



УКРАЇНА

**МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
НАЦІОНАЛЬНИЙ ТЕХНІЧНИЙ УНІВЕРСИТЕТ УКРАЇНИ
«КИЇВСЬКИЙ ПОЛІТЕХНІЧНИЙ ІНСТИТУТ імені ІГОРЯ СІКОРСЬКОГО»**

Навчально-науковий інститут
атомної та теплової енергетики

Кафедра
інженерії програмного забезпечення в енергетиці

**СУЧАСНІ АСПЕКТИ ІНЖЕНЕРІЇ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ
MODERN ASPECTS OF SOFTWARE ENGINEERING**

II МІЖНАРОДНА НАУКОВО-ПРАКТИЧНА КОНФЕРЕНЦІЯ

НАУКОВИЙ ЗБІРНИК



КИЇВ
2024

УДК 519.85+004.42

Сучасні аспекти інженерії програмного забезпечення = Modern aspects of Software Engineering : II Міжнар. наук.-практ. конф. : наук. зб. (Київ, 13 листопада 2024 р.). – Київ : КПІ ім. Ігоря Сікорського, 2024 р. – 90 с.

Матеріали науково-практичної конференції містять короткий зміст доповідей науково-дослідних робіт аспірантів, викладачів і студентів. Тези доповідей містять обговорення наукових результатів, презентації здобутків та обмін досвідом між науковцями, що проводять дослідження в галузі 121 Інженерія програмного забезпечення, а саме в області інструментарію технологій програмування та програмного забезпечення кіберфізичних систем.

Організаційний комітет

Голова організаційного комітету – Євген ГАВРИЛКО, доктор технічних наук, професор, професор кафедри інженерії програмного забезпечення в енергетиці Навчально-наукового інституту атомної та теплової енергетики

Валерій КУЗЬМІНИХ – кандидат технічних наук, доцент, доцент кафедри інженерії програмного забезпечення в енергетиці Навчально-наукового інституту атомної та теплової енергетики

Ірина ГУССВА – кандидат економічних наук, доцент кафедри інженерії програмного забезпечення в енергетиці Навчально-наукового інституту атомної та теплової енергетики

Максим САВЕЛЬЕВ – кандидат технічних наук, старший науковий співробітник, завідувач відділу інституту проблем безпеки АЕ НАН України

Шивей ЧЖУ – начальник інформаційного відділу Інституту інформаційних досліджень, Академія наук провінції Шаньдун (КНР)

Олена БАНДУРКА – PhD, старший викладач кафедри інженерії програмного забезпечення в енергетиці Навчально-наукового інституту атомної та теплової енергетики

Ксенія ОЛЕНЄВА, старший викладач кафедри інженерії програмного забезпечення в енергетиці Навчально-наукового інституту атомної та теплової енергетики

Руслан ТАРАНЕНКО, провідний інженер кафедри інженерії програмного забезпечення в енергетиці Навчально-наукового інституту атомної та теплової енергетики

Рецензенти:

Барабаш О.В., доктор технічних наук, професор, професор кафедри інженерії програмного забезпечення в енергетиці Навчально-наукового інституту атомної та теплової енергетики

Федорова Н.В., доктор технічних наук, доцент, професор кафедри інженерії програмного забезпечення в енергетиці Навчально-наукового інституту атомної та теплової енергетики

Електронне наукове видання

УДК 519.85+004.42

ISBN:

© КПІ імені Ігоря Сікорського, 2024
© Автори статей, 2024

ЗМІСТ

ПЛЕНАРНЕ ЗАСІДАННЯ

Варава І. А. Автоматизація створення додатків для Simulink-моделей.....	6
Шуклін Г. В. Методика аналізу інформації щодо кібернетичного стану енергетичного об'єкту.....	8
Свинчук О. В., Барабаш О. В. Методи забезпечення функціональної стійкості багатомодульних інформаційних систем.....	11
Пироговська Т. В., Мусієнко А. П. Моделювання сигналу дальнього поля на основі використання коефіцієнтів Чебишева.....	13
Олексій А. О., Верлань А. А. Покращений метод аналізу акустичних сигналів водного середовища на основі згорткової нейромережі SOP з застосуванням багатомасштабної згортки.....	16
Здор К. А., Шалденко О. В., Недашківський О. Л. Розбиття відео на сцени за допомогою візуальних трансформерів для відео.....	19

СЕКЦІЯ 1. Інструментарій технологій програмування

Макарчук А. В., Барабаш О. В. Застосування моделей машинного навчання для ідентифікації умов, що гарантують функціональну стійкість інформаційних систем.....	23
Царева О. С. Розподілені обчислення у контексті об'єктно-орієнтованого програмування: виклики, підходи та перспективи.....	25
Черноусов Д. І., Бандурка О. І. Засоби розробки програмного забезпечення з відкладеними періодичними задачами.....	27
Остапенко І. П. Отримання даних щодо використання ЦПУ процесом в операційних системах на основі Linux програмним чином.....	29
Гнатишин М. С., Недашківський О. Л. Архітектура програмного забезпечення для виявлення фейкової інформації в мережі Інтернет на основі нейронних мереж.....	31
Ситнік І. І., Шпурик В. В. Метод побудови імовірнісних моделей управління процесом розробки програмного забезпечення в умовах невизначеності.....	33
Логвиненко Т. С., Гагарін О. О. Оцінювання якості моделювання гідроакустичних сигналів процесу моніторингу у районах спостереження.....	35
Руй Д. І., Донець А. Г., Колумбет А. П. Аналіз інструментів створення веб-платформ підтримки екологічного менеджменту.....	37
Дембіцький В. В., Коваль О. В. Проблематика однопунктної технології ідентифікації джерел гідроакустичних сигналів.....	39
Фішер О. Є., Барабаш О. В. Методи оптимізації моделей комп'ютерного зору для розпізнавання зображень на автономних пристроях.....	41
Гейко О. О., Варава І. А. Інструменти представлення математичних формул в програмному забезпеченні.....	43
Романів Р. С., Бандурка О. І. Застосування блокчейн технології для забезпечення функціональної стійкості розподілених інформаційних систем моніторингу руху транспортних засобів.....	45
Ворвуль Д.М., Безсмертний В.Ю., Мусієнко А.П. Оцінка продуктивності готових рішень пошуково-доповненої генерації для надання рекомендації з вибору найкращого рішення....	47

СЕКЦІЯ 2. Програмне забезпечення кіберфізичних систем

Добровольський Р. В., Shiwei Zhu (Чжу Шівей), Сенченко В. Р., Коваль О. В. Програмне забезпечення перевірки семантичної сумісності онтологічних моделей для предметної області системи оцінки рівня наукової роботи викладачів кафедри.....	49
Єрмосін О. О., Gong Xiaodong (Гонг Сядун), Коваль О. В. Система побудови онтологічної моделі на основі сценаріїв АНД для використання в предметній області «оцінка рівня наукової роботи».....	51
Єрмосіна О. О., Gong Xiaodong (Гонг Сядун), Коваль О. В. Аналіз та виконання сценаріїв АНД з використанням онтологічної моделі для задачі «оцінка рівня наукової роботи».....	53
Чжан Мінцзюнь (Zhang Mingjun), Стативка Ю. І. Конвертація рейтингу неоднорідних об'єктів з рангової шкали у метричну.....	55
Voitko A. S., Kuzminyh V. O. Dynamic message routing strategy in microservice architectures for data processing.....	56
Taranenko R. A., Taranenko E. R. Knowledge model of data in the representation of information of intelligent analytical systems.....	58
Савко В. Я., Гаврилко Є. В. Розробка інтелектуального спеціального програмного забезпечення управління вентиляційних систем Конфайнменту на основі урахування інерційності системи.....	60
Максименко П. О., Гусєва І. І. Обмеження і виклики комп'ютерного зору в умовах низького освітлення при створенні карт для потреб транспортної сфери.....	63
Ярмак Д. О., Варава І. А. Засоби формування сценаріїв динамічних гідроакустичних експериментів.....	65
Лебедєв А. В., Федорова Н. В. Актуальність та проблеми використання мутаційного тестування для аналізу якості коду систем з машинним навчанням.....	67
Хоменко О. М., Коваль О. В. Підходи до аналізу каскадних ефектів в роботі електромереж.....	69
Ляшко І. І., Залевська О. В. Концептуальні основи та перспективи розробки програмного забезпечення для розпізнавання міміки у віртуальних середовищах на основі антропоморфних даних.....	72
Садовська І. І., Коваль О. В. Система розширення або злиття баз знань.....	75
Передера В. Р., Недашківський О. Л. Актуальність, проблеми та неточності існуючих методів класифікації та виявлення джерел інформації в Інтернеті на основі нейронних мереж.....	77
Євтушенко Д. М., Донець А. Г., Колумбет В. П. Застосування розширень браузерів Google Chrome та Mozilla Firefox для аналізу змісту веб-сторінок та формування статистики в інформаційній системі тестування та оцінювання знань.....	80
Клименко Я. В., Свинчук О. В. Методи і програмне забезпечення для моделювання процесів теплообміну і гідродинаміки в турбінних системах.....	83
Дащик А. С., Залевська О. В. Огляд програмного забезпечення для топологічної оптимізації.....	85



ПЛЕНАРНЕ ЗАСІДАННЯ

АВТОМАТИЗАЦІЯ СТВОРЕННЯ ДОДАТКІВ ДЛЯ SIMULINK-МОДЕЛЕЙ

Середовища MATLAB надає два інструменти побудови графічного інтерфейсу користувача: GUIDE та App Designer[1]. GUIDE наразі підтримується у MATLAB 2024a, але у майбутніх версіях MATLAB фірма MathWorks від нього відмовиться. В даний час існує можливість міграції додатків створених за допомогою GUIDE у формат App Designer. Як середовище розробки графічного інтерфейсу, GUIDE має невелику палітру основних елементів керування та інструментів (редактор коду, інспектор властивостей та переглядач об'єктів). Починаючи із версії 2019b з'явилася можливість створювати GUI за допомогою App Designer. App Designer має шаблони віконних додатків, додатків для Simulink-моделей та шаблони для кастомізації користувацьких компонентів. Крім більшої кількості основних елементів керування в App Designer доступні контейнери, меню, панельні індикатори та перемикачі, авіаційні прилади, компоненти для роботи із моделями Simulink.

Створення додатка для Simulink-моделі починається із її асоціації з графічним інтерфейсом.

Традиційний спосіб створення графічного інтерфейсу передбачає перетягування компонентів на форму та їх зв'язування із змінними Simulink-моделі.

При ручному створенні змінних для блоків необхідно виконувати наступні дії:

- викликати діалогове вікно «Block Parameters»;
- біля параметру натиснути на кнопку «Create Variable»;
- у вікні «Create New Data» ввести назву змінної.

За допомогою запропонованого розширення результат цих дій користувача буде формуватися програмним кодом.

Для автоматизації процесу створення змінних пропонується розширення панелі інструментів Simulink за рахунок розробленої закладки із кнопками аналізу моделі та створення змінних [2] (рисунок 1).

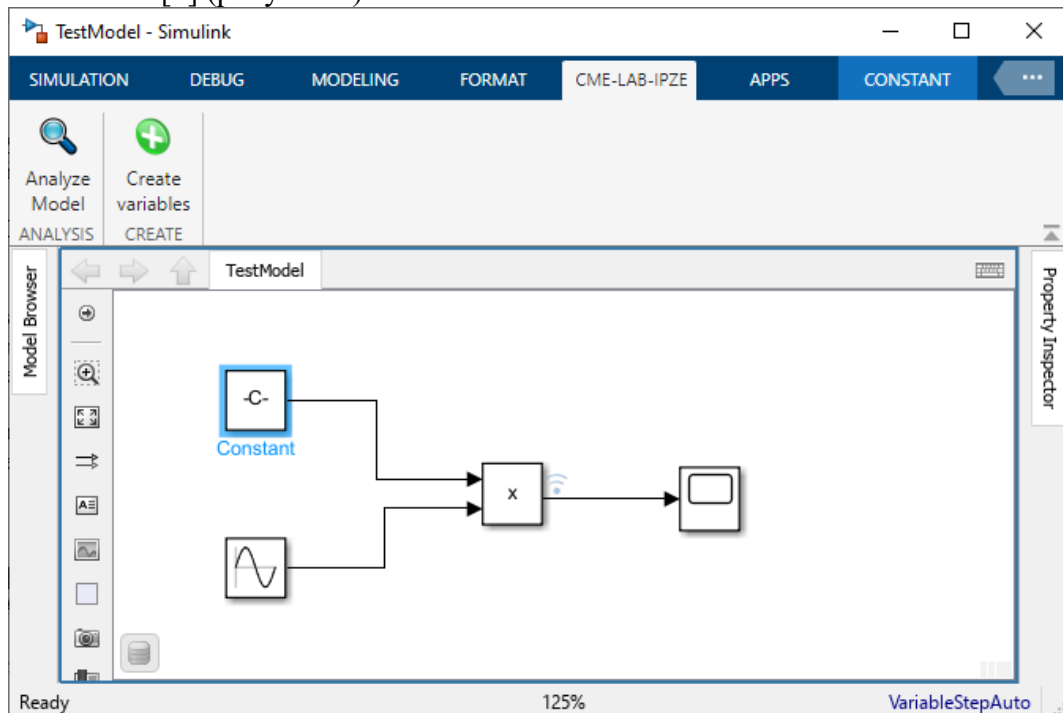


Рисунок 1 – Панель інструментів в Simulink

Зовнішній вигляд панелі конфігурується двома JSON-файлами в яких визначені піктограми та обробники, що представлені m-функціями.

Програмна зміна моделі передбачає ініціалізацію двох об'єктів:

- 1) об'єкт відображення коду для конфігурування даних і функціонального інтерфейсу для генерації коду мовою C;
- 2) робочий простір моделі.

За допомогою методу `setModelParameter` відбувається встановлення значення параметрів для моделювання. Ім'я для створюваної змінної формується на основі імені діалогового поля та суфікса `Var`. Метод `assignin` ініціалізує змінну конкретним значенням.

