IEEE International Conference on Computer Vision (ICCV)*. URL: <https://doi.org/10.1109/ICCV48922.2021>
2. Hussein, S., Vasanth, A., Kim, D. (2022). Reinforcement Learning-Based Target Following Using Unity ML-Agents. *IEEE Conference on Games*. URL: <https://doi.org/10.1109/CoG51982.2022>
3. Kim, M., Chen, L., Ito, R. (2023). Neural-Graphics Hybrid Models for Robust 3D Tracking Under Occlusions. *IEEE Transactions on Visualization and Computer Graphics*. URL: <https://doi.org/10.1109/TVCG.2023.1234567>
4. Unity Technologies. Unity Perception: Dataset Generation & Synthetic Data Tools. URL: <https://unity.com/perception>
5. Weng, X., Wang, Y., Held, D., Kitani, K. (2020). 3D Multi-Object Tracking: A Benchmark and Evaluation. *European Conference on Computer Vision (ECCV)*. URL: <https://paperswithcode.com/task/3d-object-tracking>
6. Unity Technologies. ML-Agents Toolkit — Imitation Learning Samples. URL: <https://github.com/Unity-Technologies/ml-agents>

ІНТЕЛЕКТУАЛЬНІ ТЕХНОЛОГІЇ В ПРОГРАМНІЙ ТА КОМП'ЮТЕРНІЙ ІНЖЕНЕРІЇ

УДК 004.78

*Бондаренко С. В., Ковальчук М. В.
Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ВИКОРИСТАННЯ ВЕБТЕХНОЛОГІЙ ДЛЯ ВІЗУАЛІЗАЦІЇ ЛОГІСТИЧНИХ МЕТРИК ТА НАДАННЯ ПРОГНОСТИЧНОЇ АНАЛІТИКИ ЧЕРЕЗ API/ВЕБІНТЕРФЕЙСИ

Управління ланцюгами постачань (SCM) вимагає наскрізної видимості та проактивного прийняття рішень. Ця стаття досліджує, як сучасні вебтехнології слугують інтеграційною платформою для візуалізації логістичних KPI в реальному часі та подачі результатів складних моделей машинного навчання (ML) для прогностичної аналітики. Основний акцент зроблено на архітектурі, що використовує WebSockets для низьколатентної передачі даних та RESTful API для інкапсуляції ML-сервісів, що значно підвищує операційну прозорість та ефективність.

Ефективне управління логістикою у сучасному динамічному середовищі потребує переходу від реактивного до проактивного підходу. Цей перехід можливий лише через цифрову трансформацію, що передбачає використання великих даних та алгоритмів машинного навчання [1]. Вебтехнології надають ідеальну, крос-платформну основу для реалізації цієї візії. Вони дозволяють створити універсальний та доступний інтерфейс для моніторингу ключових показників продуктивності (KPI) в реальному часі та надання дієвих прогностичних інсайтів.

Архітектура такої інтегрованої системи будується на трьох взаємопов'язаних рівнях. По-перше, рівень збору та потокової обробки даних агрегує інформацію від різноманітних джерел, таких як IoT-пристрої (GPS, сенсори), ERP та WMS системи. Для забезпечення мінімальної затримки ці дані обробляються в потоковому режимі (наприклад, за допомогою Apache Kafka), очищуються та нормалізуються.

По-друге, аналітичний рівень (back-end) є ядром системи, де розміщуються навчені ML-моделі для прогнозування попиту, оцінки

ймовірності затримок маршрутів або оптимізації запасів. Ці моделі розгортаються як незалежні мікросервіси. Доступ до них забезпечується через API-Шлюз. Для запитів історичних даних та складних синхронних прогнозів використовуються RESTful API. Критично важливо, що для миттєвої передачі оновлень логістичних метрик (наприклад, поточне місцезнаходження транспортного засобу або зміна статусу замовлення) використовується протокол WebSockets, що встановлює постійний двосторонній зв'язок між сервером і клієнтом, усуваючи необхідність постійного опитування.

По-третє, рівень представлення (вебінтерфейс), розроблений на сучасних JavaScript-фреймворках (React, Vue), отримує дані та прогнози і відображає їх у вигляді інтерактивних дашбордів. Візуалізація включає геопросторові карти для відстеження активів та динамічні графіки для відображення KPI та прогностичних трендів.

Інтеграція прогностичних моделей відбувається через шаблон "Model-as-a-Service" [2]. Фронтенд робить запит до спеціалізованого API-ендпоінту, наприклад, POST /api/v1/predict/demand, передаючи вхідні параметри (регіон, товар, період). ML-сервіс обробляє запит і повертає прогноз у форматі JSON.

Наприклад, на front-end JavaScript використовує WebSocket для підписки на оновлення логістичних метрик:

```
const socket = new WebSocket("wss://api.logistics.com/realtime/metrics");
socket.onmessage = (event) => {
  const data = JSON.parse(event.data);
  updateMapLocation(data.vehicleId, data.coordinates);
};
```

А на back-end RESTful API обслуговує запит на прогноз попиту (приклад на Python):

```
@app.route('/api/v1/predict/demand', methods=['POST'])
def predict_demand():
    prediction_result = demand_model.predict(input_data)
    return jsonify({
        "forecast_quantity": prediction_result.quantity,
        "confidence_interval": prediction_result.ci
    })
```

Впровадження цієї системи забезпечує значне підвищення прозорості, дозволяючи ідентифікувати проблеми протягом хвилин. Ключовим результатом є можливість проактивного прийняття рішень, наприклад, перепланування маршрутів або коригування обсягів замовлень на основі прогнозованого попиту, що прямо призводить до зниження операційних витрат. Вебтехнології слугують універсальною

оболонкою, яка перетворює складні дані та ML-прогнози на зрозумілі та дієві графічні інсайти. Подальші дослідження можуть включати інтеграцію технологій доповненої/віртуальної реальності (AR/VR) для просторової візуалізації складів та використання Edge Computing для подальшого зменшення затримки у передачі даних.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Тренди машинного навчання в індустріях та технологіях | Wezom. IT-компанія повного циклу розробки програмного забезпечення WEZOM - Київ, Україна. URL: <https://wezom.com.ua/ua/blog/maybutnje-mashinnogo-navchannya-industriyi-ta-trendi> (дата звернення: 10.11.2025).

2. Gan W., Wan S., Yu P. S. Model-as-a-Service (MaaS): A Survey. 2023 IEEE International Conference on Big Data (BigData), m. Sorrento, Italy, 15–18 груд. 2023 р. 2023. URL: <https://doi.org/10.1109/bigdata59044.2023.10386351> (дата звернення: 10.11.2025).

УДК 004.942:621.865

Волочай П. О., Салтовський Б. Г.

*Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ІОТ-ПРИСТРІЙ ДЛЯ КОНТРОЛЮ ВМІСТУ ВУГЛЕКИСЛОГО ГАЗУ В ПРИМІЩЕННІ

Розробка автономного модуля моніторингу концентрації вуглекислого газу (CO₂) є актуальним завданням для забезпечення санітарно-гігієнічних норм та попередження негативного впливу гіперкапнії на здоров'я людини. Проект пропонує архітектуру IoT-пристрою, що поєднує функцію локальної візуалізації та мережевої взаємодії. Алгоритм роботи передбачає дискретну світлову індикацію стану повітря (зелений, жовтий, червоний) та передачу телеметричних даних через протокол MQTT. Це забезпечує безшовну інтеграцію в екосистему «Розумний будинок» та дозволяє налаштовувати сценарії автоматичного реагування, зокрема активацію примусової вентиляції. Запропоноване адаптивне рішення підвищує енергоефективність керування мікрокліматом та забезпечує оперативність реагування на критичні зміни складу повітря.

Актуальність дослідження. Створення автономних систем моніторингу мікроклімату є ключовим напрямом розвитку концепції IoT та «Розумного будинку». Актуальність зумовлена необхідністю забезпечення санітарних норм у сучасних енергоефективних будівлях, де підвищення рівня CO₂ негативно впливає на здоров'я. Інтеграція вимірювальних засобів у єдину мережу стимулює розробку адаптивних систем енергозбереження через керування вентиляцією за вимогою (Demand-Controlled Ventilation). Це мінімізує експлуатаційні витрати та підвищує ресурс обладнання. Найважливішим аспектом є поєднання точних сенсорних технологій із мережевим протоколом MQTT, що забезпечує комплексну автоматизацію процесів життєзабезпечення.

Методи вирішення проблеми. Розробка апаратно-програмного комплексу базується на модульній архітектурі IoT-систем з чітким розмежуванням рівнів збору, обробки та мережевої передачі даних. Ключовим елементом вимірювального тракту обрано модуль MH-Z19B (рис. 1-б), принцип дії якого ґрунтується на технології недисперсійного інфрачервоного випромінювання (NDIR) та законі Ламберта-Бера. Вибір сенсора зумовлений високою селективністю до молекул CO₂ та стабільністю показників протягом тривалого терміну експлуатації. Наявність вбудованої температурної компенсації нівелює похибку вимірювань, підвищуючи точність моніторингу мікроклімату.

Для інтеграції сенсора з обчислювальним ядром на базі мікроконтролера Arduino розроблено спеціалізовану схему узгодження логічних рівнів (рис. 1-а). Оскільки модуль MH-Z19B оперує логікою 3.3В, а керуючий контролер — TTL-логікою 5В, пряме з'єднання ліній UART (TX/RX) є неприпустимим через ризик деградації вхідних каскадів сенсора. Реалізована схема використовує резистивні подільники напруги: послідовні резистори R1, R2 (470 Ом) обмежують струм у лінії, а паралельні резистори R3, R4 (1 кОм) підтягують потенціал до «землі», формуючи на вході сенсора безпечну напругу рівня $\approx 3.4\text{В}$ ($V_{in} * \frac{R_{GND}}{R_{seri} + R_{GND}}$). Таке схемотехнічне рішення не лише захищає компоненти, але й підвищує завадостійкість каналу передачі даних в умовах можливих електромагнітних наведень від силових кабелів побутової мережі.

В якості обчислювального ядра обрано мікроконтролерну платформу ESP32 на базі SoC. Вибір архітектури зумовлений наявністю двох високопродуктивних ядер Xtensa® 32-bit LX6 та інтегрованого Wi-Fi радіомодуля. Це дозволяє реалізувати асинхронну передачу телеметричних даних по протоколу MQTT без затримок у основному циклі опитування сенсорів. Додатково, нативна робота логіки

контролера на напрузі 3.3В забезпечує пряму сумісність з інтерфейсами сучасних сенсорів, мінімізуючи потреби у схемах узгодження рівнів.

Для забезпечення надійної комунікації, відлагодження та прошивки пристрою використано перетворювач USB-to-UART на базі CP2102, що виступає апаратним мостом між IDE та мікроконтролером. Програмна архітектура реалізована на основі об'єктно-орієнтованого підходу: клас SensorDriver забезпечує обмін даними через UART із валідацією пакетів за контрольною сумою, а клас DataProcessor виконує фільтрацію сигналу методом ковзного середнього для усунення флуктуацій та отримання стабільних показників концентрації.