Процедура аналізу моделі передбачає виявлення всіх блоків та їх діалогових параметрів, що можуть бути використані в графічному інтерфейсі користувача. Стандартні діалоги `Block Parameters` мають надлишкові налаштування, що не пов'язані безпосередньо із параметризацією моделі.

Користувачеві надається можливість обрати лише необхідні параметри для яких необхідно створити змінні.

Підготована за допомогою інструмента `Simulink`-модель далі використовується для побудови додатка без необхідності мануального створення змінних (рисунк 2).



Рисунок 2 – Додаток на основі `Simulink`-моделі

Для управління параметрами моделі із бібліотеки компонентів `App Designer` на тло додатка переносяться елементи керування (зазвичай це `Edit Field(Numeric)`) та зв'язуються зі змінними `Simulink`-моделі. Виклик процедури моделювання та її зупинка здійснюється за допомогою компонента `Simulation Controls`, який має відповідні кнопки.

Таким чином розроблений інструмент дозволяє зменшити кількість операцій при створенні програм на основі `Simulink`-моделей в `App Designer`. Подальшим розвитком інструменту передбачається автоматизація створення сигналів для відображення в компоненті `Time Scope`.

Список використаних джерел

1. Craig S. Lent (2022) Learning to Program with MATLAB: Building GUI Tools (2th ed.). Wiley.
2. Create Custom Simulink Toolstrip Tabs – MATLAB & Simulink. URL: <https://www.mathworks.com/help/simulink/ug/create-custom-simulink-toolstrip-tabs.htm>.

Шуклін Г.В., к.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

МЕТОДИКА АНАЛІЗУ ІНФОРМАЦІЇ ЩОДО КІБЕРНЕТИЧНОГО СТАНУ ЕНЕРГЕТИЧНОГО ОБ'ЄКТУ

Активний розвиток і впровадження інтелектуальних технологій на підприємствах стало передумовою появи такого напрямку досліджень, як кібернетичний стан. Кібернетичний стан – сфера досліджень, пов'язана із застосуванням методів штучного інтелекту в кібербезпеці, спрямована на підвищення обізнаності щодо можливих ситуацій порушень кібербезпеки та автоматичне виявлення кіберзагроз. Під терміном «кібернетичний стан енергетичного об'єкту» розумітимемо обізнаність про стан кіберсередовища енергетичних об'єктів, яка включає інформацію: про критичні вразливості енергетичних об'єктів з погляду кібербезпеки; про кіберзагрози, які ініціюють ці критичні вразливості, а також про техногенні загрози енергетичній безпеці, спричинені кіберзагрозами.

Методика моделювання сценаріїв екстремальних ситуацій в енергетиці, викликаних реалізацією кіберзагроз, визначається за допомогою формування вузлів та їх взаємозв'язків у графі безсховної мережі довіри моделі і включає побудову графа, що складається з вершин декількох типів, відповідно до розробленої структури сценарію (рисунок 1), і взаємозв'язків між ними на основі побудованих правил виду «ЯКЩО – ТО» (формула 1).

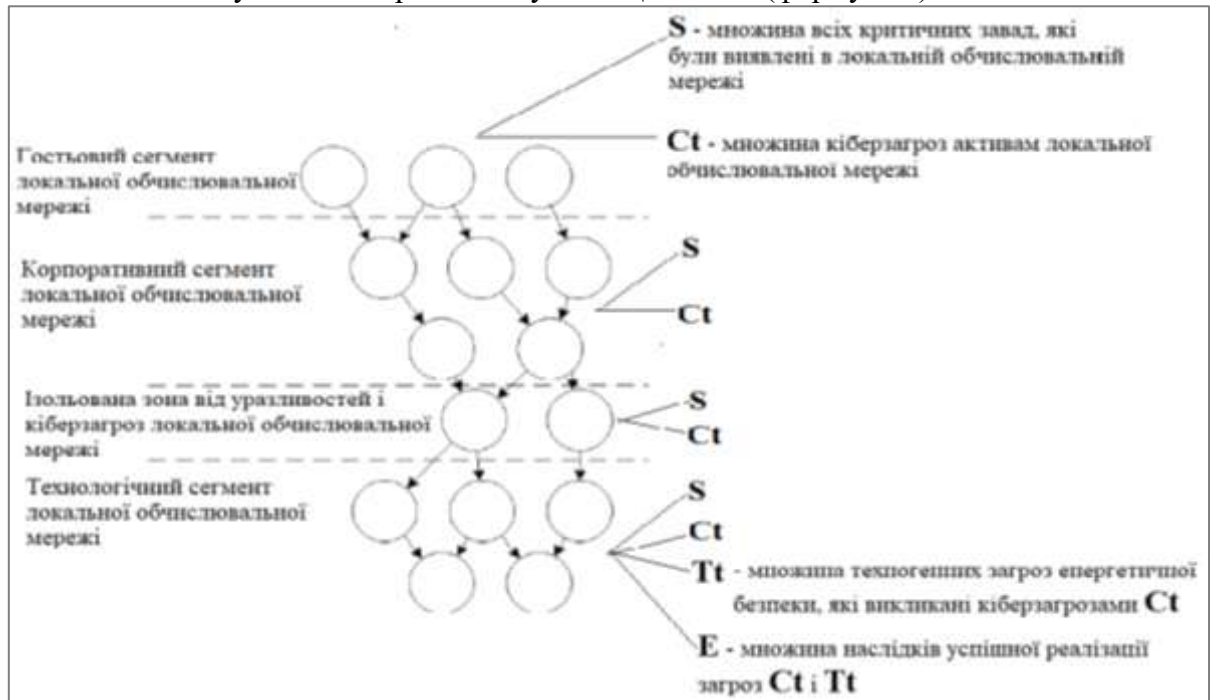


Рисунок 1 – Графічне представлення сценаріїв нештатних інцидентів, які виникають завдяки кіберзагрозам

Структура сценарію нештатних ситуацій в енергетиці, спричинених реалізацією кіберзагроз, розроблена відповідно до моделі, яку представлено на малюнку 1 і включає такі типи компонентів.

Структура сценарію також охоплює класифікацію концептів за типами, представленими вище, і за сегментами локальної обчислювальної мережі, яку розглядають: гостьовими, корпоративними, ізольованими зонами, технологічними сегментами.

Така структура сприяє візуальному відображенню можливих атак на технологічний сегмент локальної обчислювальної мережі із різних її сегментів.

1. Визначення графа бейєвської мережі довіри

$\Omega = \{k_1, k_2, \dots, k_n\}$ – множина вершин ациклічного, орієнтованого і зваженого графа, а $U = \{u_1, u_2, \dots, u_m\}$ – множина дуг цього графа. На даному графі введемо функцію

$$f : \Omega \rightarrow A, \quad (1)$$

де $A = \{A_1, A_2, \dots, A_n\}$ – множина матриць ваг (умовних ймовірностей) кожної вершини графа. Для кожної вершини k_i цієї матриці є ймовірності відповідної випадкової величини ξ_i .

Графом G бейєвської мережі довіри називається набір, який складається з трьох компонентів представлених у формулі 2:

$$G = (\Omega, U, f). \quad (2)$$

2. Визначення ймовірнісних характеристик сценарію

Кожній вершині k_i графа G відповідає єдина випадкова величина ξ_i така, що

$$\xi_i \in \{\xi(S) \cup \xi(Ct) \cup \xi(Tt) \cup \xi(E)\}, \quad (3)$$

де $\xi(S) \rightarrow S$, $\xi(Ct) \rightarrow Ct$, $\xi(Tt) \rightarrow Tt$, $\xi(E) \rightarrow E$.

Загальний вид таблиці умовних ймовірностей для відображення безумовних ймовірностей випадкової величини ξ_i у випадку, коли у відповідній вершини k_i відсутні пращури представлено в таблиці 1.

Таблиця 1 – Загальний вигляд умовних ймовірностей, коли у k_i відсутні пращури

ξ_i	$\bar{\xi}_i$
True	False
p	$1 - p$

Загальний вид таблиці умовних ймовірностей для відображення безумовних ймовірностей випадкової величини ξ_i у випадку, коли у відповідній вершини k_i наявні пращури наведено в таблиці 2.

Таблиця 2 – Загальний вигляд умовних ймовірностей, коли у k_i наявні пращури

$l_1(\xi_i)$	...	$l_k(\xi_i)$	ξ_i	$\bar{\xi}_i$
True	True	True	p_1	$1 - p_1$
...	...	...	...	...
True	True	True	...	...
False	False	False	...	...
...	...	...	...	...
False	False	False	p_k	$1 - p_k$

$$p(\xi) = (p_1(\xi), p_2(\xi), \dots, p_n(\xi)), \quad p_i(\xi) = p(\xi_i / l(\xi_i)). \quad (4)$$

3. Реалізація ймовірнісного експерименту виконується за формулою 5:

$$p(\xi_i / c) = \frac{\sum_{j=1}^m p_j(\xi, c)}{p(c)}. \quad (5)$$

де $p(\xi_i / c)$ – апостеріорна ймовірність випадкової величини ξ_i , $C = \{c_1, \dots, c_r\}$ – множина дискретних випадкових величин, що є свідками, m – кількість випадкових змінних, що залишилися і в яких не введено свідчення.

4. Аналіз альтернативних сценаріїв

Етап 1. Опис ризиків виконується за формулою 6:

$$R = \{S, Ct, Tt, E, D\}, \quad (6)$$

де D – множина збитків від наслідків сценаріїв.

Етап 2. Оцінка ризиків виконується за формулою 7:

$$R_i = p(e_i) \times D_i, \quad (7)$$

де $p(e_i)$ – ймовірність наслідку i -го сценарію, яка обчислюється за формулою (5), D_i – убуток від реалізації кіберзагроз i -го сценарію, $i = \overline{1, 2^n - 1}$.

Таблиця 3 – Умови оцінки сценаріїв нештатних інцидентів на енергетичних об'єктах, які виникають завдяки кіберзагрозам

Низький рівень загрози	Середній рівень загрози	Високий рівень загрози
$\begin{cases} 0 \leq p(e_i) \times D_i < q_1 \\ 0 \leq p(e_i) \leq 1 \\ 0 \leq D_i \leq \max D \end{cases}$	$\begin{cases} q_1 \leq p(e_i) \times D_i < q_2 \\ 0 \leq p(e_i) \leq 1 \\ 0 \leq D_i \leq \max D \end{cases}$	$\begin{cases} q_2 \leq p(e_i) \times D_i \\ 0 \leq p(e_i) \leq 1 \\ 0 \leq D_i \leq \max D \end{cases}$

У таблиці 1 q_1, q_2 – критерії кластеризації сценаріїв нештатних ситуацій, $\max D$ – максимальний убуток.

Етап 3. Ранжування активів інформаційної структури об'єкта енергетики та розробка рекомендацій.

Етап 4. Отримання інформації щодо рівня кібернетичного стану об'єкта енергетики виконується за формулою 8:

$$\eta = \frac{\lambda}{2^n - 1}, \quad (8)$$

де η – рівень кібернетичного стану об'єкта енергетики, λ – кількість розрахованих сценаріїв, $(2^n - 1)$ – кількість сценаріїв.

Список використаних джерел

1. Топал О. І., Любарець М. І., Прищепов Є. О., «Інтелектуальні системи обліку споживання енергетичних ресурсів: складники, технології радіозв'язку та перспективи впровадження в Україні» Радіотехніка і телекомунікації, Т. 31 (70), № 6, 2020, с. 27–34.
2. Барабаш О. В., Свинчук О. В., Бандурка О. І., «Програмне забезпечення контролю справного стану інформаційних систем в енергетичній галузі для забезпечення функціональної стійкості», Сучасний захист інформації, № 2 (58), 2024, с. 41–49.

Свинчук О.В., к.ф.-м.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Барабаш О.В., д.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

МЕТОДИ ЗАБЕЗПЕЧЕННЯ ФУНКЦІОНАЛЬНОЇ СТІЙКОСТІ БАГАТОМОДУЛЬНИХ ІНФОРМАЦІЙНИХ СИСТЕМ

Сучасний світ інформаційних технологій характеризується стрімким зростанням взаємозалежності та інтеграції багатомодульних інформаційних систем у всі сфери людської діяльності. Від фінансових операцій до управління критично важливою інфраструктурою – всі ці процеси залежать від стабільності та безпеки цифрових систем. Проте з розвитком та поширенням цих технологій зростає і кількість потенційних кіберзагроз. Це ставить під сумнів здатність існуючих інформаційних систем витримувати та адекватно реагувати на непередбачені втручання та збої, що можуть призвести до втрати даних, фінансових збитків, а в деяких випадках – до загрози людському життю.

Метою даної роботи є розробка програмного забезпечення для розрахунку параметрів функціональної стійкості багатомодульних інформаційних систем. Важливо використовувати різноманітні методи для оцінки можливих шляхів передачі даних і точки відмови. У роботі досліджуються класичні методи оцінки функціональної стійкості – точні та наближені. Кожен з цих методів має свої особливості, переваги та обмеження, залежно від специфіки системи, яка аналізується.

Точні методи використовують строгі математичні моделі для оцінки параметрів стійкості систем. Вони забезпечують високу точність результатів, але часто вимагають значних обчислювальних ресурсів та часу. Точні методи ідеально підходять для систем з відносно простою структурою та невеликою кількістю компонентів, де можливе детальне математичне моделювання. Проте деякі методи можуть бути використані лише для конкретних конфігурацій системи [1].

Наближені методи, в свою чергу, засновані на статистичних або евристичних підходах, що дозволяють отримати оцінку стійкості швидше і з меншими витратами ресурсів. Ці методи підходять для складних або великомасштабних систем, де точне моделювання є непрактичним. Вони дають наближену оцінку ймовірності зв'язності між елементами системи, тобто верхню і нижню оцінку. Найбільш відомими є методи Езарі-Прошана, Литвака-Ушакова та Полеського [1].

Програмне забезпечення містить інструменти для введення даних про топологію системи, а також модулі для виконання розрахунків і перегляду результатів (рисунк 1). Це допоможе користувачам легко інтерпретувати дані, отримані в результаті аналізу.

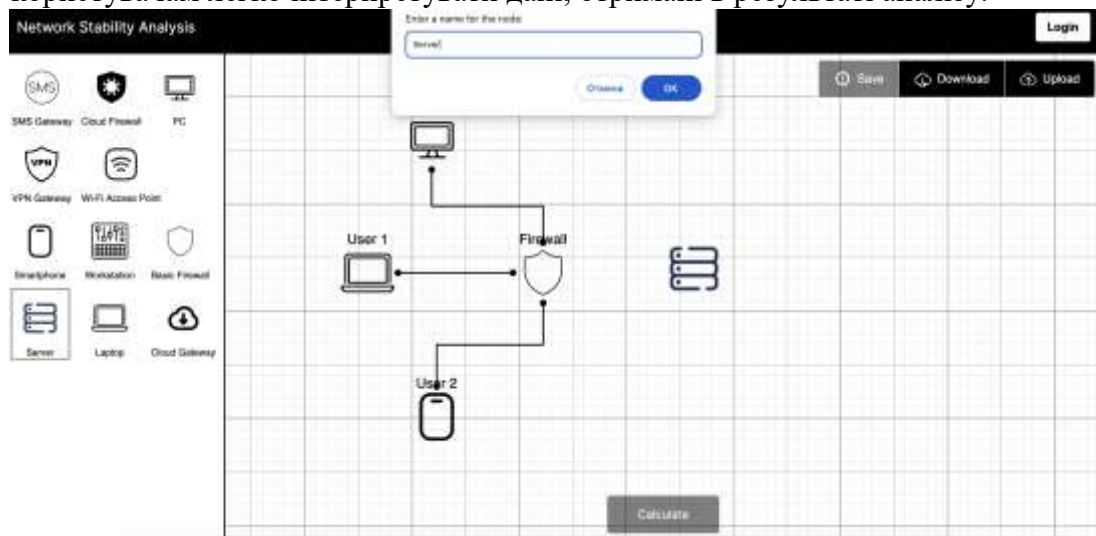


Рисунок 1 – Приклад побудованої інформаційної системи

Розроблена система містить можливість розрахунку функціональної стійкості системи. Даний інструмент аналізує елементи системи та їхні взаємозв'язки, враховуючи такі параметри, як пропускна здатність, безпека та надійність. Приклад результатів розрахунку функціональної стійкості системи для методу пошуку всіх простих шляхів та методу пошуку простих ланцюгів з урахуванням однакової ймовірності відмови представлений на рисунку 2.

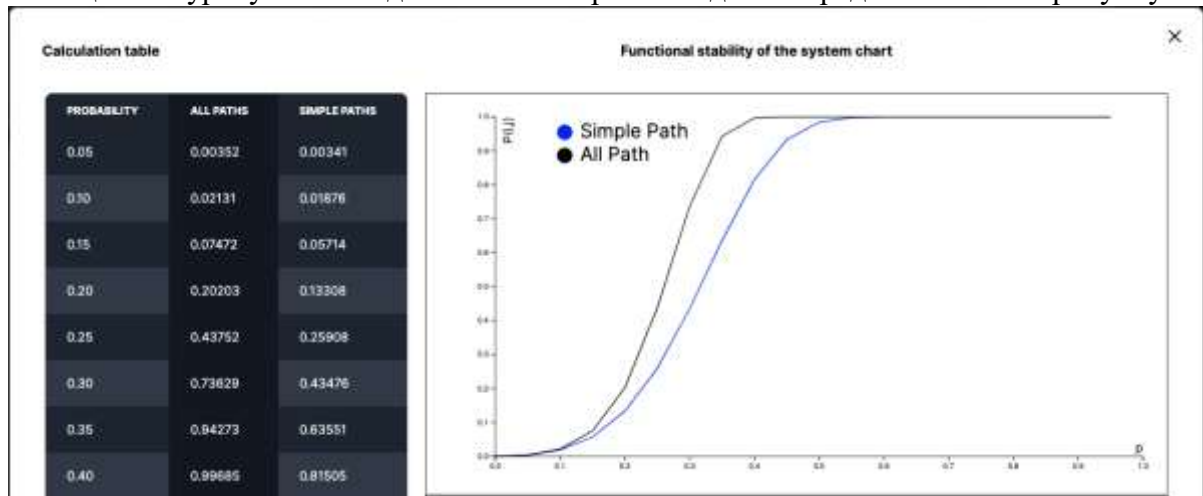


Рисунок 2 – Аналіз функціональної стійкості між двома вузлами системи

У таблиці наведено порівняльний аналіз обох методів, де кожен рядок представляє конкретну взаємодію або сценарій, за яким оцінюється функціональна стійкість. Це дозволяє порівняти ефективність обох методів за різних умов і різних навантажень. На цьому ж рисунку також показаний графік, що показує співвідношення та порівняння результатів алгоритму для різних значень ймовірності відмови. На графіку показано, як кожен алгоритм адаптується до мінливих умов і як це впливає на загальну стабільність системи. Це дає можливість виконати детальний аналіз і вибрати найкращий алгоритм для конкретної структури інформаційної системи.

Список використаних джерел

1. Собчук В. В., Барабаш О. В., Мусієнко А. П. (2022). Основи забезпечення функціональної стійкості інформаційних систем підприємств в умовах впливу дестабілізуючих факторів: монографія. Київ: Міленіум.

Пироговська Т.В., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Мусієнко А.П., д.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

МОДЕЛЮВАННЯ СИГНАЛУ ДАЛЬНЬОГО ПОЛЯ НА ОСНОВІ ВИКОРИСТАННЯ КОЕФІЦІЄНТІВ ЧЕБИШЕВА

На утворення та поширення звукових хвиль в водах океану мають вплив такі параметри системи, як погодні умови, солоність та густина водного середовища, просторові та енергетичні характеристики джерела випромінювання сигналу [1]. Дослідження утворення гідроакустичних каналів, розробка підводного зв'язку та телеметрії вимагають більшої деталізації в заданні вхідних параметрів [2]. В роботах [3-4] розглянуто генерацію імпульсів сферичним джерелом біля ідеальної границі та питання інтерференції взаємодії плоских та сферичних хвиль. Постановка хвильового рівняння має передбачати процес поширення акустичних збурень з врахуванням впливу відбитого та розсіяного гідроакустичного поля. Процес вирішення даного питання здійснювалося за використанням розробленого програмного комплексу моделювання звукового поля методом нормальних мод. Розробка, моделювання та використання такого підходу потребує введення часової залежності між змінними динамічної системи [4]. Це дозволить проводити більш глибоке дослідження процесу поширення звукових хвиль в підводному каналі, їх акустичний тиск в вертикальних перетиках середовища та встановлення коефіцієнтів збурення поля під час роботи з джерелом.

Такий підхід до моделювання потребує аналізу класичних моделей дослідження хвильових процесів. В роботі [5] розглядається двовимірний аналітичний модель хвильового рівняння в циліндричному хвилеводі. Оцінка похибки такого підходу наведена авторами в роботі [6]. Авторами [7] визначається оцінка відстані від випромінювача до початку координат та опис акустичного поля, лише через дискретні структури. В ході роботи [7] розглянутий аналітичний розв'язок для неоднорідного гідроакустичного поля на основі методу нормальних мод. Такий підхід враховує вплив донного прошарку на характеристики поля.

Найбільш ефективно застосування методу нормальних мод є для випромінювачів, що генерують звук на низькій частоті. В роботі [8] показано, що шум корабля, що віддаляється, можна віднести до низькочастотного. Під час подібних досліджень звукових хвиль та їх випромінювача, досить часто використовується тиск та акустична енергія [9], що визначені в кожній точці середовища. В [9] розглядаються спектри судового шуму та шуму перевезення з використанням методу нормальних хвиль. Для оцінки точності методу використовуються дані, отримані для профілю Мунка. Дно розглядається як шар з різною швидкістю та щільністю, тож метод можливо застосовувати для більш широкого спектру задач.

В більшості розглянутих задач, перехід до тривимірної моделі методу нормальних хвиль має певну проблематику, пов'язану зі стійкістю нелінійних нормальних мод [10]. Тривимірний модель, при застосуванні методу, не враховує в явному вигляді змінну часу. Вважається, що застосування методу нормальних мод можливе при умові, що моди є незмінними протягом певного часового періоду [11]. Проте, при дослідженні випромінювання точкового джерела циліндричного хвилеводу, врахування часової змінної є необхідним.

Розв'язуючи задачу методом нормальних мод, отримуємо функцію вигляду 1:

$$P(r, z) = \sqrt{\frac{2\rho}{h}} \sum_{m=1}^{\infty} \sin(\gamma_n z_0) \sin(\gamma_n z) H_0^1(k_m r), \quad (1)$$

$$\text{де } \gamma_n = \frac{\pi(0.5-m)}{n}, m=0, 1, 2, \dots \quad H_0^1(k_m r) \approx \sqrt{\frac{2}{\pi k_m r}} e^{i(k_m r - \frac{\pi}{4})}$$

При даному підході кожен доданок є відповідною нормальною модою, що не змінюється певний час. При введенні змінної часу доцільно було б просумувати знайдені моди. Такий підхід забезпечить зв'язок між нормальними модами, але не забезпечить вплив

кожної моди на систему в цілому. Тому постає необхідність введення коефіцієнтів, що забезпечували вплив нормальної хвилі на систему. В якості таких коефіцієнтів запропоновано використати коефіцієнти полінома Чебишева, а саме формула 2:

$$T_k(x) = \cos(k \cos^{-1}(x)), k = 0, 1, 2, \dots \quad (2)$$

Розширення можливо провести за допомогою будь-якої гладкої функції на проміжку $[-1, 1]$, тобто розв'язок задачі (1), (2) можливо зобразити у вигляді формули 3:

$$(r, z, t) = \sqrt{\frac{2\rho}{h}} \sum_{n=1}^{\infty} \sum_{k=1}^{\infty} \sum_{m=1}^{\infty} \cos(k \cos^{-1}(x)) \sin(\gamma_n z_0) \sin(\gamma_n z) H_0^1(k_m r) e^{-i\omega t}. \quad (3)$$

Маємо дискретний ряд з урахуванням часової змінної. Інтерполюємо значення отримані за допомогою (6) лінійною, квадратичною, експоненціальною функціями та сплайном для можливості встановлення неперервної функції апроксимації.

Для отримання даних для інтерполяції використаємо значення за формулою 4:

$$k_m = \sqrt{\frac{\omega^2}{c^2} - ((m + 0.5) \frac{\pi}{D})^2}, \quad (4)$$

де ω – частота поширення звукової хвилі, C – швидкість поширення звукової хвилі.

$C = 1500$ м/с; $\omega = 25$ Гц; глибина – 100 м; початкове положення випромінювача (0, 0, 50), дальність хвилеводу – 5000 м, діаметр – 50 м, тоді довжина звукової хвилі рівна 60 м. Крок по часу рівний 0,1 с, по z – 300 м.

Для здійснювання моделювання використовувався програмний комплекс моделювання гідроакустичного поля. Програмний комплекс складається з 5 пакетів: HydroAcModComp, HydroAcModCompViewModels, ModCompGeoFunction, SQLData, NormalModes. Реалізація обрахунків моделювання методом нормальних мод здійснена в пакеті NormalModes.

Результат роботи програмного забезпечення зображено на рисунку 1. Розроблене програмне забезпечення дозволяє задавати такі параметри середовища, як: батиметрія, шляхом обираючи ділянки дна для експерименту; профіль швидкості звуку, який відображає зміну відповідно до глибини; температуру, солоність, густину води у зоні експерименту; координати приймача, траєкторію джерела хвиль в експерименті.

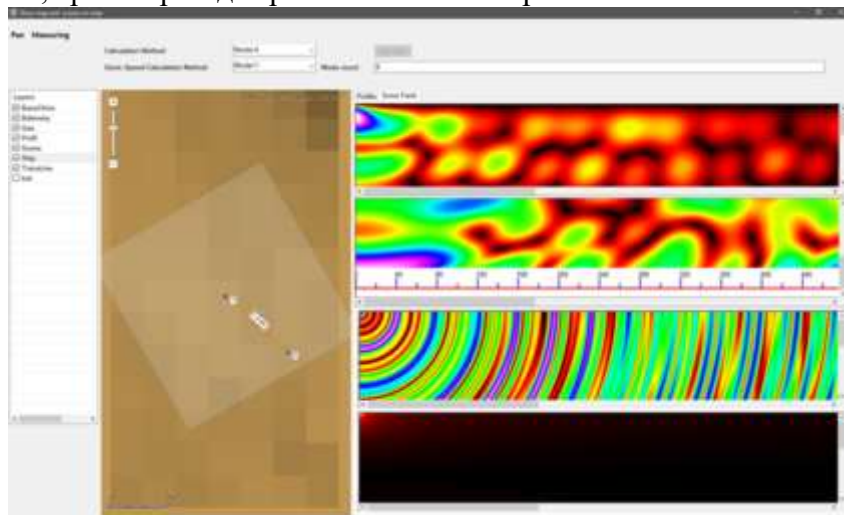


Рисунок 1 – Робота створеного програмного забезпечення

Висновки. Проведено комп'ютерне моделювання з використанням програмного забезпечення для моделювання гідроакустичного поля, отримані за результатами якого залежності дозволяють ввести часовий параметр у розв'язок хвильового рівняння методом нормальних хвиль. Це дає змогу передбачити розвиток звукової хвилі з часом і не вимагає подальших досліджень змінної часу. Наявність часової змінної дозволяє більш глибоко аналізувати фізичні та акустичні властивості звукових полів: їх поширення в підводному

каналі, визначення акустичного тиску на вертикальних розрізах середовища в будь-який момент часу, встановлення коефіцієнт збурення, в тому числі як функція часу. Перехід до тривимірної моделі рівняння нормальної хвилі з урахуванням часу дає можливість у подальших дослідженнях враховувати зміну параметрів нелінійного процесу з часом. Це дає можливість передбачити подальшу взаємодію між звуковою хвилею і передавачем.