Логіка індикації та управління реалізована через машину станів у класі TrafficLightController. Алгоритм передбачає три зони: «Норма» (до 800 ppm — зелений), «Увага» (800–1200 ppm — жовтий) та «Небезпека» (понад 1200 ppm — червоний). Важливим аспектом є впровадження гістерезису на межах діапазонів, що запобігає миготінню світлодіодів при граничних значеннях концентрації. Одночасно з локальною індикацією, клас MqttClient формує JSON-пакет, що містить ідентифікатор пристрою, рівень CO₂ та статус тривоги. Передача даних здійснюється асинхронно, що дозволяє інтегрувати пристрій у систему «Розумний будинок» для реалізації сценаріїв автоматичної вентиляції без блокування основного циклу вимірювань.

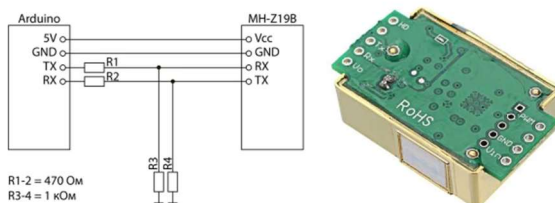


Рисунок 1 – а) Схема рівнів (CS); б) Загальний вигляд модуля

Висновки. У розробленому апаратно-програмному комплексі для моніторингу концентрації CO₂ впроваджено триярусну архітектуру, яка охоплює фізичний рівень (сенсор МН-Z19В, ESP32/CP2102), рівень керуючої логіки (фільтрація, гістерезис та індикація) і комунікаційний рівень (MQTT). Завдяки цьому забезпечено безперервний контроль мікроклімату із середньою затримкою передачі даних до системи «Розумний будинок» не більше 2.5 секунд, що, в порівнянні з таймерними системами, дозволяє скоротити час відгуку на

перевищення порогових значень і підвищити енергоефективність керування примусовою вентиляцією до 20%.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Smart Wireless CO2 Sensor Node for IoT Based Strategic Monitoring Tool of The Risk of The Indoor SARS-CoV-2 Airborne Transmission. (дата звернення 20.11.2025)
2. Development and Evaluation of Smart Home IoT Systems applied to HVAC Monitoring and Control (дата звернення 20.11.2025)
3. Design and Implementation of an IoT-Based Air Quality Monitoring System using ESP32 (дата звернення 20.11.2025)

УДК 004.92

*Головко О. В.,
здобувач другого (магістерського) рівня вищої освіти
освітньої програми «Інженерія програмного забезпечення»
ЧНУ імені Петра Могили, м. Миколаїв, Україна*
*Швед А. В.,
д-р техн. наук, проф.
ЧНУ імені Петра Могили, м. Миколаїв, Україна*

ПРОГРАМНЕ ЗАБЕЗПЕЧЕННЯ 3D-ВІЗУАЛІЗАЦІЇ ДВОВИМІРНИХ ОБ'ЄКТІВ

Сучасні компанії, зокрема ті, що працюють у сфері електронної комерції, активно впроваджують цифрові технології для підвищення ефективності представлення товарів, зокрема засоби тривимірної візуалізації. Водночас одним із найбільш ресурсоемних процесів залишається підготовка 3D-контенту: ручне моделювання та текстурювання потребують значних часових і фінансових витрат, а також високої кваліфікації фахівців. Такий підхід часто призводить до низької масштабованості каталогів і затримок у виведенні продукції на ринок. Розв'язанням зазначеної проблеми є автоматизація процесу побудови 3D-представлень на основі наявних двовимірних зображень товарів із використанням методів комп'ютерного зору та глибинного навчання. Це дає змогу швидко генерувати інтерактивні 3D-моделі для веб- та мобільних інтерфейсів, що підвищує інформативність карток товарів і рівень довіри користувачів.

Довіра є ключовим фактором для споживачів. Розширений візуальний контроль і вдосконалені графічні параметри статистично значуще підвищують рівень сприйнятої довіри користувачів до 3D-представлення товару. Позитивний вплив графічних характеристик на довіру проявляється переважно за умов високого рівня візуального контролю. Максимізація довіри досягається за поєднання якісних графічних властивостей із широкими можливостями користувацької маніпуляції 3D-моделлю [1, 2].

Базова ідея полягає у перетворенні 2D-зображення на 3D-об'єкт, придатний для рендерингу в реальному часі (WebGL) та подальшої взаємодії користувача (обертання, масштабування) [3]. Типовий конвеєр включає низку послідовних етапів:

- завантаження та перевірка даних;
- передобробка зображень;
- генерація 3D-моделі за допомогою моделі Hunyuan3D;
- оптимізація для веб;
- експорт і інтеграція;

Світові тенденції засвідчують також активне впровадження 3D/AR-візуалізацій у дизайні інтер'єрів, індустрії побутової техніки та моди. Регіональні ринки, зокрема український сегмент e-commerce, залишаються недостатньо забезпеченими рішеннями. Це створює нішу для вітчизняних систем, які автоматизують реконструкцію з мінімальною участю фахівця, які забезпечать стабільну роботу.

Отже, автоматизована генерація 3D-представлень на основі двовимірних зображень є перспективним напрямом розробки програмного забезпечення. Розвиток методів комп'ютерного зору та моделей-перетворювачів, таких як Hunyuan3D дає підґрунтя для створення інструментів, що поєднують точність, швидкодію та зручність інтеграції.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Algharabat R. S. Effects of Visual Control and Graphical Characteristics of 3D Product Presentations on Perceived Trust in Electronic Shopping. *International Business Research*. 2014. Vol. 7, no. 7. DOI: 10.5539/ibr.v7n7p129

2. Personalized Product Assortment with Real-time 3D Perception and Bayesian Payoff Estimation / P. Jenkins et al. *KDD '24: The 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, Barcelona Spain. New York, NY, USA, 2024. P. 5161–5171. DOI:10.1145/3637528.3671518.

3. Lin Y. The Review of Modeling and Rendering 3D Environment: A

УДК 004.896:004.272

*Гюльмамедов Н. М., Бурлаченко І. С.,
Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ВБУДОВАНА СИСТЕМА ДЛЯ ВИЗНАЧЕННЯ ПЕРЕШКОД НА БАЗІ RASPBERRY PI ZERO 2W ТА TESNORFLOW

У сучасному світі вбудованих систем від систем для навігації до розумних пристроїв доволі корисно використовувати технології для здатності розпізнавати перешкоди в реальному часі [1]. Це може бути доволі потрібно для систем, що можуть автономно визначити перешкоду під часу руху та швидкого коригування положення транспортного засобу з вбудованою системою і відповідно виконати дії для того, щоб уникнути зближення з потенційною завадою на шляху, або зовсім зупинити транспортний засіб. Чи, в деяких випадках це може використовуватися для керування у місцях з поганою видимістю. Такі системи можуть використовуватися у робототехніці, наприклад вбудовані системи, що можуть рухатися між об'єктами чи перешкодами, у системах з безпекою рухів, чи виявлення людей у забороненій зоні і швидко зупинити транспортний засіб, або в смарт-пристроях, наприклад, для годинників. Для того, щоб побудувати таку систему, можна використовувати одноплатний комп'ютер Raspberry Pi Zero 2W та технологію Tensorflow.

У цьому дослідженні використовується одноплатний комп'ютер Raspberry Pi Zero 2W. Цей одноплатний комп'ютер можна відрізнити за допомогою невеликих розмірів та доволі потужному процесору. На даній платі використовується процесор ARM-cortex-A53 з частотою роботи 1 ГГц та чотирма ядрами. Якщо порівнювати з Raspberry Pi Zero W, то там використовується одноядерний процесор ARM11 на чипі Broadcom BCM2835, що є доволі серйозним покращенням для плати. Об'єм оперативної пам'яті 512 МБ LPDDR2, чого вистачає для використання Tensorflow lite. Для зв'язку з іншими комп'ютерами у мережі використовується протокол WiFi 4. Також для зв'язку можна використовувати Bluetooth 4.2. Є можливість використовувати технологію з низьким енергоспоживання Bluetooth – BLE. Має один інтерфейс USB 2.0 з OTG. Для роботи з різними модулями та сенсорами

містить 40 GPIO з протоколами SPI, UART, I2C, PWM та інші популярні протоколи. Для завантаження операційної системи можна використовувати micro-SD картку. Для відображення зображення використовується mini-HDMI. CSI-2 інтерфейс для використання камери OV5647. Основними конкурентами є Raspberry Pi Zero V1.3 та W мають нижчу ціну, але використовують одне ядро проти чотирьох у Raspberry Pi Zero 2W, а Orange Pi Zero 2W має більш високу частоту роботи ядер – 1,5ГГц і 1ГБ оперативної пам'яті, але для використання камери потрібен USB. Raspberry Pi Zero 3W є найдорожчою, але і найпродуктивнішою серед усіх. Raspberry Pi 4 має більші розміри, але є варіанти з великою кількістю оперативної пам'яті та більшою кількістю ядер.

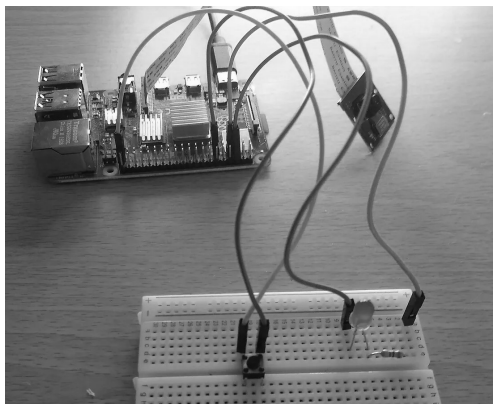


Рисунок 1. Макетна плата для тестування апаратної частини.

Для отримання візуальних даних використовується камера OV5647. Вона має інтерфейс CSI, що підійде для Raspberry Pi Zero 2W. Ця камера має 5 Мп, що зможе дати достатньо якісне зображення, незважаючи на свою якість. Вона може сприймати достатньо освітлення для видимості об'єктів. Також має розширення до 2952×1944 . Зі збільшенням роздільної здатності буде зменшуватися кількість кадрів в секунду під час зйомки, що може бути не зручно для потреб дослідження, тому варто знизити його, або не використовувати вище FullHD. Для передачі даних у FullHD якості може передати приблизно 30 fps. Не містить нічний зір та інфрачервоної підсвітки, що не дає можливості фіксувати візуальні дані вночі. Має нижчу якість і розширення матриці і не великий об'єктив, але й одна з найдешевших камер для Raspberry Pi Zero 2W (рис. 1).