Список використаних джерел

1. Brillouin L. (1960). Wave Propagation and Group Velocity. Academic Press, 166. Available at: <https://www.elsevier.com/books/wave-propagation-andgroup-velocity/brillouin/978-1-4832-3068-9>
2. Brekhovskikh L. (1976). Waves in Layered Media. Academic Press, 520. Available at: <https://www.elsevier.com/books/waves-in-layered-media/brekhovskikh/978-0-12-130560-4>
3. Mann J. A., Tichy J., Romano A. J. (1987). Instantaneous and time- averaged energy transfer in acoustic fields. The Journal of the Acoustical Society of America, № 82 (1), 17–30. DOI: <http://doi.org/10.1121/1.395562>
4. Mobarakeh P. S., Grinchenko V. T., Popov V. V., Soltannia B., Zrazhevsky G. M. (2020). Contemporary Methods for the Numerical-Analytic Solution of Boundary-Value Problems in Noncanonical Domains. Journal of Mathematical Sciences, № 247 (1), 88–107. DOI: <http://doi.org/10.1007/s10958-020-04791-4>
5. Korzhyk O., Naida S., Kurdiuk S., Nizhynska V., Korzhyk M., Naida A. (2021). Use of the pass-through method to solve sound radiation problems of a spherical electro-elastic source of zero order. EUREKA: Physics and Engineering, № 5, 133–146. DOI: <http://doi.org/10.21303/2461-4262.2021.001292>
6. Mobarakeh P. S., Grinchenko V. T. (2015). Construction Method of Analytical Solutions to the Mathematical Physics Boundary Problems for NonCanonical Domains. Reports on Mathematical Physics, № 75 (3), 417–434. DOI: [http://doi.org/10.1016/s0034-4877\(15\)30014-8](http://doi.org/10.1016/s0034-4877(15)30014-8)
7. Kazak M. S., Petrov P. S. (2020). On Adiabatic Sound Propagation in a Shallow Sea with a Circular Underwater Canyon. Acoustical Physics, № 66 (6), 616–623. DOI: <http://doi.org/10.1134/s1063771020060044>
8. Dyubchenko M. E. (1984). The influence of axisymmetric modes of vibrations on the sensitivity and directivity characteristics of a piezoceramic For reading only sphere. Acoustic Journal, № 30 (4), 477–481. Available at: http://www.akzh.ru/pdf/1984_4_477-481.pdf
9. Leiko O., Derepa A., Pozdniakova O., Starovoit Y. (2018). Acoustic fields of circular cylindrical hydroacoustic systems with a screen formed from cylindrical piezoceramic radiators. Romanian Journal of Acoustics and Vibration, № 15 (1), 41–46. Available at: <http://rjav.sra.ro/index.php/rjav/article/view/49>
10. Aronov B. (2009). Coupled vibration analysis of the thin-walled cylindrical piezoelectric ceramic transducers. The Journal of the Acoustical Society of America, № 125 (2), 803–818. DOI: <http://doi.org/10.1121/1.3056560>
11. Papkova J. I., Papkov S. O., Yaroshenko A. A. Energy characteristics of the hydroacoustic field in a nonuniform marine medium with a cylindrical body floating on the surface // Physical Oceanography – Springer Science + Business Media, 2006. – T. 16, № 3. – pp. 168–176
12. Papkova Yu. I. Method of normal modes for a three-dimensional model of a hydroacoustic waveguide // Dynamic systems. 2013. № 3–4. URL: <https://cyberleninka.ru/article/n/metod-normalnyh-mod-dlya-trehmernoy-modeli-gidroakusticheskogo-volnovoda> (дата звернення: 13.07.2022).

Олексій А.О., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Верлань А.А., д.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

ПОКРАЩЕНИЙ МЕТОД АНАЛІЗУ АКУСТИЧНИХ СИГНАЛІВ ВОДНОГО СЕРЕДОВИЩА НА ОСНОВІ ЗГОРТКОВОЇ НЕЙРОМЕРЕЖІ SOP З ЗАСТОСУВАННЯМ БАГАТОМАСШТАБНОЇ ЗГОРТКИ

CNN мають значні переваги перед іншими архітектурами завдяки своїй здатності автоматично виділяти найважливіші особливості даних, що зменшує потребу в ручному створенні ознак. Завдяки згортковим шарам CNN ефективно розпізнають просторові та частотні шаблони, а також мають меншу кількість параметрів, що дозволяє знизити обчислювальні витрати. У задачах аналізу акустичних сигналів у водному середовищі CNN особливо корисні для виділення характерних частотних компонентів, ігноруючи фоновий шум, що забезпечує високу точність класифікації навіть у складних умовах.

Для виділення ознак у підводних нестационарних сигналах ефективнішим є підхід у часово-частотній області. Перетворення з постійним Q (CQT) дозволяє перетворювати сигнал із часової області в часово-частотну область з високою частотною роздільною здатністю на низьких частотах, що надає більше деталей порівняно з короткочасним перетворенням Фур'є (STFT). У даному дослідженні CQT використовується для отримання часово-частотних ознак, які подаються на вхід запропонованої згорткової моделі. У цій роботі пропонується поєднання багатомасштабної згортки та методу пулінгу другого порядку (SOP) для захоплення часових кореляцій в ознаках, що покращує можливості розрізнення різних цілей за допомогою вхідних даних CQT.

Оригінальна архітектура складається з вхідного шару, двох згорткових шарів, шару SOP, шару нормалізації та повнозв'язного шару з активаційною функцією Softmax. Архітектура нейромережі наведено в таблиці 1.

Таблиця 1 – Архітектура оригінальної нейромережі

CQT input (64×64)
Conv. $32 \times 32 \times 8$
Conv. $16 \times 16 \times 16$
SOP $16 \times 16 \times 16$
Norm. 4096
Dense 1024
Softmax 5

Алгоритм роботи нейромережі:

1) CNN генерує карти ознак, де кожна карта відповідає певній частотній смужці у CQT-вхідному зображенні. Для кожної карти ознак формується траєкторія часових значень у кожному з каналів;

2) SOP використовує скалярний добуток (векторний добуток) між двома часовими траєкторіями для визначення кореляції між ними. Для кожного частотного біна обчислюється кореляційна матриця, яка є симетричною, позитивно напіввизначеною та відображає кореляції між усіма каналами CNN для цього біна;

3) кожен частотний бін (висота карти ознак) зберігає окремий результат SOP, тобто для кожної частоти створюється своя кореляційна матриця. Це дозволяє мережі вловлювати кореляції, зберігаючи частотні характеристики для кращого розрізнення класів;

4) отримані SOP-матриці перетворюються в один вектор, який потім нормалізується. Спочатку для кожного елемента обчислюється квадратний корінь, що знижує чутливість до великих значень. Далі застосовується l2-нормалізація, яка стандартизує вектор до фіксованої довжини;

5) після цього отриманий SOP-вектор передається через повнозв'язний шар до шару Softmax для визначення ймовірностей приналежності до класів.

До нейромережі було додано шар багатомасштабної згортки, що використовує три паралельні фільтри різних розмірів (3×3 , 5×5 , 7×7), об'єднуючи їхні результати через конкатенацію та нормалізацію. Це дозволяє виділяти ознаки на різних масштабах. Також додано середній пулінг, який зменшує кількість даних, фокусуючись на важливіших ознаках, що спрощує процес класифікації. Архітектура покращеного методу представлена в таблиці 2.

Таблиця 2 – Архітектура покращеної нейромережі

CQT input (64×64)
Multiscale conv (3×3 , 5×5 , 7×7)
Avg. pooling
SOP $16 \times 16 \times 16$
Norm. 1024
Dense 1024
Softmax 5

Покращення точності класифікації завдяки додаванню багатомасштабної згортки можливе, оскільки багатомасштабна згортка забезпечує такі переваги:

- багатомасштабна згортка використовує кілька ядер згортки різного розміру (3×3 , 5×5 , 7×7). Це дозволяє кожному ядру захоплювати різні масштаби деталей із сигналу, що особливо корисно для акустичних сигналів, де важливі як дрібні деталі, наприклад, високочастотні коливання, так і більші структурні ознаки, наприклад, низькочастотні компоненти;

- використання згорток різного розміру дозволяє моделі навчатися багатошаровим абстракціям сигналу, що допомагає покращити точність класифікації. Зокрема, такий підхід сприяє зменшенню ймовірності втрати важливих деталей, які можуть бути представлені лише на одному масштабі;

- оскільки акустичні сигнали можуть містити широке розмаїття особливостей, зокрема гармонічні та шумові компоненти, різні розміри ядер дозволяють мережі зосередитися на виявленні широкого спектру ознак у спектральному представленні сигналу;

- завдяки застосуванню нормалізації покращується збіжність, що дозволяє моделі ефективніше навчатися на великій кількості акустичних даних;

- застосування багатомасштабної згортки дозволяє моделі навчатися ігнорувати певні види шуму та одночасно виділяти суттєві ознаки на різних масштабах, що робить її більш стійкою до шумових змін у сигналі.

Було створено датасет на основі датасету ShipsEar [2]. Ці сигнали складаються з записів, що включають 4 типи кораблів і поділяються на 4 класи. Файли були сегментовані на фрагменти розміром 2 секунди кожен. Також було створено датасет з високим рівнем шуму, де середнє співвідношення сигнал – шум сягало -10. На базі датасету Storm [3] було створено аналогічний датасет. Він містить 5 класів кораблів. Штучний датасет є комбінацією декількох частот із декількома гармоніками різної амплітуди. Додатково було створено штучний датасет із додаванням штучних шумів. Для штучного датасету було створено 4 класи. Усі датасети розділені на набори для навчання, валідації та тестування у пропорції 70 %, 20 % та 10 %.

Оригінальний та покращений методи були протестовані на вказаних вище трьох оригінальних датасетах та датасетах із додаванням шумів. Результати представлені в таблиці 3.

Таблиця 3 – Результати тестування оригінальної та покращеної нейромережі шляхом застосування різні датасети

	Оригінальний метод	Покращений метод
Датасет Storm	95 %	98 %
Датасет Storm з фоновими шумами	61 %	68 %
Датасет Shiptear	95 %	99 %
Датасет Shiptear з фоновими шумами	63 %	74 %
Штучний датасет	100 %	94 %
Штучний датасет з фоновими шумами	85 %	94 %

Список використаних джерел

1. Cao X., Togneri R., Zhang X., Yu Y., «Convolutional Neural Network with Second-Order Pooling for Underwater Target Classification». In IEEE Sensors Journal, T. 19, № 8, pp. 3058-3066, 15 April 2019, DOI: 10.1109/JSEN.2018.2886368.
2. Santos-Domínguez D., Torres-Guijarro S., Cardenal-Lopez A., Pena-Gimenez A. ShipsEar: An underwater vessel noise database. Applied Acoustics. 2016; № 113, pp. 64–9.
3. Вимірювальні системи та програмне забезпечення для морських охоронних систем і дослідницьких полігонів: звіт про НДР (заключ.). НТУУ «КПІ»; кер. роб. Є. Мачуський. – К., 2012. – 104 л. + відеосюжет + CD-ROM. – Д/б № 2429-п.

Здор К.А., Національний технічний університет України
 «Київський політехнічний інститут імені Ігоря Сікорського»
 Шалденко О.В., к.т.н., Національний технічний університет України
 «Київський політехнічний інститут імені Ігоря Сікорського»
 Недашківський О.Л., д.т.н., Національний технічний університет України
 «Київський політехнічний інститут імені Ігоря Сікорського»

РОЗБИТТЯ ВІДЕО НА СЦЕНИ ЗА ДОПОМОГОЮ ВІЗУАЛЬНИХ ТРАНСФОРМЕРІВ ДЛЯ ВІДЕО

Виявлення зміни сцен у візуальних медіа відіграє важливу роль для таких сфер як аналіз відео контенту, відеоспостереження та організація цифрового контенту. Математичні методи розв'язання цієї задачі мають недостатню точність через нездатність визначати контекст закладений в кожен окрему сцену, що призводить до низької точності [1]. Наявні підходи, ґрунтовані на застосуванні нейронних мереж, краще справляються з задачею розпізнавання концептуальних змін між зображеннями, але потребують величезних обчислювальних потужностей [2]. Також викликом є аналіз сучасного відео контенту, який стає з кожним роком все більш динамічним та має більш складні види переходів між сценами [3].

Розглядаючи різні архітектури, треба також приділити увагу їх швидкодії через те, що при застосуванні цих підходів для задач реального світу, треба швидкість аналізу, яка дозволяє аналізувати контент в режимі реального часу. Розглядаючи згорткові нейронні мережі, ми стикаємося з проблемою узагальнення через те, що не можна передати декілька кадрів для аналізу. Цю проблему можна обійти, використовуючи кодування зображень на основі згорткових нейронних мереж, а далі на основі цих закодованих значень визначати зміну плану. Використовуючи сіамські нейронні мережі, можна написати детектор зміни сцени, але проблемою цього підходу буде відсутність знань моделі про попередні розрахунки. Також для розв'язання цієї задачі можна застосовувати рекурентні нейронні мережі. Вони мають механізм пам'яті, що дозволяє їм ефективно зберігати інформації про попередні кадри та використовувати це для точного визначення зміни сцени. Недоліком цього підходу буде низька швидкість обчислень. Було вирішено використовувати архітектуру візуального трансформера для відео, яку показано на рисунку 1, оскільки вона дозволяє ефективно визначати концептуальні особливості для кожного кадру, може аналізувати декілька кадрів одночасно, що дозволяє ефективно їх порівнювати не тільки на візуальному, але і на концептуальному рівні та має кращу швидкість за архітектури на основі загорткових та рекурентних нейронних мереж [10, 11].

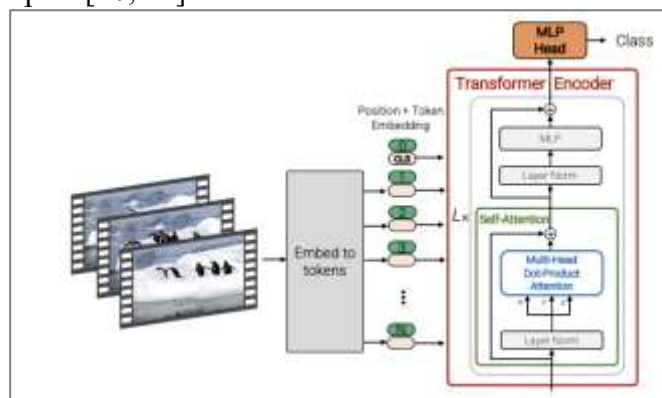


Рисунок 1 – Приклад архітектури візуального трансформера для відео

Також для покращення швидкодії було вирішено використовувати лише частину кадрів з відео як показано на рисунку 2. Але при простому пропусканні кадрів втрачається можливість визначення кадру на якому відбулася зміна сцени, тому перед аналізом зміни сцени було вирішено застосовувати алгоритм знаходження зміни планів. Для цієї задачі було застосовано нейро-математичний підхід для виявлення змін планів у відео послідовностях.

Цей підхід складається з двох етапів, перший це визначення різниці між кадрами, другий це передача різниці в рекурентну нейронну мережу для визначення зміни плану [9]. Для визначення різниці між планами алгоритм розбиває кожен кадр на блоки [4], створює різні представлення у декількох кольорових просторах [5], обраховує контури за допомогою фільтра Собеля-Фельдмана [8] та на основі отриманих результатів розраховуються гістограми [6]. Далі відстань між сусідніми кадрами розраховується за допомогою Евклідової відстані між гістограмами [7]. Розраховані відстані передаються в рекурентну нейронну мережу, яка визначає зміну плану. Використавши цей підхід ми можемо зменшити кількість оброблювальних кадрів для визначення зміни сцен до кількості планів на відео або до кількості кратній кількості планів.



Рисунок 2 – Приклад зміни сцени використовуючи лише частину кадрів

Для навчання моделі на основі архітектури візуального трансформера для відео було вирішено взяти класичні датасети OS VSD, Rai та BBC Planet Earth. Через те, що переважна кількість статей вимірювала власну точність на датасеті BBC Planet Earth було вирішено використовувати його як валідаційний сет, а датасети OS VSD та Rai для тренування. В результаті було використано 31 відео для тренування та 11 відео для валідації результатів навчання. Кожне відео в датасеті було проаналізовано за допомогою алгоритму розбиття відео на плани та на основі отриманих результатів було зроблено таблицю відношення планів до сцен. Даний підхід дозволив отримувати ключові кадри для кожного плану та передавати їх в модель для визначення зміни сцен. Через те, що тренувальний набір даних вийшов недостатньо великим для ефективного навчання було вирішено застосовувати методи штучного розширення датасету. Під час навчання кадри випадковим чином обрізались, змінювався їх розмір, змінювалась яскравість кадру, додавався гаусовий шум та інші трансформації. Окрім цього обирались не тільки сусідні плани в сцені, а будь-які плани які належать одній сцені та брати з них випадковий кадр. Сусідні сцени також могли міняти місцями. Застосування цих підходів дозволило ефективно розширити датасет, досягнути стабільних результатів та гарного узагальнення даних.

Під час розробки моделі тестувались різні варіанти подачі фреймів у модель та оптимальна кількість фреймів необхідна для визначення концептуальних залежностей між сценами. В результаті експериментів було визначено, що найбільша точність досягається при застосуванні 4 кадрів, це дозволяє моделі визначати концепт кожного кадру при цьому не передаючи надлишкову інформацію, яка може виступати як шум.

Запропонований підхід представляє новий спосіб визначення змін сцен і надає можливість подальшого покращення результатів шляхом збільшення обсягу навчальних даних, експериментів із архітектурою трансформерів та прискорення розробленого методу завдяки застосуванню алгоритму прунінгу. Цей підхід також дозволяє створювати ефективні архітектури для аналізу відео з можливістю горизонтального масштабування та ефективного використання хмарних інфраструктур. Крім того, точне визначення планів та сцен у відео сприяє ефективному виконанню завдань аналізу відео, таких як пошук та узагальнення.

Висновки. Було розроблено новий підхід який дозволяє ефективно визначати зміни сцен на відео. При розробки цього підходу було інтегровано метод розбиття відео на плани який оснований на поєднанні математичного підходу з нейронною мережею. Це дозволило ефективно розділити відео на плани та значно зменшити необхідну кількість даних для визначення зміни сцени. Також було досліджено та модифіковано для цієї задачі датасети Rai, OS VSD та Discovery Planet Earth. За допомогою визначення сцен для цих відео та додаткових трансформацій вдалось вирішити проблему недостатньої кількості навчальних даних, також Було застосовано архітектуру візуального трансформера для відео, що дозволило досягнути передових результатів внаслідок здатності моделі ефективно визначати контекст фреймів, переданих на аналіз, та порівнювати їх між собою. Також застосування архітектури візуальних трансформерів для відео відкриває шлях для подальшого покращення точності моделі шляхом створення нових датасетів для навчання, що було недоступно для попередніх рішень на основі математичних підходів, та дає можливість будувати системи, які можуть масштабуватись горизонтально та застосовуватись в хмарній інфраструктурі.

Список використаних джерел

1. Weiyao L., Ming-Ting S., Hongxiang K., Hai-Miao H.. (2010). A New Shot Change Detection Method Using Information from Motion Estimation. 264-275. 10.1007/978-3-642-15696-0_25.
2. Tomáš S., Jakub L. (2020). TransNet V2: An effective deep network architecture for fast shot transition detection.
3. H. Abdulhussain, Sadiq & Ramli, Abd Rahman & Saripan, M Iqbal & Mahmmod, Basheera & Al-Haddad, Syed Abdul Rahman & Jassim, Wissam. (2018). Methods and Challenges in Shot Boundary Detection: A Review. Entropy. 20. 10.3390/E20040214.
4. Soyoung P., Jeongwoo S., Sun-Joong K. (2016). Study on the effect of frame size and color histogram bins on the shot boundary detection performance. 1–2. 10.1109/ICCE-Asia.2016.7804726.
5. Ibrahim Z., Khaled E., Eid E.. (2016). Abrupt Cut Detection in News Videos Using Dominant Colors Representation. 10.1007/978-3-319-48308-5_31.
6. Jordi M., Gabriel F.. Video shot boundary detection based on color histogram.
7. Partha M., Sanjoy S., Bhabatosh Ch.. (2012). A Model-Based Shot Boundary Detection Technique Using Frame Transition Parameters. Multimedia, IEEE Transactions on. 14. 223-233. 10.1109/TMM.2011.2170963.
8. Zhao H., Li X., Yu L.. (2008). Shot Boundary Detection Based on Mutual Information and Canny Edge Detector. 1124–1128. 10.1109/CSSE.2008.939.
9. Shaldenko O., Zdor K. (2024). Neuro-mathematical fusion for shot change detection in video sequences. Actual Issues of Modern Science. European Scientific e-Journal, 29, 15–24. Ostrava: Tukulart Edition, European Institute for Innovation Development.
10. Bo-Kai R., Hong-Han S., Wen-Huang Ch.. (2022). Vision Transformers: State of the Art and Research Challenges. 10.48550/arXiv.2207.03041.
11. Arnab A., Dehghani M., Heigold G., Sun Ch., Lucic M., Schmid C. (2021). ViViT: A Video Vision Transformer. 10.48550/arXiv.2103.15691
12. Baraldi L., Grana C., Cucchiara R. (2015). Shot and Scene Detection via Hierarchical Clustering for Re-using Broadcast Video. 9256. 10.1007/978-3-319-23192-1_67

Секція 1

ІНСТРУМЕНТАРІЙ ТЕХНОЛОГІЙ ПРОГРАМУВАННЯ

Макарчук А.В., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Барабаш О.В., д.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

ЗАСТОСУВАННЯ МОДЕЛЕЙ МАШИННОГО НАВЧАННЯ ДЛЯ ІДЕНТИФІКАЦІЇ УМОВ, ЩО ГАРАНТУЮТЬ ФУНКЦІОНАЛЬНУ СТІЙКІСТЬ ІНФОРМАЦІЙНИХ СИСТЕМ

На даний момент сформульовано чимало показників функціональної стійкості та вимог до інформаційної системи, виконання яких функціональну стійкість забезпечує [1]. Однак, головною проблемою більшості з цих показників та умов є те, що їх важко використовувати із-за високої обчислювальної складності. Тому виникає необхідність в розробці методів їх оцінки, які б вимагали менше обчислень.

Як відомо, зараз підвищується інтерес до машинного навчання, в тому числі, й в контексті оптимізації обчислень [1, 2]. Якщо говорити про оцінку показників функціональної стійкості за допомогою моделей машинного навчання, то такі спроби вже проводилися [3-5]. Однак, спроб оптимізувати ідентифікувати виконання тих чи інших вимог, виконання яких забезпечуватиме функціональну стійкість, майже не проводилося. В даній роботі проводилася спроба оптимізувати встановлення умови (формула 1):

$$\chi(G) \geq 2 \wedge \lambda(G) \geq 2, \quad (1)$$

де G – граф, що описує топологію розглядуваної інформаційної системи, $\chi(G)$ – степінь вершинної зв'язності графа G , $\lambda(G)$ – степінь реберної зв'язності графа G , за допомогою трьох моделей машинного навчання: дерев рішень, випадкових лісів та нейронних мереж прямого розповсюдження.

У межах дослідження було вибрано 6 параметрів, які подаватимуться на вхід моделі: мінімальна, середня та максимальна кількість ліній зв'язку, проведених до машини, кількість машин і ліній зв'язку в системі, а також L_{22} – норма матриці суміжності графа, що описує топологію інформаційної системи. На виході очікується бінарна величина, рівна одиниці, якщо при заданих параметрах умова (1) виконується, і рівна нулю в протилежному випадку. Начальна та тестова вибірки було згенеровано випадковим чином.

При навчанні та тестуванні вище зазначених моделей машинного навчання на основі згенерованої вибірки було отримано наступні результати (таблиця 1).

Таблиця 1 – Результат тренування та подальшого тестування вибраних моделей машинного навчання

Модель	Відсоток правильних відповідей	Достовірність правильної ідентифікації виконання умови (1)	Достовірність правильної ідентифікації невиконання умови (1)
Дерево рішень	86,2	0,88	0,847
Випадковий ліс	86,84	0,882	0,856
Нейронна мережа прямого розповсюдження	85,5	0,868	0,843

Із таблиці 1 чітко видно, що при зазначеному вище наборі параметрів, що передаються на вхід, серед вибраних моделей машинного навчання найкращу точність дає випадковий ліс. Таку точність можна отримати, якщо для випадкового лісу поставити наступні вимоги:

- мінімальна кількість об'єктів у листі дерева рівна 1;

- мінімальна кількість об'єктів в листі, достатня для розбиття, рівна 4;
- кількість дерев рішень у випадковому лісі повинна бути рівна 16.

При більш детальному вивченні сформованої вибірки було виявлено, що деякі з вище зазначених параметрів за заданих умов можуть бути попарно корельованими. Це, в свою чергу, нашоухує на думку, що, видаливши деякі з цих параметрів та перекалібрувавши вибрані моделі машинного навчання можна отримати приблизно таку ж, а, можливо, й кращу точність. Так, серед вибраних параметрів, які подаються на вхід, потенційно придатними на виключення є середня і максимальна кількість ліній зв'язку, проведених до машини, та кількість машин в системі. При проведенні експериментів виявлено, що із вибраного списку параметрів можна виключити перші два із зазначених і, при цьому, мати приблизно таку ж точність при відповідній перекалібровці тобто, другими словами, на вхід розглянутих моделей можна подавати лише мінімальну кількість ліній зв'язку, проведених до машини, кількість машин і ліній зв'язку в системі, а також L_{22} – норма матриці суміжності графа, що описує топологію інформаційної системи, і при необхідній калібровці отримати точність, майже таку ж, як і в табл. 1 (таблиця 2). В інших ситуаціях буде різко падати точність прогнозу виконання або невиконання умови (1).

Таблиця 2 – Точність вибраних моделей машинного навчання після видалення з тренувальної вибірки середньої та максимальної кількості ліній зв'язку, проведених до машини

Модель	Відсоток правильних відповідей	Достовірність правильної ідентифікації виконання умови (1)	Достовірність правильної ідентифікації невиконання умови (1)
Дерево рішень	85,9	0,89	0,83
Випадковий ліс	86,7	0,881	0,852
Нейронна мережа прямого розповсюдження	86,5	0,886	0,847

Як можна бачити на таблицях 1 і 2, можна побачити, що серед вибраних моделей машинного навчання найоптимальнішою в плані точності є модель, основана на випадковому лісі. Можна бачити, що нейронні мережі можуть давати подібну точність в контексті розглядуваної задачі, але при відповідному калібруванні вони можуть вимагати значно більше обчислень.

Список використаних джерел

1. Барабаш О. В. (2004). Побудова функціонально стійких розподілених інформаційних систем. Київ: НАОУ.
2. Барабаш О., Лукова-Чуйко Н., Мусієнко А., Собчук В. (2018). Забезпечення функціональної стійкості інформаційних мереж на основі розробки методу протидії DDoS-атакам. Сучасні інформаційні системи, № 2 (1), С. 55–63.
3. Cen J., Chen X., Xu M., Zou Q. (n.d.). Deep finite volume method for high-dimensional partial differential equations. Guangzhou. (Preprint).
4. Zhang X., Helwig J., Lin Y., Xie Y., Fu C., Wojtowysch S., Ji Sh. (n.d.). SineNet: Learning temporal dynamics in time-dependent partial differential equations. Texas. (Preprint).
5. Dulny A., Heinisch P., Hotho A., Krause A. (n.d.). GrIND: Grid Interpolation Network for Scattered Observations. Würzburg: University Würzburg. (Preprint).