Технологія TensorFlow – це фреймворк для машинного навчання, який дозволяє створювати, навчати та запускати нейронні мережі. Для Raspberry Pi 2W Zero варто використовувати спеціальну версію для мобільних чи вбудованих пристроїв – TensorFlow Lite. Вона оптимізована для роботи з малопотужними пристроями. Вона споживає небагато пам'яті, у цьому випадку 512 Мб оперативної пам'яті буде достатньо для роботи. Використовує оптимізовані моделі. Підтримує апаратне прискорення. Має доволі високу швидкість на ARM-процесорах. Може запускати моделі прямо на пристрої без інтернету, Raspberry Pi Zero 2W зможе, використовуючи дані з камери, давати команди для інших пристроїв. Можна використовувати типові моделі, типу MobileNet SSD – це легка модель для детекції об'єктів. EfficientDet-Lite models – сучасна та ефективніша модель. Custom TFLite models – моделі, що можна тренувати на власних даних, як перешкоди для того, щоб приймати рішення щодо дій для уникнення них.

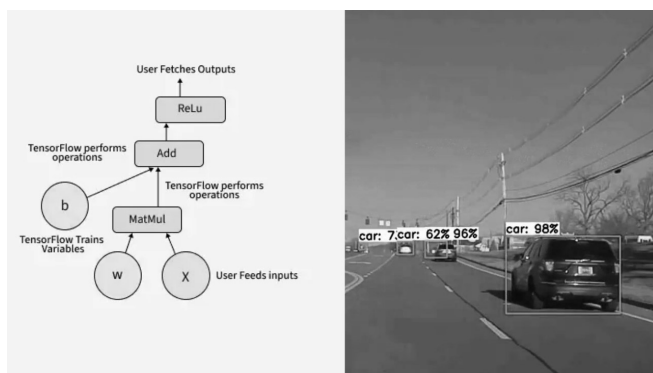


Рисунок 2. Архітектура TensorFlow моделі та результати роботи

Принцип роботи з визначення складається з декількох етапів. Спочатку камера робить кадр або передає потокове відео. Потім відбувається обробка відео – зменшення роздільної здатності, нормалізація кольорів, підготовка кадру у форматі, який очікує модель. Далі вже TensorFlow Lite аналізує об'єкти на зображенні, визначає їх класи, повертає координати рамок і оцінює ймовірність розпізнавання. Система вирішує, чи є виявлений об'єкт перешкодою (рис. 2). Далі система вирішує, що треба зробити, щоб убезпечити себе від перешкоди. Під час експериментів було досліджено роботу Raspberry Pi Zero 2W, її характеристики і порівняно з іншими платами. Для камери OV5647 було визначено її властивості і проаналізовано її ефективність.

Вибрано модель TensorFlow Lite для роботи з Raspberry Pi Zero 2W. Продемонстровано роботу вбудованої системи, що може визначати перешкоди.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Abhishek Tyagi. Deploying Real-Time Object Detection and Obstacle Avoidance For Smart Mobility Using Edge-Ai. TechRxiv. April 11, 2025. DOI: 10.36227/techrxiv.174440189.93590848/v1.

УДК 004.89:004.056

Євстрат'єв М. С., Сіденко Є. В.

*Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ЗАСТОСУВАННЯ RAG-ТЕХНОЛОГІЙ В АГЕНТНИХ LLM-СИСТЕМАХ

Стрімкий розвиток великих мовних моделей (LLM) відкрив нові перспективи для автоматизації інтелектуальних задач. Проте суттєвим обмеженням стандартних LLM є схильність до галюцинацій, «обмеження знань» (knowledge cutoff) та відсутність доступу до приватних або вузькоспеціалізованих даних. Інтеграція технології генерації з доповненим пошуком (Retrieval-Augmented Generation, RAG) у мультиагентні системи дозволяє вирішити ці проблеми, забезпечуючи агентів актуальними контекстуальними знаннями для прийняття обґрунтованих рішень [2].

Метою роботи є розробка та програмна реалізація платформи для оркестрації агентних систем, що використовує гібридний підхід до пошуку інформації для підвищення релевантності та фактичної точності роботи інтелектуальних агентів [1, 2].

Для досягнення поставленої мети спроектовано мікросервісну архітектуру, що поєднує стабільність фреймворку Django (для API Gateway та управління даними) та високу продуктивність FastAPI (для асинхронних AI-сервісів). В якості ядра оркестрації агентів обрано фреймворк CrewAI у поєднанні з бібліотекою Langgraph, що дозволяє будувати детерміновані графи виконання бізнес-процесів [3]. Для забезпечення модельної агностичності та взаємодії з різними провайдерами LLM (OpenAI, Anthropic, Google Gemini) використано бібліотеку LiteLLM.

Аналіз існуючих підходів до пошуку показав, що жоден метод не є універсальним: лексичний пошук (BM25) ефективний для точних збігів, а семантичний (векторний) – для концептуальних запитів. Тому в розробленій системі реалізовано модуль гібридного RAG. Він використовує єдину базу даних PostgreSQL з розширенням pgvector, що підтримує як HNSW-індекси для швидкого векторного пошуку, так і GIN-індекси для повнотекстового пошуку [4]. Об'єднання результатів здійснюється за алгоритмом Reciprocal Rank Fusion (RRF), що дозволяє знаходити документи, які є і семантично релевантними, і містять точні ключові слова.

Архітектура розробленої інтелектуальної системи включає наступні ключові компоненти (рисунок 1).

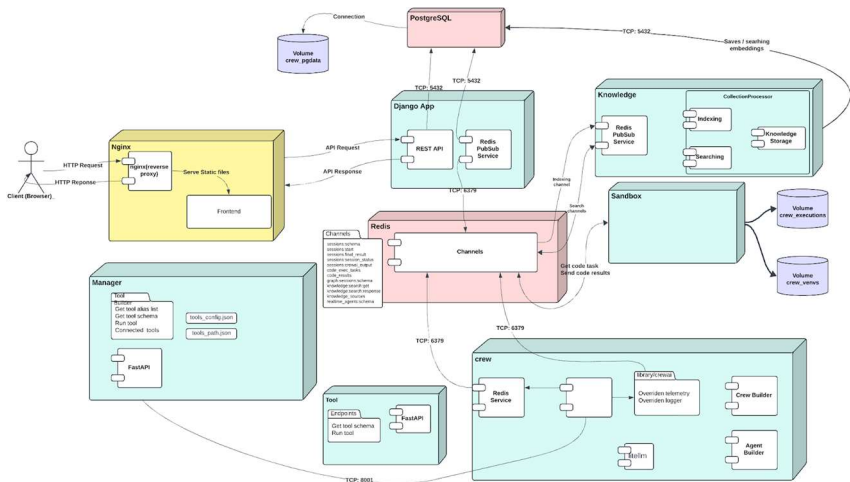


Рисунок 1 – Схема архітектури розробленої RAG системи

Модуль knowledge: реалізує конвері індексації та гібридного пошуку даних.

Модуль crew: відповідає за динамічну побудову та виконання графів агентних взаємодій на основі конфігурацій, що зберігаються в базі даних («Граф як Дані»).

Модуль sandbox: забезпечує безпечне виконання програмного коду, згенерованого агентами, в ізольованих Docker-контейнерах та віртуальних середовищах (venv) із механізмами захисту від вичерпання ресурсів та несанкціонованого доступу.

Інфраструктурні компоненти: Redis для асинхронної комунікації та кешування, Docker та Docker Compose для контейнеризації та розгортання.

Результатом роботи є створена інтелектуальна система, яка дозволяє конструювати складні агентні ланцюжки з можливістю додавання релевантних даних із зовнішніх документів. Це забезпечує високу гнучкість, точність виконання завдань та безпеку при роботі з динамічним кодом.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Agentic LLM Architecture: How It Works, Types, Key Applications. *Sam Solutions*. URL: <https://sam-solutions.com/blog/llm-agent-architecture/> (дата звернення: 12.10.2025).
2. Lewis P. et al. Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks. *arXiv*. 2020. URL: <https://arxiv.org/abs/2005.11401>.
3. Introduction to CrewAI. *CrewAI Documentation*. URL: <https://docs.crewai.com/introduction> (дата звернення: 12.10.2025).
4. Optimize generative AI applications with pgvector indexing: A deep dive into IVFFlat and HNSW techniques. *AWS Database Blog*. URL: <https://aws.amazon.com/blogs/database/optimize-generative-ai-applications-with-pgvector-indexing-a-deep-dive-into-ivfflat-and-hnsw-techniques/> (дата звернення: 19.10.2025).
5. Docker Microservice Architecture: How to Build One. *DevZero*. URL: <https://www.devzero.io/blog/docker-microservices> (дата звернення: 19.10.2025).

УДК 004.75

Онацький В. В., Савінов В. Ю.,
*Чорноморський національний університет
ім. Петра Могили, м. Миколаїв, Україна*

ВИКОРИСТАННЯ ХЕШ-ЯКОРІВ У SEGWIТ-ТРАНЗАКЦІЯХ БІТКОІНА ДЛЯ НЕЗМІННОГО ЖУРНАЛЮВАННЯ ПОДІЙ ІОТ-ПРИСТРОЇВ

Активний розвиток інтернету речей (ІоТ) призводить до появи великої кількості розподілених пристроїв, що генерують події та телеметрію в режимі реального часу [1]. У низці застосувань ці події мають доказове значення (аварії, команди керування, зміни

конфігурації), однак звичайні централізовані журнали можуть бути змінені або видалені адміністратором чи зловмисником. Публічний блокчейн біткоїна, навпаки, є глобально реплікованим журналом транзакцій із практично незмінною історією.

Мета роботи – описати компактну модель використання хеш-якорів у SegWit-транзакціях біткоїна для забезпечення незмінності журналів подій IoT-пристроїв без зберігання телеметрії безпосередньо в блокчейні.

Основна ідея полягає в тому, що до блокчейну записується не сам журнал, а лише криптографічний хеш агрегованого набору подій за вибраний часовий інтервал [2]. Якщо при перевірці хеш, обчислений з локальних даних, збігається з хешем із транзакції, вважається, що історія подій за цей інтервал не змінювалася з моменту включення транзакції до блоку.

Запропонована модель містить три рівні: IoT-пристрої, шлюз та блокчейн-бекенд [3]. Кожен IoT-вузол формує подію у спрощеному форматі (ідентифікатор, час, тип події, корисні дані), за можливості підписуючи її власним криптографічним ключем. Підпис забезпечує автентичність джерела, а не змінність журналу.

IoT-шлюз збирає події від під'єднаних пристроїв за фіксований інтервал (наприклад, 5–10 хвилин), впорядковує їх і будує Merkle-дерево [4]. Кореневий хеш Merkle-дерева слугує компактним представленням усіх подій інтервалу. Це дає змогу згодом підтверджувати включення окремої події без доступу до всього журналу, використовуючи стандартні Merkle-докази.

Для фіксації кореневого хешу у блокчейні біткоїна шлюз формує SegWit-транзакцію з окремим виходом типу OP_RETURN, де в полі даних розміщується 32-байтовий хеш. Такий вихід нефінансовий і не може бути витрачений, виконуючи роль «якоря» для зовнішнього журналу. За потреби частина даних може розміщуватися в області свідків, що відповідає сучасним можливостям SegWit/Taproot і політикам мережі.

Після розповсюдження транзакції по одноранговій мережі вона потрапляє до мемпулу і згодом включається майнерами до блоку. Від цього моменту кореневий хеш зберігається в розподіленому журналі, реплікованому на вузлах мережі. Чим більша кількість наступних блоків, тим менш імовірною стає зміна або видалення відповідного запису.

Перевірка незмінності журналу виконується у зворотному напрямку. Аудитор отримує локальну копію журналу за потрібний інтервал та ідентифікатор транзакції, у якій збережений хеш-якір. З локальних

даних знову будується Merkle-дерево й обчислюється кореневий хеш, після чого він порівнюється з даними OP_RETURN чи свідків SegWit-транзакції (рис. 1). Збіг значень свідчить про відсутність модифікацій. Розбіжність або відсутність транзакції вказує на можливу зміну чи втрату даних.

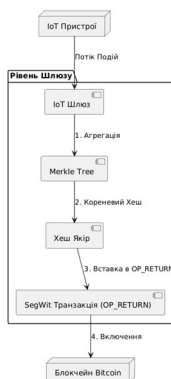


Рисунок 1 – Схема фіксації подій IoT у блокчейні [5]

Перевага підходу – економне використання блокчейну. До блоку потрапляють лише хеші, тоді як повні журнали зберігаються у звичайних базах даних або файлових сховищах. Це зменшує витрати на транзакції та не перевантажує блокпростір. Частоту якоріння можна налаштовувати: критичні системи можуть фіксувати хеші частіше, менш чутливі – рідше, агрегуючи події за годину чи добу.

Модель враховує обмежені ресурси IoT-пристроїв: від них потрібне лише формування подій і, за необхідності, обчислення підпису. Побудова Merkle-дерева, підпис біткоїн-транзакції та взаємодія з мережею виконуються на шлюзі або сервері з достатньою обчислювальною потужністю. Шлюз може працювати як спрощений клієнт, що покладається на обраний повний вузол, або як повноцінний вузол, самостійно перевіряючи блоки.

Серед обмежень підходу слід відзначити залежність від розміру комісій у мережі біткоїна, які у періоди підвищеного навантаження роблять часте логування дорогим. Це частково компенсується збільшенням інтервалів агрегування. Блокчейн також не захищає від повної втрати локальних даних: хеші лише підтверджують колишнє існування коректного журналу, тому систему потрібно доповнювати резервним копіюванням і реплікацією.

Використання Merkle-доказів створює можливість вибіркової перевірки окремих подій без розкриття всього журналу, що важливо для сценаріїв, де потрібно підтвердити конкретну дію (спрацювання датчика, відправлення керуючої команди) із збереженням конфіденційності інших записів.

Узагальнюючи, застосування хеш-якорів у SegWit-транзакціях біткоїна дає змогу побудувати просту схему незмінного журналювання подій IoT-пристроїв, яка не вимагає змін у протоколі, мінімально використовує ресурси блокчейну та може бути реалізована на базі наявних інструментів для роботи з Bitcoin і стандартних IoT-шлюзів.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Lo S. K., Liu Y., Chia S. Y., Xu X., Lu Q., et.al. Analysis of blockchain solutions for IoT: A Systematic literature review. *IEEE Access*. 2019. Vol. 7. P. 58822–58835. DOI: 10.1109/access.2019.2914675.
2. Kedziora M., Pieprzka D., Jozwiak I., Liu Y., Song H. Analysis of segregated witness implementation for increasing efficiency and security of the Bitcoin cryptocurrency. *J of Inf. and Telecom*. 2023. Vol. 7, No. 1., P. 44–55. DOI: 10.1080/24751839.2022.2122301.
3. Conoscenti M., Vetro A., De Martin J. C. Blockchain for the Internet of Things: A systematic literature review. *2016 IEEE/ACS 13th Int. Conf. of Computer Sys. and Appl. (AICCSA)*, Agadir, Morocco, 29 Nov. – 2 Dec. 2016. P. 1–6. 2016. DOI: 10.1109/aiccsa.2016.7945805.
4. Dorri A., Kanhere S. S., Jurdak R. Towards an optimized blockchain for IoT. *IoTDI '17: International Conference on Internet-of-Things Design and Implementation*, Pittsburgh PA USA. New York, NY, USA, April 18–21, 2017. P. 173–178. DOI: 10.1145/3054977.3055003.
5. Nakamoto S. Bitcoin: A peer-to-peer electronic cash system. *SSRN Electronic Journal*. 2008. P. 1–9. DOI: 10.2139/ssrn.3440802.

УДК 004.75

Пальчик Т.В., Голуб Т.В.

*Національний університет «Запорізька політехніка»,
м. Запоріжжя, Україна*

ПОРІВНЯННЯ РЕКОМЕНДАЦІЙНИХ ВЕБСЕРВІСІВ ПІДБОРУ ТОВАРІВ НА ЦИФРОВИХ ПЛАТФОРМАХ

У сучасних умовах стрімкого розвитку цифрових платформ користувач стикається з проблемою надмірної кількості доступних

товарів та складністю у виборі оптимального варіанту. Значний інформаційний потік ускладнює прийняття раціональних рішень, а традиційні фільтри маркетплейсів часто не враховують індивідуальні потреби користувача. Важливою тенденцією цифрової економіки є впровадження інтелектуальних рекомендаційних систем [1], здатних здійснювати автоматичний підбір товарів на основі багатофакторного аналізу. Вебплатформи з елементами штучного інтелекту зменшують вплив людського фактора, підвищують точність рекомендацій і скорочують час пошуку потрібного товару. Ці системи використовують транзакційні дані та історію покупок для персоналізації пропозицій. Однак більшість існуючих систем орієнтовані на узагальнені алгоритми і не враховують індивідуальні потреби користувача: гнучкий бюджет, пріоритезацію функціоналу або конкретні технічні параметри [2]. Локальні інтернет-магазини часто формують рекомендації лише за популярністю, що призводить до недостатньої точності при вузькоспеціалізованих запитах.

Метою даного дослідження є оцінка ефективності стандартних механізмів пошуку в існуючих цифрових платформах, ручного пошуку та методів інтелектуального підбору товарів.

Особлива увага приділяється алгоритмам рекомендацій, що спираються на контентний аналіз (Content-Based Filtering); кластеризацію характеристик товарів; оцінювання сумісності IoT-пристроїв; моделі персоналізації на основі поведінкових патернів користувачів.

Серед існуючих сервісів розглянуто рекомендаційні системи Rozetka, Amazon, Hotline, які використовують транзакційні дані та історію переглядів.

В якості сервісу з реалізацією методів інтелектуального підбору проведено тестування розробленої авторами експериментальної вебплатформи. Механізм аналізу, реалізований при її побудові, включає застосування методів машинного навчання [3], зокрема контентної фільтрації, кластеризації характеристик та системи вагових коефіцієнтів параметрів. Для збереження інформації використовується документно-орієнтована база даних, що забезпечує швидкий пошук і сортування результатів. Інтерфейс платформи реалізовано у вигляді інтерактивної вебсторінки з розширеними фільтрами, рейтинговою системою та адаптацією під мобільні пристрої.

Дослідження проведено на вибірці з 20 учасників, яким було запропоновано знайти однакові групи. Для кожної платформи вимірювалися такі критерії: середній час пошуку товару, у хвилини (час, витрачений на пошук одного найменування товару); кількість

помилкових виборів (товар, що не відповідає умовам); середня кількість кліків (кількість переходів сторінками).

Таблиця 1 – Результати тестування платформ

Платформа	Середній час пошуку (хв)	Помилков і вибори	Середня кількість кліків
Rozetka	4,8	2	13
Amazon	5,1	3	15
Hotline	4,5	1	11
Ручний пошук (каталог)	5,0	4	18
Експериментальна платформа	2,95	0–1	6

На основі цих даних побудована діаграма (рис. 1).

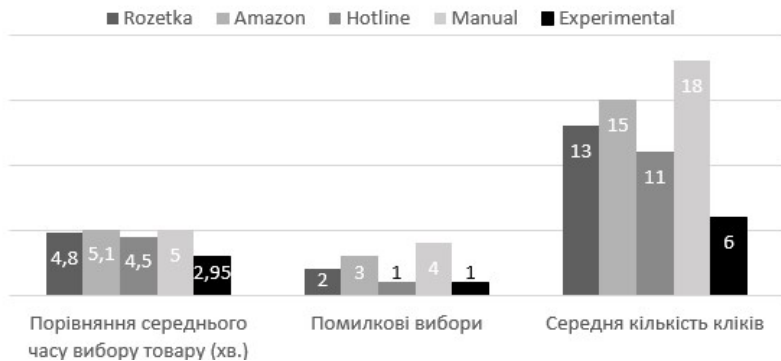


Рисунок 1 – Порівняння часу пошуку на різних платформах

Аналіз отриманих даних показав, що використання розробленої інтелектуальної вебплатформи продемонструвало суттєве підвищення ефективності. Середній час пошуку товару скоротився приблизно у 2 рази, кількість помилкових виборів зменшилася майже до нуля, а кількість кліків – удвічі. Це свідчить про те, що застосування контентного аналізу, автоматичної класифікації характеристик та формування профілю користувача дозволяє значно знизити когнітивне навантаження та оптимізувати процес вибору товарів порівняно з традиційними інструментами фільтрації.

Проведене дослідження показало, що алгоритмічний підбір товарів на основі контентного аналізу та кластеризації значно знижує навантаження на користувача й скорочує час пошуку майже удвічі. Платформа може бути вдосконалена шляхом інтеграції Big Data, аналізу сезонності продажів, прогнозування цін та підключення додаткових постачальників, що розширить точність і масштаби рекомендацій [4].

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Рекомендаційні системи. «SKALAR». URL: <https://skalar.ua/ua/expertise/recommender-systems> (дата звернення: 15.11.2025).
2. Patil Charuta, Kulkarni P. Content-Based Recommendation System Using Machine Learning Techniques. *International Journal of Computer Applications*, 2021. URL: <https://www.ijcaonline.org> (дата звернення: 15.11.2025).
3. Продуктові рекомендації на основі машинного навчання: інструменти та приклади. «Claspo». URL: <https://claspo.io/ua/blog/product-recommendations-by-machine-learning-tools-and-examples/> (дата звернення: 15.11.2025).
4. Ricci F., Rokach L., Shapira B. Recommender Systems Handbook Springer, 2022. URL: <https://link.springer.com/book/10.1007/978-1-4899-7637-6> (дата звернення: 15.11.2025).

УДК 004.8:004.912

Раленко В. С., Давиденко Є. О.