РОЗПОДІЛЕНІ ОБЧИСЛЕННЯ У КОНТЕКСТІ ОБ'ЄКТНО-ОРІЄНТОВАНОГО ПРОГРАМУВАННЯ: ВИКЛИКИ, ПІДХОДИ ТА ПЕРСПЕКТИВИ

Об'єктно-орієнтоване програмування (ООП) є потужним підходом для створення гнучких та добре структурованих систем. Дуже важливо навчати студентів правильно застосовувати принципи та підходи ООП в розподіленому програмуванні, адже сучасні розподілені системи стикаються з низкою викликів. Це високі навантаження на мережу, необхідність надійної синхронізації даних між вузлами та забезпечення стійкості до збоїв.

Принципи ООП [1] дозволяють забезпечити логічну та ієрархічну структурування компонентів у розподілених системах. Інкапсуляція дозволяє захистити внутрішній стан об'єктів. Її застосування дозволяє ізолювати внутрішню реалізацію об'єктів, завдяки чому можна оптимізувати передачу лише необхідних даних. Успадкування спрощує розробку ієрархій сервісів шляхом створення узагальнених класів та розширення їх специфічними властивостями. Також завдяки успадкуванню можна легко реалізувати стратегії відновлення, наслідуючи базовий клас для обробки помилок, що дає змогу повторно використовувати код та зменшити ймовірність помилок при реалізації логіки відновлення. Поліморфізм забезпечує реалізацію динамічних інтерфейсів для забезпечення гнучкої роботи з різними типами запитів у системах.

Не варто забувати, що інтеграція ООП у мікросервісну та хмарну архітектуру суттєво спрощує створення, тестування та підтримку масштабованих систем. Мікросервісна архітектура забезпечує ізоляцію окремих функціональних модулів як незалежних об'єктів із визначеними методами. Своєю чергою, хмарні обчислення також використовують принципи ООП для структурування ресурсів.

У розподілених системах забезпечення коректності й узгодженості даних є надзвичайно важливим. ООП пропонує механізми для організації роботи з даними для зменшення ризику виникнення конфліктів. За допомогою об'єктів для представлення стану системи дозволяє визначити, які частини даних можуть бути оновлені асинхронно, а які потребують синхронізації. Серед популярних інструментів для розробки розподілених систем на основі ООП варто відзначити Apache Thrift, gRPC від Google, Java RMI

Попри переваги, розподілені системи на базі ООП стикаються з технічними труднощами [2]. Серйозним викликом є високий обсяг міжпроцесного зв'язку (IPC), як і синхронізація об'єктів на різних вузлах системи. Особливо критичним є питання безпеки використання чутливих даних у розподілених середовищах. ООП пропонує механізми інкапсуляції, що дозволяє захистити дані від несанкціонованого доступу, надаючи контроль над тим, які частини програми можуть взаємодіяти з даними об'єкта.

Проте об'єктно-орієнтоване програмування надає систематизований і структурований підхід до проектування розподілених систем, що дозволяє ефективно вирішувати виклики, пов'язані з високими навантаженнями на мережу, синхронізацією даних та забезпеченням стійкості до збоїв. Цей підхід не тільки підвищує надійність і гнучкість систем, а й дозволяє зосередитися на реалізації функціональності, не турбуючись про деталі реалізації. Структурованість ООП дозволяє розробникам і тестувальникам працювати над різними компонентами одночасно, що підвищує ефективність команди та скорочує час виходу на ринок.

Перспективи використання ООП в розподілених системах базуються на можливості легкої модифікації і розширення системи. Завдяки принципам успадкування і поліморфізму розробники можуть створювати нові об'єкти, наслідуючи існуючі, що значно спрощує процес масштабування системи.

Ще однією перспективою є інтеграція ООП і функціонального програмування для підвищення гнучкості та продуктивності систем, а також дослідження нових технологій, які зможуть підвищити ефективність ООП у розподілених середовищах.

Список використаних джерел

1. Коваль А. І., Яшина О. М., Радельчук Г. І., Форкун Ю. В. Порівняння об'єктно-орієнтованої та функційної парадигм програмування у проектуванні програмного забезпечення. Вісник Хмельницького національного університету. 2021. № 3 (297). С. 34–38.
2. Минайленко Р. М., Мелешко Є. В. Проблеми розподілених обчислень та шляхи їх вирішення. Центральноукраїнський науковий вісник. 2021. № 4 (35). С. 3–7.

Черноусов Д.І., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Бандурка О.І., д.ф., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

ЗАСОБИ РОЗРОБКИ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ З ВІДКЛАДЕНИМИ ПЕРІОДИЧНИМИ ЗАДАЧАМИ

Існує багато рішень для створення відкладених періодичних задач, кожна з яких має свої особливості, переваги та сферу використання і застосовується у більшості застосунків, що використовують мережу інтернет. Прикладом таких задач є регулярна email розсилка. Буде розглянуто популярні рішення заради розуміння суті роботи таких програм, визначення того як використовувати та комбінувати ці програми, щоб оптимізувати використання комп'ютерних ресурсів серверної частини інформаційної системи.

Деякі з цих програмних засобів є застарілими, серед них Celery. Celery — це популярний інструмент для асинхронного виконання задач у Python. Він призначений для розподілених обчислень і підходить для обробки відкладених періодичних задач. Завдяки механізму черг Celery дозволяє виконувати задачі у фоновому режимі, розвантажуючи основний сервер і переносючи операції на сервери робітники, так звані воркери [1].

Сильним мінусом є те, що стандартні воркери Celery не мають вбудованої підтримки *asuncio*, що унеможливорює виконання асинхронних функцій. Тобто, кожна задача займає весь потік виконання до завершення. Наприклад, при надсиланні HTTP запиту на інший сервер задача буде очікувати весь час відповіді від серверу, не вивільняючи ресурси на виконання інших задач.

Більш розвинутим інструментом є Cron jobs, або заплановане завдання cron, – це потужний інструмент для автоматизації завдань, який використовується в Unix-подібних операційних системах. Cron job, або демон cron, написаний низькорівневою мовою програмування C, може бути використаний з іншими мовами[2]. Він є частиною ядра операційної системи Linux. Cron дозволяє запускати команди та скрипти в певний час або з певною періодичністю. За замовчуванням він підтримує асинхронність.

Як працює cron job:

- 1) користувачі зберігають свої cron-завдання в текстовому файлі під назвою *crontab*. Цей файл містить рядки, кожен з яких описує одне завдання;
- 2) демон cron – це системна програма, яка постійно працює в фоновому режимі та перевіряє *crontab*-файли на наявність завдань, які потрібно виконати;
- 3) демон порівнює час на даний момент з часом вказаним у розкладі. Якщо вони співпадають, cron запускає відповідну команду;
- 4) команда виконується у відділеному процесі. Результат виконання може бути записом у журнал або відправлення повідомлення по електронній пошті користувачеві.

Структура даних, яка використовується cron, зазвичай полягає у внутрішній таблиці розкладу. Ця таблиця зберігає всі розклади завдань та інші відомості, необхідні для їх виконання. Кожен запис таблиці містить такі дані:

- 1) поля, що вказують, коли завдання повинно бути виконане. Зазвичай це поля, які відповідають за хвилини, години, дні місяця, місяці та дні тижня;
- 2) команда або скрипт, який повинен бути виконаний;
- 3) інші параметри, такі як шлях до виконуваного файлу, середовище тощо.

Фонові завдання є невід'ємною частиною Android-додатків, дозволяючи їм виконувати операції навіть коли додаток не використовується безпосередньо[3]. Це і є однією з причин чому більшість сучасних смартфонів мають багатоядерні процесори, де більшість ядер є значно слабшими за основне ядро. Це важливо, адже головний потік, відповідальний за взаємодію з користувачем та оновлення екрану, не повинен бути перевантажений довготривалими процесами. До таких фонових завдань належать:

- 1) завантаження аналітичних звітів;

- 2) синхронізація інформації;
- 3) автоматичне створення резервних копій даних;
- 4) обробка геолокаційних даних.

Існують обмеження щодо фонових завдань, щоб економити заряд батареї та ресурси. Система Android розподіляє додатки на різні категорії очікування залежно від частоти їх використання. Додатки, які використовуються часто, отримують більше свободи для фонових завдань, тоді як ті, які використовуються рідко, обмежуються.

Поширені способи обробки фонових завдань в Android:

Служби – це довготривалі фонові процеси, які можуть виконувати завдання навіть коли користувач виходить з програми. Вони є гарним вибором для таких речей, як відтворення музики або синхронізація даних.

Приймачі мовлення – це керовані подіями прослуховувачі, які реагують на системні трансляції, такі як зміни доступності мережі або рівень заряду батареї. Їх можна використовувати для запуску фонових завдань, коли відбуваються певні події.

WorkManager – це новіший за JobScheduler, більш гнучкий інструмент планування фонових завдань. Він дозволяє визначати обмеження, такі як необхідність підключення Wi-Fi або зарядки, перш ніж завдання буде запущено.

Foreground Service – це вид сервісу, який показує повідомлення користувачеві та має високий пріоритет, що дозволяє йому продовжувати працювати, навіть якщо інші додатки видаляються з пам'яті.

Хоча WorkManager є потужним інструментом для планування фонових завдань в Android, він має певні обмеження, про які слід знати:

1) WorkManager не дозволяє запускати завдання частіше, ніж раз на 15 хвилин. Це обмеження встановлено на системному рівні, щоб запобігти надмірному навантаженню на ресурси та економити заряд батареї;

2) WorkManager використовує систему черг та пріоритетів для планування завдань. Це означає, що виконання завдання не гарантується в певний час, і воно може бути відкладено або перенесено, якщо система завантажена або не вистачає ресурсів;

3) WorkManager дозволяє вам визначати обмеження для виконання завдань, такі як доступність Wi-Fi, заряд батареї або час доби. Проте, існують певні обмеження щодо того, які обмеження ви можете використовувати. Наприклад, не можна вказати точний рівень заряду батареї, який потрібно для запуску завдання.

Таким чином застосування правильних інструментів дозволяє ефективно використовувати серверні ресурси при роботі з асинхронними задачами за допомогою CronJob. Більше того, серверні відкладені періодичні задачі можливо перенести сторону мобільного застосунку за допомогою WorkManager, проте слід зважати на вищезгадані обмеження.

Список використаних джерел

1. Gogoi Harshjeet Ch., Snehathatha N., Amudha S. (2024). Enhancing Visual Creativity with Neural Style Transfer using Celery Backend Architecture, 56–70. <https://ieeexplore.ieee.org/document/10575471>
2. R. Parlika; P. Atmaja (2020). Realtime monitoring of Bitcoin prices on several Cryptocurrency markets using Web API, Telegram Bot, MySQL Database, and PHP-Cronjob, 115–122. <https://ieeexplore.ieee.org/document/9321109>
3. M. Wambua (2023). Modern Android 13 Development Cookbook: Over 70 recipes to solve Android development issues and create better apps with Kotlin and Jetpack Compose. <https://ieeexplore.ieee.org/document/10251241>

ОТРИМАННЯ ДАНИХ ЩОДО ВИКОРИСТАННЯ ЦПУ ПРОЦЕСОМ В ОПЕРАЦІЙНИХ СИСТЕМАХ НА ОСНОВІ LINUX ПРОГРАМНИМ ЧИНОМ

Ядро Linux надає зручний спосіб для отримання інформації щодо використання ресурсів ОС та процесів, а саме: за допомогою утиліт `top` та `htop`, за допомогою файлів в директорії `/proc`.

`top` – консольна утиліта надає інформацію про загальний стан ОС та процесів користувачів за наступними параметрами: використана віртуальна пам'ять; використана фізична пам'ять; загальний обсяг пам'яті, яку процес ділить з іншими; поточний стан процесу (запущений, призупинений, зомбі); відсоток/значення використовуваного часу ЦПУ; відсоток/значення використаної оперативної пам'яті (ОЗП); тривалість роботи процесу. Надає інтерфейс для сортування процесів по параметрам. Надає можливість завершити або змінити пріоритет процесу.

`htop` – консольна утиліта, що була створена на заміну `top`, оскільки розширює її можливості, а саме:

1) надає інформацію про всі процеси (і системні в тому числі, на відміну від утиліти `top`);

2) має кращий користувацький інтерфейс, що проявляється в: можливості використання миші; фільтрувати, сортувати за багатьма параметрами; кольорове відображення інформації, гнучкість у налаштуванні візуальних тем; наявність деревоподібного режиму перегляду, що дозволяє легко визначити батьківські та дочірні процеси.

Директорія `/proc` – це змонтована псевдофайлова система, що надає інтерфейс до структур даних ядра Linux. Більшість з них доступні лише для читання, але деякі файли доступні на редагування, що дозволяє змінювати змінні ядра.

Для моніторингу використання ресурсів використовують наступні файли:

1) `/proc/stat` – файл, що містить поточний стан ядра/системи. Перший рядок містить агреговану інформацію про ЦПУ, наступні рядки, для кожного ядра ЦПУ. Перелік параметрів, що містяться в рядках опису ЦПУ:

- 1.1) `name` (якщо «сри» – то агреговані дані, якщо «сриN» – то для конкретного ядра ЦПУ;
- 1.2) `user` – час в користувацькому режимі;
- 1.3) `nice` – час в користувацькому режимі з низьким пріоритетом;
- 1.4) `idle` – час очікування задач;
- 1.5) `iowait` – час очікування завершення введення – виведення;
- 1.6) `irq` – час на переривання процесів;
- 1.7) `softirq` – час обслуговування;

2) `/proc/[PID]/stat`, де `[PID]` – ідентифікатор процесу – файл, що містить поточний стан процесу `[PID]`. Перелік параметрів, що містяться у файлі і мають відношення до використання ЦПУ:

- 2.1) `utime` – час в користувацькому режимі;
- 2.2) `stime` – час в режимі ядра;
- 2.3) `cutime` – час в користувацькому режимі дочірніх процесів;
- 2.4) `cstime` – час в режимі ядра дочірніх процесів.

Найпростіший програмний спосіб отримання даних щодо використання ЦПУ ОС або процесом – це зчитувати та парсити дані з файлів директорії `/proc`.

У кожного поля, що описані вище є своя позиція у файлі, кожне значення розділене або символом переходу на новий рядок ('`\n`'), або пробільним символом ('').

Існують показники, які можна порахувати на основі отриманих даних з відповідних файлів:

1) загальний час роботи ЦПУ: сума всіх агрегованих значень “сру”, що описані для /proc/stat;

2) відсоткове відношення завантаження ЦПУ за формулою 1:

$$\frac{\text{загальний час роботи ЦПУ } (\square 1) - \text{час очікування задач } \square\square\square\square(\square 1.4)}{\text{загальний час роботи ЦПУ } (\square 1)} * 100\%; \quad (1)$$

3) загальний час роботи процесу: сума всіх описаних в пункті 2 /proc/[PID]/stat значень;

4) відсоткове відношення використання ЦПУ процесом за формулою 2:

$$\frac{\text{загальний час роботи ЦПУ } (\square 1) - \text{загальний час роботи процесу } (\square 3)}{\text{загальний час роботи ЦПУ } (\square 1)} * 100\%. \quad (2)$$

Отже, ОС на основі Linux надають можливість переглядати інформацію про використання ресурсів ОС та процесів. Також мають легкий доступ отримання цієї інформації програмно використовуючи файлову систему ОС.

Список використаних джерел

1. proc(5) – Linux man page. URL: <https://man7.org/linux/man-pages/man5/proc.5.html>.

Гнатишин М.С., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Недашківський О.Л., д.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

АРХІТЕКТУРА ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ ДЛЯ ВИЯВЛЕННЯ ФЕЙКОВОЇ ІНФОРМАЦІЇ В МЕРЕЖІ ІНТЕРНЕТ НА ОСНОВІ НЕЙРОННИХ МЕРЕЖ

В умовах сучасного інформаційного суспільства проблема дезінформації набуває критичного значення. З поширенням соціальних мереж та онлайн-платформ люди стикаються з величезною кількістю новин, з яких значна частина є недостовірною або навмисно викривленою. Це створює ризики як для окремих користувачів, так і для суспільства загалом, оскільки дезінформація може впливати на громадську думку, політичні рішення та соціальну стабільність [1]. Тому розробка інструментів для автоматичного виявлення неправдивої інформації є надзвичайно актуальною. Використання нейронних мереж у цьому контексті надає нові можливості для аналізу великих обсягів даних, виявлення прихованих закономірностей і підвищення точності оцінки достовірності інформації [2]. Дане дослідження спрямоване на створення архітектури системи, яка здатна аналізувати публікації в реальному часі та надавати користувачам релевантну оцінку правдивості інформації.

На рисунку 1 представлено архітектуру програмного забезпечення для виявлення неправдивої інформації в мережі Інтернет, зокрема у соціальній мережі Twitter, із використанням нейронних мереж.

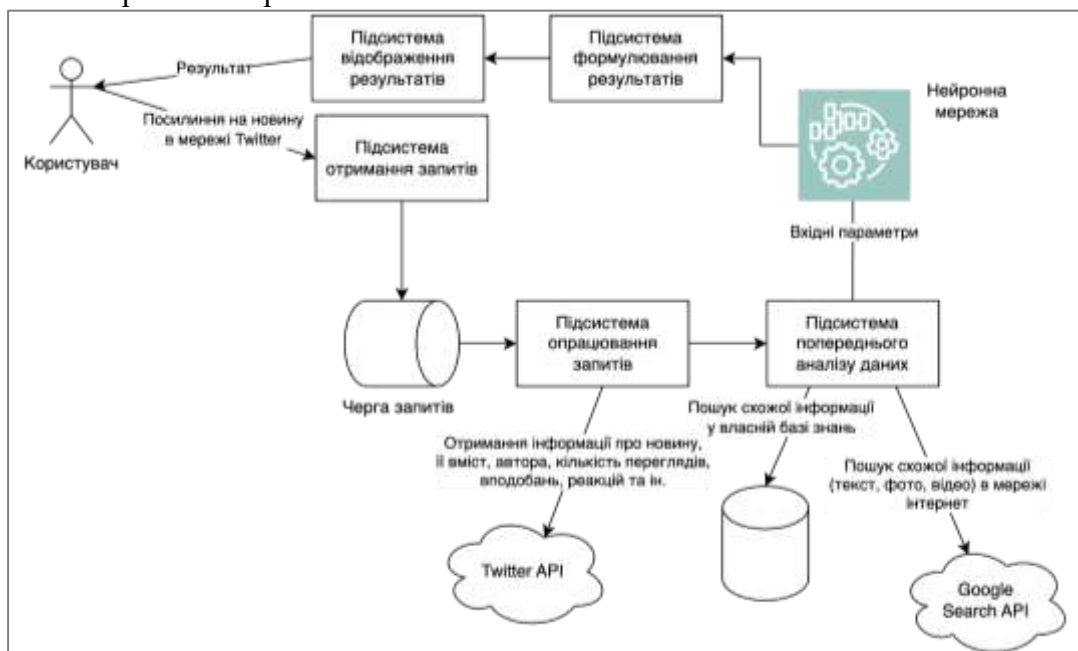


Рисунок 1 – Модель архітектури програмного забезпечення на прикладі опрацювання новини з соціальної мережі Twitter

Архітектура складається з кількох взаємопов'язаних підсистем, що забезпечують процес збору, аналізу та інтерпретації даних, а також формування й подачу результатів користувачу. Основним елементом є нейронна мережа, яка виконує кінцевий аналіз вхідних даних.

Опис компонентів системи:

- 1) підсистема отримання запитів відповідає за отримання запитів від користувача. Запити містять посилання на новини або повідомлення в соціальній мережі Twitter. Після цього запити зберігаються в черзі для подальшої обробки;
- 2) черга запитів слугує буфером для зберігання запитів, що надійшли. Це дозволяє

забезпечити стабільність системи, оскільки запити обробляються послідовно і в оптимальний час, що знижує ризик перевантаження системи;

3) підсистема обробки запитів виконує первинну обробку отриманих запитів. Залучає Twitter API для отримання детальної інформації про новини, включаючи текст, автора, кількість переглядів, вподобання, реакції тощо [3]. Цей етап дозволяє підготувати необхідні дані для подальшого аналізу;

4) підсистема попереднього аналізу даних виконує пошук схожої інформації в базі знань та через Google Search API для аналізу інших джерел у мережі Інтернет. Цей крок важливий для визначення, чи існує підтвердження чи спростування отриманої інформації з боку інших джерел;

5) нейронна мережа це основний компонент системи, який використовує вхідні параметри (зокрема інформацію з Twitter і пошукових запитів) для здійснення глибокого аналізу інформації. Нейронна мережа навчається на основі великої кількості даних для визначення ознак правдивості або неправдивості інформації [4];

6) підсистема формулювання результатів відповідає за обробку результатів аналізу, здійсненого нейронною мережею. Вона формує підсумковий висновок, який може включати ймовірність неправдивості інформації та інші метрики надійності;

7) підсистема відображення результатів це останній компонент архітектури, який надає користувачу результати аналізу. Виводиться підсумкова інформація про надійність новини або публікації, а також можливі посилання на додаткові джерела, що підтверджують чи спростовують її правдивість.

Розроблена архітектура є модульною і забезпечує комплексний підхід до виявлення фейкової інформації, об'єднуючи різні джерела даних і методи аналізу для покращення точності результатів. В рамках проведення подальших досліджень планується порівняння ефективності конкретних інструментів для забезпечення найкращої швидкодії та коректності результатів.

Список використаних джерел

1. Аналіз методів та програмного забезпечення для виявлення неправдивої інформації у мережі інтернет на основі нейронних мереж / Гнатишин М. С., Недашківський О. Л. // Зв'язок – 2024 – Вип. 4. – С. 52–57 – URL: <https://doi.org/10.31673/2412-9070.2024.045257>
2. How does fake news spread? Understanding pathways of disinformation spread through APIs / Ng L. H. X., Taeihagh A. // Policy Internet. – 2021. – Вип. 13. – С. 560–585. – URL: <https://doi.org/10.1002/poi3.268>.
3. Fake news detection on Twitter / Sharma S., Saraswat M., Dubey A. K. // International Journal of Web Information Systems. – 2022. – Вип. 18. – С. 388–412. – URL: <https://doi.org/10.1108/IJWIS-02-2022-0044>.
4. Dynamic graph neural network for fake news detection / Song C., Teng Y., Zhu Y., Wei S., Wu B. // Neurocomputing. – 2022. – Вип. 505. – С. 362–374. – URL: <https://doi.org/10.1016/j.neucom.2022.07.057>

Ситнік І.О., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Шпурик В.В., к.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

МЕТОД ПОБУДОВИ ІМОВІРНІСНИХ МОДЕЛЕЙ УПРАВЛІННЯ ПРОЦЕСОМ РОЗРОБКИ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ В УМОВАХ НЕВИЗНАЧЕНОСТІ

Статистика успішності проєктів розробки програмного забезпечення (ПЗ) за останні роки демонструє доволі низький відсоток успішності, що формує простір для вдосконалення технологій управління проєктами розробки [1] [3].

Підвищення успішності проєктів передбачає визнання важливості надійного менеджменту протягом усього процесу розробки. Однак, процес розробки ПЗ є складним із великою кількістю нюансів, що можуть, як впливати на його результати, так і один на одного. Запропоноване в цій роботі рішення – це використання ймовірнісного підходу для керування проєктами розробки ПЗ.

Ймовірнісний підхід надає моделі процесу розробки ймовірнісний характер, для пошуку ймовірності величин, які є найбільш суттєвими для прийняття рішення керівником проєкту.

Вірогіднісні міркування – це підхід, в якому релевантні знання предметної області використовується для прийняття рішень в умовах невизначеності. Системи вірогіднісних міркувань мають дві складові: модель та алгоритм виводу.

Алгоритм виводу може бути не унікальним для даної системи і, наразі, існує багато таких алгоритмів, що використовуються на практиці, тому при створенні системи можливо використання вже існуючих алгоритмів виводу.

Ймовірнісні моделі відображають релевантні знання про предметну область, вірогідності подій та зв'язки між ними. Ці моделі є унікальними для кожної предметної області. У випадку проєкту розробки ПЗ ця модель може містити загальні знання про технології, що можуть використовуватися, вірогіднісні залежності між ними та бюджетом, термінами, метриками продуктивності команди тощо. Тому, алгоритм побудови моделі – головне питання у спробі виразити процес розробки ПЗ у вигляді системи вірогіднісних міркувань. На даний момент існує проблема, яка полягає у відсутності практичних ймовірнісних моделей процесу створення ПЗ.

Системи вірогіднісних міркувань надають можливість виконувати міркування трьох видів:

- 1) прогнозування найімовірнішого розвитку подій: алгоритм виводу використовує модель та факти про поточну ситуацію для розрахунку ймовірностей;
- 2) вивід ймовірних причин подій: використовується та ж інформація та алгоритм виводу, що і при прогнозуванні, але, також, подаються результати подій для виводу вірогідностей їх причин;
- 3) навчання моделі: враховуються факти та результати минулого досвіду для більш точного підрахунку ймовірностей [2].

Оскільки між метриками продуктивності команди, різними рішеннями в ході проєкту та результатами розробки ПЗ існує статистична відповідність, але відсутній причинно-наслідковий зв'язок, то вважається доцільним представити процес розробки ПЗ у вигляді ймовірнісної моделі.

Вірогіднісне програмування, в свою чергу, надає гнучкі засоби реалізації вірогіднісних міркувань та мов програмування для розробки систем вірогіднісних міркувань, їх моделей та використання цих систем на практиці. В порівнянні з іншими мовами представлення (такі як мережі Байєса та Марковські мережі), мова програмування є набагато гнучкішою та більш динамічною, а елементи мов програмування загального призначення надають засоби інтеграції з існуючими системами менеджменту, такі як Atlassian Jira.

Отже, дослідження зосереджене на розробці алгоритму створення ймовірнісних моделей розробки ПЗ. Які саме можна виділити елементи в даному процесі, які між ними є залежності та що можна віднести до знань моделі або до її параметрів, які визначаються в ході навчання – є одними із головних питань, а використання вірогіднісного програмування надасть практичні засоби реалізації та використання таких моделей.

Запропонований ймовірнісний підхід зможе надати можливості знаходження оптимального шляху розробки ПЗ та з'ясування можливих причин провалу проєкту. В результаті дослідження очікується створення методу побудови ймовірнісних моделей розробки ПЗ та створення архітектури ПЗ, що використовує елементи ймовірнісного програмування.

Список використаних джерел

1. Palumbo S., Rehberg B., Li H. (2024, 30 квітня). Software Projects Don't Have to Be Late, Costly, and Irrelevant. BCG Global. <https://www.bcg.com/publications/2024/software-projects-dont-have-to-be-late-costly-and-irrelevant>
2. Pfeffer A. (2016). Practical Probabilistic Programming. Manning Publications
3. Software Survival in 2024: 2023 Statistics and Quality Assurance. (б. д.). Beta Breakers. <https://www.betabreakers.com/project-failure-statistics-and-the-role-of-quality-assurance/>

Логвиненко Т.С., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Гагарін О.О., к.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

ОЦІНЮВАННЯ ЯКОСТІ МОДЕЛЮВАННЯ ГІДРОАКУСТИЧНИХ СИГНАЛІВ ПРОЦЕСУ МОНІТОРИНГУ У РАЙОНАХ СПОСТЕРЕЖЕННЯ

У наслідок науково-технічного прогресу стрімко змінюються та вдосконалюються технології пов'язані з використанням моделювання як інструменту отримання нових знань щодо об'єктів які досліджуються. Одним з таких напрямків досліджень є системи моделювання та проведення комп'ютерних експериментів дослідження гідроакустичних процесів (ГАП) які потребують наявності процедур оцінювання якості моделей у процесі аналізу результатів проведених експериментів та мають яскраво виражену специфічну спрямованість або за ціллю експерименту або за видом моделей що використовуються.

Основною привабливою властивістю таких систем є здатність накопичення, покрокове вироблення та повторне використання знань (формалізованих посилок, висновків, закономірностей, що відносяться до досліджуваного об'єкта) в режимі реального процесу моделювання поведінки складної технічної системи (СТС).