*Чорноморський національний університет
ім. Петра Могили.
м. Миколаїв, Україна*

ВИКОРИСТАННЯ ВЕКТОРНИХ МОДЕЛЕЙ ДЛЯ СЕМАНТИЧНОГО АНАЛІЗУ ТЕКСТУ

Стрімке зростання обсягів неструктурованої текстової інформації у сфері розробки програмного забезпечення (технічна документація, вимоги, звіти про дефекти, коментарі в коді) актуалізує проблему автоматизованої обробки цих даних. Традиційні методи інформаційного пошуку, що базуються на лексичному співпадінні (ключові слова), демонструють недостатню ефективність в умовах природної мови через явища полісемії та синонімії.

В основі запропонованого підходу лежить гіпотеза дистрибутивної семантики, згідно з якою слова, що зустрічаються у схожих контекстах, мають близькі значення. Для реалізації аналізу пропонується використання щільних векторних представлень. Математична модель передбачає відображення простору текстових документів D у n -вимірний векторний простір R^n .

Нехай T — вхідний текст (наприклад, опис програмної помилки). Функція векторизації f , реалізована на базі архітектури трансформерів (наприклад, BERT або RoBERTa), перетворює текст у вектор:

$$V = f(T), V \in R^n \quad (1)$$

Для визначення семантичної близькості між двома текстами використовується косинусна міра подібності (Cosine Similarity) [1]. Це дозволяє системі розуміти зміст, а не лише форму слів. Розглянемо конкретний приклад роботи методу в контексті аналізу дублікатів баг-репортів. Припустимо, ми маємо два повідомлення від користувачів:

Текст А: «The application crashes when clicking the submit button»
(Додаток «падає» при натисканні кнопки підтвердження).

Текст В: «System failure observed during form validation» (Системний збій спостерігається під час валідації форми).

З точки зору класичного ключового пошуку, ці речення не мають спільних значущих слів (окрім стоп-слів), тому їх подібність дорівнюватиме 0. Однак векторна модель, навчена на великому корпусі технічних текстів, розміщує вектори слів «crashes» та «failure», а також «submit» та «validation» у близьких областях простору.

Розрахунок косинусної подібності для векторів цих речень V_A та V_B дає результат:

$$Sim(V_A, V_B) = \frac{V_A * V_B}{||V_A|| * ||V_B||} \approx 0,82 \quad (2)$$

Отримане значення 0.82 свідчить про високу семантичну спорідненість, що дозволяє алгоритму автоматично класифікувати ці повідомлення як дублікати однієї помилки, незважаючи на різну лексику [2]. Алгоритмічна реалізація методу включає етапи попередньої обробки (токенізація, лематизація), генерації ембедінгів та індексації векторів. Експериментальні дослідження показали, що використання контекстно-залежних векторних моделей дозволяє підвищити точність (Precision) пошуку релевантної документації на 15-20% у порівнянні з методами TF-IDF. Застосування відображеного підходу є перспективним для задач автоматизованої трасуальності вимог (Requirements Tracing), кластеризації дефектів та

інтелектуального пошуку в базах знань (Knowledge Bases) програмних продуктів.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Mikolov T., Chen K., Corrado G., Dean J. Efficient estimation of word representations in vector space. Proceedings of the International Conference on Learning Representations (ICLR). Scottsdale, 2021. URL: <https://arxiv.org/abs/1301.3781>.

2. Reimers N., Gurevych I. Sentence-BERT: Sentence Embeddings using Siamese BERT-Networks. Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing. Hong Kong, Association for Computational Linguistics, 2021. P. 3982-3992. in vector space.

УДК 004.65

Трезубова І. А., Вайцеховський О. С.

*Державний університет інтелектуальних технологій і зв'язку,
м. Одеса, Україна*

СТРАТЕГІЇ СТРИМАННЯ ВИСОКОГО НАВАНТАЖЕННЯ БАЗ ДАНИХ ВЕБ-СЕРВІСІВ

Інформаційні технології та інтернет займає велику частину у повсякденному житті багатьох людей, що призвело до розширення аудиторії користувачів ставлять жорсткі вимоги до надійності та продуктивності систем управління базами даних. База даних виступає фундаментальним компонентом будь-якої інформаційної системи, а забезпечення її масштабування за умов інтенсивного навантаження є одним з найактуальніших викликів у сучасній архітектурі веб-сервісів.

Вертикальне масштабування це просте рішення для збільшення потужності сервера через розширення ЦП та ОЗУ, зберігаючи інфраструктуру. Однак воно обмежене фізично та економічно, роблячи систему вразливою до відмов. Це робить його недостатнім для високодоступних веб-додатків. Горизонтальне ж масштабування розподіляє навантаження по кількох машинах.

Реплікація це ключовий крок у горизонтальному масштабуванні БД, дозволяючи створювати копії для розділення читання та запису. Веб сервіси переважно виконують операції читання, тож репліки знижують

навантаження на основний сервер. При суворій узгодженості, такі як фінансові системи уникають або синхронізують копії.

Шардування є технікою оптимізації, яка полягає у фізичному розрізанні вихідної бази даних на кілька менших, незалежних фрагментів (шардів). Це забезпечує необмежене горизонтальне масштабування, яку дозволяє паралельно виконувати операції різних користувачів. Існують такі варіанти:

Шардування за Діапазоном, розподіл даних відповідно до певного діапазону значень.

Шардування за Хешем, використання хеш-функції для рівномірного розподілу ключів між вузлами.

Зниження безпосереднього навантаження на основні бази даних досягається за допомогою двох високопродуктивних архітектурних рішень: кешування та поліглотної персистенції.

Кешування є найбільш ефективним методом мінімізації навантаження, оскільки воно базується на принципі зберігання часто використовуваних даних у швидше доступному місці, існують такі стратегії:

Бекенд-Кешування (In-Memory Stores): Використання високошвидкісних сховищ, таких як Redis або Memcached, які використовуються для зберігання об'єктів, сесій або попередньо обчислених результатів.

CDN (Content Delivery Network): Використовується для кешування статичних ресурсів.

Кешування на Рівні Застосунку/Клієнта: Використання спеціалізованих бібліотек, для кешування даних на стороні фронтенду, що знижує загальну кількість запитів до API та бекенду.

Ефективність кешування нерозривно пов'язана зі стратегією інвалідації, коли дані в БД оновлюються, але кеш ще не скинутий, підкреслює перевагу архітектур, що допускають тимчасову неузгодженість - BASE-моделей.

Поліглотна персистенція визначається як використання кількох технологій зберігання, обраних відповідно до потреб застосунку. Цей підхід відходить від монолітного використання RDBMS, розвантажуючи його, оскільки NoSQL-бази (Key-Value, Document, Graph) ефективніші для нереляційних завдань. Це дозволяє "вибрати правильний інструмент для роботи", мінімізуючи витрати на розробку.

Горизонтальне масштабування в умовах мережевої комунікації підпорядковане CAP-теоремі (Узгодженість (C), Доступність (A), Стійкість до Розподілення Мережі (P)). Оскільки P є обов'язковим, архітектори обирають між ACID (сувора C, обмежена A) та BASE

(висока А, кінцева С). У високомасштабованих соціальних мережах, архітектура компенсує це агресивним багаторівневим кешуванням.

У високодоступних системах із мікросервісною архітектурою підтримка узгодженості даних є складною проблемою. Для атомарної фіксації історично використовувався протокол Двофазної Фіксації (Two-Phase Commit, 2PC). Проте 2PC є блокувальним протоколом, який знижує доступність і перешкоджає незалежному масштабуванню. Більшість систем відмовляються від 2PC на користь моделей BASE.

Для забезпечення надійності та незалежного масштабування мікросервісів, що гарантують кінцеву узгодженість (BASE):

Stateless Microservices: Архітектура будується на незалежних, безстатусових мікросервісах. Це дозволяє автономно розробляти та впроваджувати функції, мінімізуючи ризик збою усієї системи.

Transactional Outbox: Цей шаблон проектування гарантує, що зміна в базі даних фіксується атомарно разом із повідомленням про цю зміну, що дозволяє іншим мікросервісам асинхронно оновити свій стан.

Перехід до BASE та асинхронних патернів є необхідним наслідком вимоги високої доступності та стійкості до розподілення в умовах мережевої комунікації.

Жодні архітектурні зміни не можуть повністю компенсувати неефективну роботу бази даних без коректної програмної оптимізації. Програмна оптимізація фокусується на аналізі виконання та ефективному використанні індексів. Індеси пришвидшують пошук даних у таблиці, завдяки виділення унікальних значень. Однак це додає витрати на кожну операцію вставки та оновлення, що сповільняє БД. Також розбиття складних запитів на менші, атомарні шматки, які можна виконувати послідовно або паралельно.

Ефективність оптимізації запитів доведена практичними прикладами компаній, таких як Netflix. Після міграції своєї реляційної інфраструктури на Amazon Aurora та оптимізації, досягли 50%–75% зниження середньої затримки в критичних додатках [4]. Це підтверджує, що горизонтальне масштабування повинне бути підкріплено програмною оптимізацією.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Fowler M. NoSQL: A Brief Guide to the Emerging World of Polyglot Persistence. 2012. [Електронний ресурс] – Режим доступу: <https://www.bigdata.ir/wp-content/uploads/2016/08/6C594096C6E33AE41A6D365F2C4588C6.pdf>
2. Netflix consolidates relational database infrastructure on Amazon Aurora. AWS Blog, 2023. [Електронний ресурс] – Режим доступу:

<https://aws.amazon.com/ru/blogs/database/netflix-consolidates-relational-database-infrastructure-on-amazon-aurora-achieving-up-to-75-improved-performance/>

3. The Four Meta Secrets of Scaling at Facebook. Highscalability.com, 2017. [Електронний ресурс] – Режим доступу: <https://highscalability.com/the-four-meta-secrets-of-scaling-at-facebook/>

4. Kleppmann M. (2017). “Designing Data-Intensive Applications” [Електронний ресурс] – Режим доступу: <https://www.oreilly.com/library/view/designing-data-intensive-applications/9781491903063/>

УДК 004.45

Шишкін М.С.

*Харківський національний економічний
університет імені Семена Кузнеця, м. Харків, Україна*

АГЕНТНИЙ ШІ В ЖИТТЄВОМУ ЦИКЛІ РОЗРОБКИ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ: ФАЗА РОЗРОБКИ

Вступ.

Фаза розробки - життєвого циклу розробки програмного забезпечення (Software development lifecycle, SDLC) традиційно зосереджена на ручному кодуванні, експертних оцінках та ітеративному налагодженні, що виконується розробниками-людьми за підтримки статичних інструментів. Поява агентного штучного інтелекту привела до зміни парадигми в цьому процесі, дозволивши системам штучного інтелекту діяти автономно, досягати цілей, координувати підзадачі та адаптувати рішення на основі зворотного зв'язку з навколишнього середовища. На відміну від традиційних інструментів автодоповнення коду, які працюють як пасивні помічники, агентні системи штучного інтелекту характеризуються орієнтацією на цілі, контекстним мисленням, делегуванням завдань та здатністю самостійно виконувати багатоетапні робочі процеси. Це перетворює фазу розробки з виключно людської діяльності на середовище для співпраці, де розробники-люди контролюють цілі високого рівня, тоді як агенти штучного інтелекту обробляють завдання рівня виконання.

Застосування агентного штучного інтелекту у фазі розробки.

У щоденній діяльності з розробки програмного забезпечення агентні системи штучного інтелекту працюють як автономні співробітники, здатні перетворювати абстрактні вимоги на робочі компоненти

програмного забезпечення. Після отримання історій користувачів або технічних специфікацій ці агенти виконують семантичний аналіз для інтерпретації бізнес-намірів, визначення функціональних меж та розкладання високорівневих цілей на структуровані завдання розробки. На основі цього аналізу агенти генерують початкові скелети коду, визначають інтерфейси модулів та створюють заглушки сервісів, які відповідають вибраним архітектурним шаблонам, таким як мікросервіси або багаторівневі дизайни додатків. Контракти API визначаються програмно, що забезпечує узгодженість між визначеннями сервісів, моделями запитів та схемами відповідей, тоді як залежності від сторонніх розробників вибираються та інтегруються відповідно до міркувань сумісності та безпеки.

Агентний ШІ може орієнтуватися та розуміти складні існуючі кодові бази, щоб забезпечити відповідність нових змін встановленим стандартам розробки та конвенціям проекту. Агенти аналізують історію репозиторію, рекомендації щодо кодування та архітектурну документацію, щоб відповідати правилам іменування, стандартам форматування та межах модулів. Під час впровадження нових функцій автономні агенти здатні знаходити відповідні точки розширення в кодовій базі, впроваджувати логіку без порушення принципів інкапсуляції та виконувати локалізований рефакторинг, коли виявляється надвисока складність коду або шаблони дублювання. Цей безперервний процес рефакторингу покращує зручність обслуговування та підтримує зменшення технічного боргу одночасно з розгортанням функцій, а не відкладає такі зусилля на фазі після релізу [1].

Автоматизоване тестування та інтеграція безпеки.

Робочий процес тестування також тісно інтегрується з агентною підтримкою розробки. Агенти ШІ динамічно генерують модульні та інтеграційні тести, узгоджені з новим кодом, і можуть автоматично виконувати ці тести в конвеєрах розробки. Коли з'являються збої тестування або попередження статичного аналізу, агенти інтерпретують журнали, виявляють потенційні першопричини та ітеративно змінюють код, доки не будуть виконані критерії стабільності. Паралельно агенти, орієнтовані на безпеку, сканують вихідні файли на наявність вразливостей, таких як недоліки ін'єкцій, небезпечні залежності або неправильні конфігурації, та пропонують безпечні шаблони кодування, перш ніж недоліки поширяться в нижчі середовища. Цей цикл зворотного зв'язку в режимі реального часу забезпечує «безпеку зсуву вліво», де структурні проблеми та проблеми відповідності вирішуються безпосередньо на етапі розробки, а не після випуску.

Вплив на продуктивність та співпрацю.

Співпраця агентів забезпечує вимірне підвищення продуктивності. Розробники витрачають менше часу на повторювані завдання кодування та генерацію шаблонів і більше часу на моделювання рішень, проектування бізнес-логіки та перевірку архітектурних рішень. Цикли розробки функцій скорочуються, оскільки агенти ШІ працюють безперервно без перерв, а якість коду покращується завдяки вбудованому забезпеченню узгодженості та автоматизації тестування. Розробники-люди все частіше виконують роль супервайзерів, які керують стратегічними намірами, перевіряють результати та вирішують складні граничні випадки, що вимагають обґрунтування предметної області, що виходить за межі машинної інтерпретації.

Ризики та виклики управління.

Незважаючи на суттєві переваги в продуктивності та якості, що з'являються завдяки агентному ШІ на етапі розробки SDLC, його впровадження також створює значні технічні, організаційні та етичні проблеми, які вимагають ретельного управління. Одним з центральних ризиків є тенденція до надмірного делегування, коли команди розробників все більше покладаються на автономних агентів для прийняття рішень щодо проектування та логіки впровадження. Хоча таке делегування прискорює розробку, воно може поступово зменшувати безпосередню взаємодію розробників із семантикою програмного забезпечення, архітектурним мисленням та низькорівневими практиками налагодження. З часом ця залежність може підірвати важливі інженерні компетенції, роблячи команди менш здатними критично перевіряти рішення, створені ШІ, або ефективно реагувати на складні збої, які перевищують можливості автоматизованого мислення.

Складність управління зростає, оскільки відповідальність за розробку розподіляється між людьми та автономними агентами. Традиційні моделі відповідальності припускають простежуване людське авторство для інженерних рішень, тоді як агентна розробка впроваджує частково непрозорі ланцюжки рішень, керовані оперативною інженерією, внутрішніми міркуваннями та ймовірнісними механізмами генерації. Визначення відповідальності за дефекти, вразливості безпеки або невідповідність нормативним вимогам стає складнішим, особливо в середовищах, що підлягають суворим вимогам до аудиту або сертифікації. Ця проблема вимагає впровадження систем управління, які забезпечують простежуваність коду, згенерованого штучним інтелектом, збереження підказок генерації, відстеження версій моделі та документовані контрольні точки перевірки для забезпечення

перевіреного нагляду. Ці обмеження підкреслюють необхідність структурованих процесів перевірки та формальних робочих процесів перевірки для підтримки підзвітності та довіри[2].

Висновок.

На завершення, агентний ШІ фундаментально змінює фазу розробки SDLC, переміщуючи робочі процеси, що вимагають виконання, до автономних систем, одночасно піднімаючи розробників-людей до супервайзерів та архітектурних ролей. Найефективніша модель впровадження поєднує агентну автономію з людським управлінням, гарантуючи, що підвищення продуктивності збалансовано забезпеченням якості, гарантіями безпеки та етичною відповідальністю. У міру розвитку агентних фреймворків та їхньої глибшої інтеграції в середовища розробки очікується, що вони пришвидшать розробку, підвищать надійність коду та переосмислять практики спільної розробки програмного забезпечення в хмарних та корпоративних екосистемах.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Russell S., Norvig P. Artificial Intelligence: A Modern Approach. – 4-те вид. – Hoboken : Pearson, 2021. – 1152 с.
2. NIST. Artificial Intelligence Risk Management Framework (AI RMF 1.0). – Gaithersburg : National Institute of Standards and Technology, 2023. – 128 с.

УДК 621.398:004.71/.72

Рудий Р.О., Воробець Г.І.

*Чернівецький національний університет
імені Юрія Федьковича, м. Чернівці, Україна*

ІСРАРХІЧНА МОДЕЛЬ ЗАХИСТУ ПЕРСОНІФІКОВАНИХ ДАНИХ В ТЕЛЕКОМУНІКАЦІЙНИХ МЕРЕЖАХ ТА ЇЇ ПРИКЛАДНА РЕАЛІЗАЦІЯ

Сучасний розвиток інформаційно-комунікаційних технологій (ІКТ), значне збільшення кількості і підвищення якості сервісів і пристроїв для обміну даними як в побуті, так і для виробничих задач потребує вирішення питань захисту персоніфікованої інформації користувачів послугами ІКТ. При цьому захисту потребують як персональні дані, що стосуються паролів доступу до різних сервісів, медичних даних конкретної особи, приватна інформація про членів сім'ї, майнові

питання, тощо, так і службова інформація, що стосується професійної діяльності окремих працівників, колективів, організацій. Більша онлайн присутність і активність користувачів у мережах зумовлює зростання кіберзагроз для персоналізованої інформації цих же користувачів.

Наразі існують різні підходи щодо захисту таких даних на програмному та апаратному рівнях в телекомунікаційних і комп'ютерних мережах. Можна розглядати і певну ієрархічну модель щодо застосування необхідного рівня захисту різних даних враховуючи їх важливість, вимоги до конфіденційності та використовуваних технічних програмних та апаратних засобів зберігання і захисту персоналізованої інформації.

Для зберігання персональних даних часто використовуються програми, які зберігають дані на пристрої користувача, сервіси, на яких дані зберігаються віддалено, або спеціальні пристрої.

Метою даних досліджень стала розробка уніфікованого підходу для реалізації захисту даних на основі ієрархічної моделі вимог до конфіденційності інформації, а також демонстрація такого підходу шляхом створення і випробування апаратно-програмного рішення для зберігання облікових записів на основі платформи Arduino Uno.

В роботі використано сучасні методи і технології зберігання даних в сховищах з різним рівнем доступу та моделі опису взаємодії технічних пристроїв і прикладних сервісних програм.

Об'єктом та предметом дослідження є принципи побудови пристроїв для зберігання цифрової інформації та створення пристрою для зберігання облікових записів на платформі Arduino Uno.

Для розробки ієрархічної моделі і критеріїв важливості захисту даних основою стали різні методи доступу до захищених даних. Зокрема розглядаються системи з прямим та опосередкованим доступом, серверний та мережевий доступ у вигляді мережевих систем зберігання даних (NAS) і мереж зберігання даних (SAN) [1-5].

Для практичної демонстрації запропонованого підходу застосовано спеціалізовану платформу на основі малогабаритного мікроконтролера [6-8], який дозволяє з'єднуватися з пристроєм користувача та обробляє його запити (рис. 1). Програмну реалізацію системи виконано із застосуванням платформи Windows Forms у вигляді користувацького інтерфейсу (рис. 2).

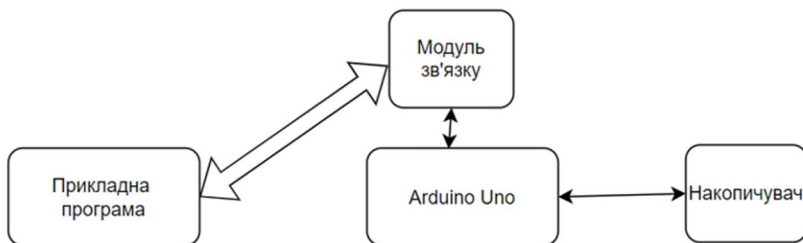


Рисунок 1 – Модель пристрою для зберігання конфіденційних записів даних

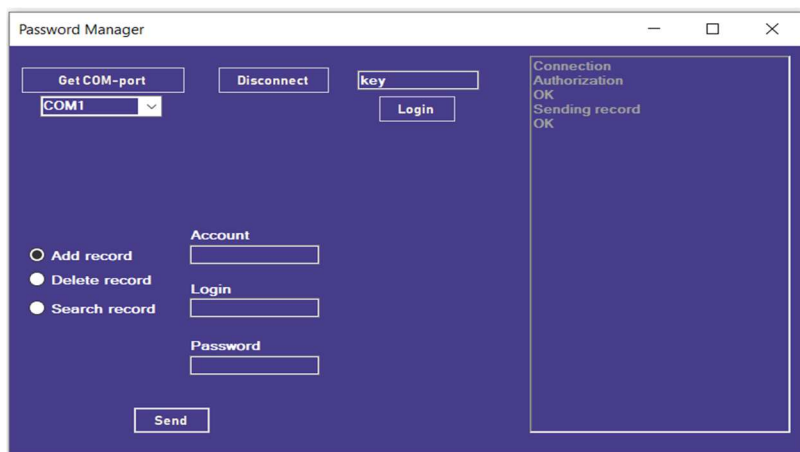


Рисунок 2 – Віконний інтерфейс користувача для контролю запису і рівня захисту даних.

В залежності від критеріїв важливості та конфіденційності інформації різним даним присвоюється спеціальний індекс рівня доступу та визначається ієрархічний рівень відповідності середовища для їх зберігання і додаткового кодування/шифрування [9]. Дані завжди зберігаються на певному типі носія чи в певному середовищі (SSD, флеш-пам'ять; гібридне сховище; хмарні сховища; гібридне хмарне сховище), тому доступ до них можливий тільки користувачеві.

Програма складається з набору файлів, що генеруються у Windows Forms. Основними є файли Program.cs, Form1.cs та StringChipher.cs, які власне забезпечують створення відповідної форми, запити на доступ до запису/зчитування/зберігання даних та обміни ними між користувачем та відповідними середовищами.

Наукова новизна проведених досліджень і запропонованих рішень полягає в розробці уніфікованої інформаційної ієрархічної моделі та критеріїв захисту персоніфікованих та конфіденційних даних в телекомунікаційних системах, створенні та апробації апаратно-програмного рішення на платформі Arduino Uno, що дозволяє зберігати інформацію на віддалених пристроях чи середовищах, розробці системи захисту даних на пристрої у вигляді авторизації користувача та шифрування даних, реалізації способу комунікації між пристроєм та програмою користувача.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Damjanovski V. CCTV Networking and Digital Technology / Vlado Damjanovski. – Sydney: Elsevier, 2013. – 614 с.
2. Tate J. Introduction to Storage Area Networks / J. Tate, P. Beck, L. Miklas: RedBooks, 2017. – 280 с.
3. NAS vs SAN. URL: <https://hypertecdirect.com/knowledge-base/nas-vs-san/>.
4. What is data storage? URL: <https://www.ibm.com/topics/data-storage>
5. Memory. URL: <https://www.arduino.cc/en/Tutorial/Foundations/Memory>
6. Overview of the Arduino UNO Components. URL: <https://docs.arduino.cc/tutorials/uno-rev3/intro-to-board>
7. Bayle J. C Programming for Arduino / Julien Bayle. – Birmingham: Packt Publishing, 2013. – 488 с.
8. SD. URL: <https://www.arduino.cc/reference/en/libraries/sd/>
9. Paar C. Understanding Cryptography / C. Paar, J. Pelzl. – Bochum: Springer, 2013. – 372 с.

ЗМІСТ

ІНТЕЛЕКТУАЛЬНІ ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ

<i>Антипова П. Б., Болюбаи Н. М.</i> АНАЛІЗ ДИНАМІКИ ТА СТАБІЛЬНОСТІ ЕЛЕКТРОННОГО ВРЯДУВАННЯ КРАЇН СВІТУ ІЗ ВИКОРИСТАННЯ АЛГОРИТМУ FURRY C-MEANS.....	4
<i>Артим В. І., Козлов О. В.</i> ІНТЕЛЕКТУАЛЬНА СИСТЕМА БАГАТОКАНАЛЬНОЇ ПІДТРИМКИ КЛІЄНТІВ У СФЕРІ ПОДОРОЖЕЙ.....	7
<i>Банар А. Ю., Воробець Г. І.</i> ПОБУДОВА ТА УПРАВЛІННЯ SDN МЕРЕЖ З ВИКОРИСТАННЯМ МЕТОДІВ ШТУЧНОГО ІНТЕЛЕКТУ.....	10
<i>Безощук С. С., Калініна І. О.</i> ІНТЕЛЕКТУАЛЬНА СИСТЕМА МОНІТОРИНГУ ПАЛИВНО-МАСТИЛЬНИХ МАТЕРІАЛІВ.....	12
<i>Горшколепов І. В., Сіденко Є. В.</i> ПРОГНОЗУВАННЯ ПОТРЕБ У РЕСУРСАХ ТА ОПТИМІЗАЦІЯ ЛОГІСТИЧНИХ МАРШРУТІВ НА ОСНОВІ НЕЙРОННИХ МЕРЕЖ.....	15
<i>Денисенко Н. В.</i> МОДЕЛЮВАННЯ ІНТЕЛЕКТУАЛЬНОЇ СИСТЕМИ КЕРУВАННЯ МІКРОКЛІМАТОМ ТЕПЛИЧНОГО ГОСПОДАРСТВА	18
<i>Кисса А. Ф., Гунченко Ю. О.</i> АНАЛІЗ ТА ПОРІВНЯННЯ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ РОЗПІЗНАВАННЯ ДОРОЖНІХ ЗНАКІВ.....	20
<i>Кісельов Б. Г.</i> МЕТОДИ КЛАСИФІКАЦІЇ РУДОВИХ МАТЕРІАЛІВ ЗА СЕНСОРНИМИ ПОКАЗНИКАМИ ПРИ НИЗЬКІЙ КОНТРАСТНОСТІ ДАНИХ.....	23
<i>Кулішова Є. С., Калініна І. О.</i> ІНТЕЛЕКТУАЛЬНА СИСТЕМА КЛАСИФІКАЦІЇ ЕКОЛОГІЧНИХ ПОКАЗНИКІВ З УРАХУВАННЯМ ДИСБАЛАНСУ КЛАСІВ.....	26
<i>Пяткіна Є. П., Сіденко Є. В.</i> ІНФОРМАЦІЙНА СИСТЕМА РОЗПІЗНАВАННЯ ЕМОЦІЙ З ВИКОРИСТАННЯМ НЕЙРОННОЇ МЕРЕЖІ.....	30

Романов А. В., Гожий В. О. ІНФОРМАЦІЙНА СИСТЕМА АНАЛІЗУ ТЕНДЕНЦІЙ ЗМІНИ КЛІМАТУ ЗА СУПУТНИКОВИМИ ДАНИМИ.....	33
Скиба О. М., Гожий О. П. ІНТЕЛЕКТУАЛЬНА СИСТЕМА ГЕНЕРАЦІЇ ЗОБРАЖЕНЬ ПО ТЕКСТОВОМУ ОПИСУ.....	36
Стипаненко С. В., Кондратенко Г. В., Сіденко Є. В. СИСТЕМА ІНТЕРПРЕТАЦІЇ МЕДИЧНИХ ДІАГНОСТИЧНИХ МОДЕЛЕЙ НА ОСНОВІ МЕТОДІВ ХАІ.....	40
Фоменко О. В., Козлов О. В. ІНТЕЛЕКТУАЛЬНА СИСТЕМА КЕРУВАННЯ ДОСТАВКОЮ ВАНТАЖІВ В ДИНАМІЧНОМУ СЕРЕДОВИЩІ ЗА ДОПОМОГОЮ БПЛА НА БАЗІ UNITY.....	43
Чупина В. Є., Сіденко Є. В. ІНТЕЛЕКТУАЛЬНА СИСТЕМА КАРТОГРАФУВАННЯ ВИБУХОНЕБЕЗПЕЧНИХ ПРЕДМЕТІВ НА ОСНОВІ ДАНИХ БПЛА ТА ГЛИБИННОГО НАВЧАННЯ.....	46
Шавріна І. О., Кулаковська І. В. ІНФОРМАЦІЙНА СИСТЕМА З ОЦІНКИ ГОТОВНОСТІ СТАРТАПУ ДО МАСШТАБУВАННЯ.....	49
Shvets Y. O., Malakhov E. V. DYNAMIC RECONFIGURATION AND TASK REDISTRIBUTION TECHNOLOGY.....	52
Шиян С. І. ОСОБЛИВОСТІ ПОБУДОВИ ІНТЕЛЕКТУАЛЬНОЇ СИСТЕМИ КОНТРОЛЮ ПАРАМЕТРІВ ПММ.....	54
Яслик О. М. ІНТЕЛЕКТУАЛЬНА СИСТЕМА ПЕРВИННОЇ ДІАГНОСТИКИ COVID- 19 НА ОСНОВІ NAIVE BAYES.....	57
Девятеріков Р. Г., Мельничук С. В., Воробець Г. І. ОСОБЛИВОСТІ АНАЛІЗУ І СИНТЕЗУ АДАПТИВНОЇ ПОДІЄВО-ОРІЄНТОВАНОЇ SDN- МОДЕЛІ ПІДВИЩЕННЯ ПРОДУКТИВНОСТІ ТА ВІДМОВОСТІЙКОСТІ ІОТ-СИСТЕМ.....	59

МАШИННЕ НАВЧАННЯ ТА ШТУЧНИЙ ІНТЕЛЕКТ

Воїнова О. В., Гунченко Ю. О. МОДЕЛЮВАННЯ НЕЙРОМЕРЕЖЕВОЇ СИСТЕМИ РОЗПІЗНАВАННЯ ТА ПЕРЕКЛАДУ МОВИ ЖЕСТІВ.....	63
Гожий В. О. ОСОБЛИВОСТІ ВПРОВАДЖЕННЯ ТЕХНОЛОГІЙ ШТУЧНОГО ІНТЕЛЕКТУ В МЕНЕДЖМЕНТІ ІТ ПРОЄКТІВ.....	66

Гуркліс І. В., Петрушина Т. І. ІНТЕГРАЦІЯ НЕЙРОМЕРЕЖЕВИХ КОМПОНЕНТІВ У БАГАТОШАРОВІ АРХІТЕКТУРИ ТРАНСФОРМАЦІЇ ДАНИХ.....	69
Димо В. В., Гожий О. П. ДОСЛІДЖЕННЯ ЗГОРТКОВИХ НЕЙРОННИХ МЕРЕЖ ЗА ДОПОМОГОЮ МЕТОДІВ ПОЯСНЮВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ.....	72
Зачиняєв П. Д., Калініна І. О. ІНТЕЛЕКТУАЛЬНА СИСТЕМА ПРОГНОЗУВАННЯ ВАРТОСТІ КРИПТОВАЛЮТ.....	75
Дирда І. Ю., Сіденко Є. В. ІНТЕЛЕКТУАЛЬНА СИСТЕМА КЛАСИФІКАЦІЇ РАКОВИХ КЛІТИН ЗА ГІСТОЛОГІЧНИМИ ЗОБРАЖЕННЯМИ НА ОСНОВІ МЕТОДІВ МАШИННОГО НАВЧАННЯ.....	78
Кіпер С. Ю. МЕТОДИ ВИЯВЛЕННЯ ТА РОЗПІЗНАВАННЯ ОБЛИЧ: СУЧАСНИЙ СТАН, ВИКЛИКИ ТА ПЕРСПЕКТИВИ.....	81
Ковалів О. П., Сіденко Є. В., Кондратенко Г. В. ОПТИМІЗАЦІЯ ПРОГНОЗУВАННЯ ІНФЕКЦІЙНИХ ЗАХВОРЮВАНЬ ЗА ДОПОМОГОЮ НЕЙРОННИХ МЕРЕЖ З ВИКОРИСТАННЯМ ЕКЗОГЕННИХ ЗМІННИХ	84
Рябов Д. А, Пенко В. Г. ВИКОРИСТАННЯ НАВЧАННЯ З ПІДКРІПЛЕННЯМ ДЛЯ НАВІГАЦІЇ ДРОНІВ.....	88
Рябов Д. М., Пенко В. Г. МЕТОД ПІДВИЩЕННЯ ЕФЕКТИВНОСТІ НАВЧАННЯ З ПІДКРІПЛЕННЯМ ЗАВДЯКИ КЛАСТЕРИЗАЦІЇ ПРОСТОРУ СТАНІВ.....	90
Трунов О. М., Савчук А. О. ПОРІВНЯЛЬНИЙ АНАЛІЗ НЕЙРОМЕРЕЖЕВИХ МЕТОДІВ СЕГМЕНТАЦІЇ ЗОБРАЖЕНЬ ІНТЕРФЕЙСІВ АСК ДЛЯ АВТОМАТИЗОВАНОГО КОГНІТИВНОГО АНАЛІЗУ.....	92
Салютін М. О., Кондратенко Ю. П., Сіденко Є. В. МЕТОДИ UNDER-TA OVERSAMPLING У ЗАДАЧІ РОЗПІЗНАВАННЯ ВІЙСЬКОВИХ ОБ'ЄКТІВ.....	97
Самойленко М. М., Гожий О. П. ІНТЕЛЕКТУАЛЬНА СИСТЕМА КЛАСИФІКАЦІЇ ПАЛЬНОГО НА ОСНОВІ МЕТОДІВ МАШИННОГО НАВЧАННЯ.....	101
Стрюк О. С. РОЗГОРТАННЯ ТА ІНЖЕНЕРНА ОПТИМІЗАЦІЯ ГЕНЕРАТИВНИХ ЗМАГАЛЬНИХ МЕРЕЖ НА КОРДОННИХ СИСТЕМАХ.....	104

Хіньов Д. О., Кулаковська І. В. ІНТЕЛЕКТУАЛЬНА СИСТЕМА КЛАСИФІКАЦІЇ МУЗИЧНИХ ЖАНРІВ НА ОСНОВІ СПЕКТРОГРАМ ІЗ ВИКОРИСТАННЯМ ЗГОРТКОВИХ НЕЙРОННИХ МЕРЕЖ.....106

Цимбал А. Ю., Калініна І. О. ІНТЕЛЕКТУАЛЬНА СИСТЕМА ІДЕНТИФІКАЦІЇ ФІШИНГУ.....109

Чернова О. Ю., Антоненко О. С. ОСОБЛИВОСТІ ПРЕДСТАВЛЕННЯ ЧАСОВИХ РЯДІВ ДЛЯ МОДЕЛІ МАШИННОГО НАВЧАННЯ В ПРОГНОЗУВАННІ ПОГОДИ.....112

Шеремет К. О., Сіденко Є. В. ГЕНЕРАЦІЯ СИНТЕТИЧНИХ РЕНТГЕНІВСЬКИХ ЗНІМКІВ НА ОСНОВІ ОПИСІВ ІЗ ВИКОРИСТАННЯМ НЕЙРОННИХ МЕРЕЖ.....115

ІНТЕЛЕКТУАЛЬНІ КОМП'ЮТЕРНО-ІНТЕГРОВАНІ ТЕХНОЛОГІЇ

Василенко В. І., Пенко В. Г. ПІДХІД ДО ЕФЕКТИВНОГО ВИКОРИСТАННЯ РЕСУРСІВ ЗА ДОПОМОГОЮ НАВЧАННЯ З ПІДКРІПЛЕННЯМ.....119

Кондрашов Д. А. Музика І. О. ІНТЕЛЕКТУАЛЬНЕ КЕРУВАННЯ ПРОЦЕСОМ ПРЕЦИЗІЙНОГО СОРТУВАННЯ РУД НА ОСНОВІ МАШИННОГО НАВЧАННЯ.....120

Купін А. І., Косей М. П. МАТЕМАТИЧНА МОДЕЛЬ РОЙОВОГО ІНТЕЛЕКТУ ДЛЯ КАРТОГРАФУВАННЯ ТЕРИТОРІЇ ЗА ДОПОМОГОЮ БПЛА.....124

Моторнюк С. О. ПРОБЛЕМАТИКА ДОСТОВІРНОСТІ ТА ПОВНОТИ ВІДКРИТИХ ДАНИХ: ВИКЛИКИ ТА НАПРЯМИ ІНТЕЛЕКТУАЛЬНОЇ ОБРОБКИ.....126

Романенко О. О., Купін А. І. КОНЦЕПЦІЯ ЦИФРОВОГО ДВІЙНИКА ДЛЯ СИСТЕМ ТЕХНІЧНОГО ОБСЛУГОВУВАННЯ.....131

Сбітнєв О. Ю., Волощук Л. О. ПОРІВНЯЛЬНИЙ АНАЛІЗ МЕТАЕВРИСТИЧНИХ АЛГОРИТМІВ ДЛЯ ЗАДАЧІ РОЗМІЩЕННЯ СЕРВІСІВ У ТУМАННИХ ОБЧИСЛЮВАЛЬНИХ СИСТЕМАХ.....133

Харламенко В. Ю., Тимохін Є. В. ОГЛЯД ПЕРСПЕКТИВНИХ МЕТОДІВ ПРОГНОЗНОГО КЕРУВАННЯ ВОДНИМИ РЕСУРСАМИ ЗБАГАЧУВАЛЬНОЇ ФАБРИКИ.....136

<i>Гиляка В. О., Кулаковська І. В.</i> ІНТЕЛЕКТУАЛЬНА СИСТЕМА СУПРОВОДУ ОБ'ЄКТА В 3D-СЕРЕДОВИЩІ НА БАЗІ UNITY ТА НЕЙРОННИХ МЕРЕЖ.....	139
---	-----

ІНТЕЛЕКТУАЛЬНІ ТЕХНОЛОГІЇ В ПРОГРАМНІЙ ТА КОМП'ЮТЕРНІЙ ІНЖЕНЕРІЇ

<i>Бондаренко С. В., Ковальчук М. В.</i> ВИКОРИСТАННЯ ВЕБТЕХНОЛОГІЙ ДЛЯ ВІЗУАЛІЗАЦІЇ ЛОГІСТИЧНИХ МЕТРИК ТА НАДАННЯ ПРОГНОСТИЧНОЇ АНАЛІТИКИ ЧЕРЕЗ API/ВЕБІНТЕРФЕЙСИ.....	143
---	-----

<i>Волочай П. О., Салтовський Б. Г.</i> ІОТ-ПРИСТРІЙ ДЛЯ КОНТРОЛЮ ВМІСТУ ВУГЛЕКИСЛОГО ГАЗУ В ПРИМІЩЕННІ.....	145
--	-----

<i>Головка О. В., Швед А. В.</i> ПРОГРАМНЕ ЗАБЕЗПЕЧЕННЯ 3D-ВІЗУАЛІЗАЦІЇ ДВОВИМІРНИХ ОБ'ЄКТІВ.....	148
---	-----

<i>Гюльмамедов Н. М., Бурлаченко І. С.</i> ВБУДОВАНА СИСТЕМА ДЛЯ ВИЗНАЧЕННЯ ПЕРЕШКОД НА БАЗІ RASPBERRY PI ZERO 2W ТА TENSORFLOW.....	150
--	-----

<i>Євстратьєв М. С., Сіденко Є. В.</i> ЗАСТОСУВАННЯ RAG-ТЕХНОЛОГІЇ В АГЕНТНИХ LLM-СИСТЕМАХ.....	153
---	-----

<i>Онацький В. В., Савінов В. Ю.</i> ВИКОРИСТАННЯ ХЕШ-ЯКОРІВ У SEGWIТ-ТРАНЗАКЦІЯХ БІТКОЇНА ДЛЯ НЕЗМІННОГО ЖУРНАЛЮВАННЯ ПОДІЙ ІОТ-ПРИСТРОЇВ.....	155
---	-----

<i>Пальчик Т. В., Голуб Т. В.</i> ПОРІВНЯННЯ РЕКОМЕНДАЦІЙНИХ ВЕБСЕРВІСІВ ПІДБОРУ ТОВАРІВНА ЦИФРОВИХ ПЛАТФОРМАХ.....	158
---	-----

<i>Раленко В. С., Давиденко Є. О.</i> ВИКОРИСТАННЯ ВЕКТОРНИХ МОДЕЛЕЙ ДЛЯ СЕМАНТИЧНОГО АНАЛІЗУ ТЕКСТУ.....	161
---	-----

<i>Трезубова І. А., Вайцеховський О. С.</i> СТРАТЕГІЇ СТРИМАННЯ ВИСОКОГО НАВАНТАЖЕННЯ БАЗ ДАНИХ ВЕБ- СЕРВІСІВ.....	163
--	-----

<i>Шишкін М. С.</i> АГЕНТНИЙ ШІ В ЖИТТЄВОМУ ЦИКЛІ РОЗРОБКИ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ: ФАЗА РОЗРОБКИ.....	166
---	-----

<i>Рудий Р. О., Воробець Г. І.</i> ІЄРАРХІЧНА МОДЕЛЬ ЗАХИСТУ ПЕРСОНІФІКОВАНИХ ДАНИХ В ТЕЛЕКОМУНІКАЦІЙНИХ МЕРЕЖАХ ТА ЇЇ ПРИКЛАДНА РЕАЛІЗАЦІЯ.....	169
--	-----

ДЛЯ НОТАТОК

ДЛЯ НОТАТОК

В авторській редакції
Комп'ютерна верстка *В. Шевченко*
Друк *С. Волинець*. Фальцювально-палітурні роботи *О. Мішалкіна*.

Підп. до друку 30.12.2025.
Формат 60х84 ¹/₁₆. Папір офсет.
Гарнітура «Times New Roman». Друк ризограф.
Ум. друк. арк. 10,5. Обл.-вид. арк. 8,2.
Тираж 10 пр. Зам. № 7156.

Видавець і виготовлювач: ЧНУ ім. Петра Могили.
54003, м. Миколаїв, вул. 68 Десантників, 10.
Тел.: 8 (0512) 50–03–32, 8 (0512) 76–55–81,
e-mail: rector@chmnu.edu.ua.
Свідоцтво суб'єкта видавничої справи ДК № 6124 від 05.04.2018.