Для систем такого класу потрібно створити різнопланові засоби такі як: інструментарій використання різнопланових моделей, розробити систему управління на основі гнучкого налаштування умов та цілій експерименту, побудувати середовища ідентифікації умов проведення експерименту, наприклад, район проведення, параметри акваторії, набір приладів виявлення сигналів (їх розташування, характеристики). При цьому однією з найбільш складних та найменш вивченою є проблема оцінювання якості моделей та отриманих з їх допомогою якісних результатів.

Тому великого значення набуває процес оцінювання якості отриманих результатів та виявлення показників такого оцінювання. Слід відмітити що проблема оцінювання ускладнюється ще тим що потрібно оцінювань сам процес моделювання та схеми (сценарії) його проведення, а також собственно самі моделі що його реалізують. За результатами такого оцінювання формується узагальнена якісна оцінка проведених експериментів.

Для упорядкування та ефективної роботи з даними показниками якості потрібно створити кластерний простір з показників оцінювання. Цей простір побудуємо у вигляді трьох вимірного за основним напрямками оцінювання сховища, а саме кластери: придатність G, оптимальність O, перевага S.

Групи властивостей показників оцінювання якості процесу моделювання та моделей які використовуються при проведенні експериментів формують загальну оцінку якості як суму часткових показників за групами. Пошук такої оцінки може бути побудовано на основі використання методів які базуються на побудові моделей у вигляді нейронних мереж.

Пошук показника необхідної якості об'єкту задається умовами або вимогами, яким повинні задовольняти можливі значення показника його якості.

Критерії перерахованих вище груп G, O та S можна сформулювати таким чином [1].

Критерій Придатності G – формула 1:

$$G : \cap_{i=1}^m (k_{ij} \in \{k_i^A\}) \approx U, j \in (1, 2, \dots, n), \quad (1)$$

де U – достовірна подія.

Критерій Оптимальності O – формула 2:

$$O : \bigcap_{i=1}^m (k_{ij} \in \{k_i^A\}) \approx U, \bigcup_{\forall l \in \{l\} m0} (k_{lj} = k_l^{opt}) \approx U(2) j \in (1, 2, \dots, n), m0 \in [1(1)m]. \quad (2)$$

Критерій Переваги S – формула 3:

$$S : \bigcap_{i=1}^m (k_{ij} \in \{k_i^D\}) \bigcap_{j=l}^n \bigcap_{i=l}^m (k_{il} \geq k_{ij}) \approx U, l \in (1, 2, \dots, n) . \quad (3)$$

Дани всіх трьох груп, які визначають показники оцінювання якості є висловлюваними формами (або умовними операторами визначення значення якогось показника якості) , які стають висловами (дійсними) після привласнення змінним деяких значень. Якщо одержаний вислів є істинним (значення дворівнево 1), то вважають, що об'єкт за якістю задовольняє даному критерію; якщо вислів буде помилковим (значення дворівнево 0), то об'єкт не задовольняє за якістю обраному критерію.

Для вирішення задачі розрахунку комплексної оцінки якості системи моделювання за вказаними умовами потрібно сформувати та включити до відповідної кластерної групи набори даних. Ці процедури виконуються спеціальним додатковим пошуковим запитами та наповнюють відповідні набори даних.

У роботі магістра пропонується побудувати веб платформу, яка формує дані кластерів та має засоби для їх використання в розрахункових алгоритмах.

Список використаних джерел

1. Соколов Б. В., Юсупов Р. М. Концептуальні основи оцінювання й аналізу якості моделей полімодельних комплексів // Теорія та системи управління. – 2004. – № 6. – С. 5–16.

Руй Д.І., Національний технічний університет України
 «Київський політехнічний інститут імені Ігоря Сікорського»
 Донець А.Г., к.ф.-м.н., Національний технічний університет України
 «Київський політехнічний інститут імені Ігоря Сікорського»
 Колумбет В.П., д.ф., Національний технічний університет України
 «Київський політехнічний інститут імені Ігоря Сікорського»

АНАЛІЗ ІНСТРУМЕНТІВ СТВОРЕННЯ ВЕБ-ПЛАТФОРМ ПІДТРИМКИ ЕКОЛОГІЧНОГО МЕНЕДЖМЕНТУ

Постановка проблеми. Екологічні питання стали одними з найактуальніших викликів, з якими стикається сучасне суспільство, що потребує термінового вирішення для забезпечення сталого розвитку. Дослідження вказують на те, що екологічні проблеми можуть суттєво загрожувати економічному розвитку, оскільки забруднення навколишнього середовища та неефективне використання ресурсів призводять до збільшення витрат на охорону здоров'я, зниження продуктивності та витрат на відновлення екосистем. Впровадження новітніх екологічних технологій у промисловість є необхідним кроком для досягнення сталого розвитку. Програми з утилізації та переробки відходів, які реалізуються на рівні підприємств і держави, сприяють не лише зменшенню забруднення, а й економії ресурсів. Застосування інноваційних рішень у цій сфері забезпечує не тільки екологічні переваги, але й економічну вигоду, що підтверджує доцільність інвестицій у зелені технології [1].

У підсумку, екологічні питання вимагають комплексного підходу, що включає в себе активну участь усіх рівнів суспільства. Створення сучасних розподілених програмних комплексів вимагає інтеграцій з іншими інформаційними ресурсами і передбачає необхідність у використанні функціоналу готових програмних модулів та мережевих служб, які надають різні партнерські ІТ-компанії. Враховуючи потенційні можливості API та Web API для веброзробки є необхідність у створенні та впровадженні таких сервісів компаніям і державним закладам у своїх програмних продуктах. Дана технологія є досить актуальною і важливою для взаємодії, обміну й обробки даних між різним програмним забезпеченням.

Аналіз останніх досліджень. Прикладний програмний інтерфейс (API) – це програмний інтерфейс (сервісний контракт), що дозволяє двом програмним компонентам обмінюватися даними один з одним, використовуючи набір визначених функцій та протоколів. Даний інтерфейс визначає як компоненти взаємодіють один з одним, використовуючи запити та відповіді. На відміну від інтерфейсу користувача, який приймає дані від користувачів, передає їх додатку для обробки та повертає результати користувачеві, API не взаємодіє з користувачем, а виконує обмін та обробку інформації між різними частинами програмного забезпечення, отримуючи дані від одного модуля програми і передаючи результати назад в інший модуль. Аналізуючи різноманітні архітектури API, було виконано їх порівняння за визначеними критеріями, що наведено в таблиці 1. Для виконання поставленого завдання в даній роботі, враховуючи легку складність, роботу з ресурсами, підтримкою HTTP-методів та великою різноманітністю можливих форматів даних, найбільш доцільним буде використовувати рішення на основі архітектури REST. [2, 3]

Таблиця 1 – Порівняння архітектурних стилів API

Критерій API	SOAP	GraphQL	RPC	REST
Принцип специфікації	Передача повідомлення у визначеному форматі	Передача схеми та системи типів	Виклик процедури	Передача повідомлення у вигляді ресурсу
Формат даних	XML	JSON	JSON, XML, Protobuf, Thrift, FlatBuffers	XML, JSON, YAML, HTML, текст

Продовження таблиці 1

Критерій API	SOAP	GraphQL	RPC	REST
Складність вивчення	Важка	Легка	Середня	Легка
Спільнота користувачів	Мала	Велика	Зростаюча	Зростаюча
Застосування	Платіжні системи, приватні API, фінансові та телеком-сервіси, старі системи	Мобільні застосунки, складні системи, мікросервіси	Комбіновані API, мікросервісні системи	Публічні API, Вебсайти, хмарні сервіси

ASP.NET Web API використовується для створення простих вебслужб HTTP, створення вебсервісів, які є легкими, придатними для обслуговування та масштабованими, а також підтримують обмежену пропускну здатність, доступ до службових даних у веб-застосунках. [3]

Роль технологій та інновацій у вирішенні екологічних питань набуває особливого значення в умовах сучасного світу, де екологічні виклики стають все більш складними. Інновації в сферах енергетики, транспорту, сільського господарства та управління відходами відкривають нові горизонти для ефективного використання ресурсів і зменшення забруднення. Наприклад, розробка та впровадження смарт-технологій у системи енергетичного управління дозволяє не лише оптимізувати споживання електроенергії, але й інтегрувати відновлювальні джерела енергії, такі як сонячні панелі та вітрові електростанції, у загальну енергетичну мережу. Це сприяє не тільки зменшенню викидів парникових газів, але й знижує залежність від традиційних видів пального.

Висновки. В наведених тезах продемонстровано, що на цей час самим поширеним і масовим засобом інтегрувати такі програмні модулі та мережеві служби в межах розширення екологічного проекту як глобальної цінності, що надаються різними підприємствами, є механізм веб-служб платформи WebAPI, доступних за стандартними мережевими протоколами HTTP/HTTPS.

Як приклад ефективного застосування механізму веб-служб в кваліфікаційної роботи створено базову версію веб-служби WebAPI з підтримкою основних операцій управління бази даних MSSQL, що містить описи проведення хімічних аналізів речовин, а також клієнт до цієї веб-служби у вигляді веб-програми на ASP.NET Core.

Окрім того, інновації в галузі інформаційних технологій, такі як великі дані та штучний інтелект, стають потужними інструментами для моніторингу екологічного стану та прогнозування змін у навколишньому середовищі. Завдяки цим технологіям, підприємства можуть отримувати точну інформацію про вплив своїх дій на екологію, виявляти ризики та своєчасно реагувати на них.

Список використаних джерел

1. Сліпченко В. Г., Полягушко Л. Г., Круш О. Є. Система комплексного еко-енерго-економічного моніторингу для оптимізації управлінських рішень (області, району та міста). Вісник східноукраїнського національного університету імені Володимира Даля. 4 (268), 2021. С. 13–20. DOI: <https://doi.org/10.33216/1998-7927-2021-268-4-13-20>
2. Порівняння архітектурних стилів SOAP, REST, GraphQL, RPC [Електронний ресурс] – Режим доступу до ресурсу: <https://www.altexsoft.com/blog/soap-vs-rest-vs-graphql-vs-rpc/>
3. Вступ до ASP.NET Web API [Електронний ресурс] – Режим доступу до ресурсу: <https://dotnettutorials.net/lesson/web-api-architecture/>

Дембіцький В.В., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Коваль О.В., д.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

ПРОБЛЕМАТИКА ОДНОПУНКТНОЇ ТЕХНОЛОГІЇ ІДЕНТИФІКАЦІЇ ДЖЕРЕЛ ГІДРОАКУСТИЧНИХ СИГНАЛІВ

Проблематика однопунктної ідентифікації відіграє важливу роль у вдосконаленні систем виявлення та локалізації джерел гідроакустичних сигналів. Однак, щоб забезпечити ефективність і точність цих систем, необхідно детально вивчити ключові обмеження технології, такі як обмеженість просторової інформації, стійкість до завад та вплив середовища на поширення звуку. Однопунктна технологія дозволяє здійснювати ідентифікацію джерела за допомогою лише одного сенсора, що спрощує конструкцію системи. Розуміння проблематики та її вирішення сприятиме розробці нових підходів і алгоритмів, які підвищать точність виявлення та адаптацію системи до реальних умов.

Обмеження просторової інформації. Найбільша проблема однопунктної технології полягає в тому, що використання лише одного датчика чи гідрофона не дозволяє визначити точні координати джерела сигналу на основі простої триангуляції. У системах з кількома сенсорами, кожен з яких отримує сигнал з різних точок, можливе визначення напряму і відстані до джерела. Але однопунктна система позбавлена цієї можливості, тому їй доводиться покладатися виключно на аналіз часових, частотних та інших характеристик сигналу, що робить процес ідентифікації та визначення місцезнаходження значно складнішим і менш точним.

Слабка стійкість до завад. Підводне середовище характеризується високим рівнем акустичних завад і шумів, зокрема від природних (течії, хвилі, морські тварини) та техногенних (промислові судна, підводні пристрої) джерел. В умовах таких завад стає складно точно виділити сигнал від потрібного джерела, особливо якщо воно є малопотужним або віддаленим. Це потребує розробки складних алгоритмів шумопридушення та обробки сигналу.

Вплив фізичних факторів на поширення звуку. У воді звукові хвилі піддаються впливу температури, тиску, солоності, що змінює їхню швидкість і напрямок. Це викликає заломлення сигналу, а також його відбивання від дна і поверхні, що призводить до спотворень. Такі зміни ускладнюють точне визначення місцезнаходження джерела і вимагають корекції моделей поширення звуку для відображення цих змін.

Обмежені можливості класифікації. У випадках, коли джерела сигналів подібні між собою, наприклад, мають близькі частотні характеристики, однопунктна система може мати труднощі у точному їх розрізненні. Для підвищення точності класифікації використовуються додаткові методи, наприклад, машинне навчання, однак і вони потребують великих обсягів навчальних даних, що не завжди можливо в реальних умовах.

Високі вимоги до точності сигналів і алгоритмів обробки. Відсутність просторової інформації накладає жорсткі вимоги до якості обробки сигналу, що включає спектральний аналіз, часово-частотний аналіз та інші методи. Кожен із цих методів потребує точного налаштування для досягнення потрібної чутливості до змін у сигналі, а також обчислювальних ресурсів для їх реалізації у реальному часі.

Складність інтеграції в існуючі системи. Однопунктна технологія ідентифікації може потребувати специфічних налаштувань і модифікацій для роботи з уже існуючими системами моніторингу та контролю. Це вимагає додаткових досліджень та адаптації технології для забезпечення її сумісності з іншими сенсорами, джерелами даних та апаратними платформами.

Список використаних джерел

1. Кулішенко О. П. Методи ідентифікації акустичних джерел на основі обробки сигналів: монографія. – Харків: Наукова думка, 2018. – 340 с.
2. Мельник В. Ю., Сидоренко Л. М. Аналіз поширення звукових хвиль у підводному середовищі для систем ідентифікації. // Гідроакустичні дослідження та технології. – 2020. – № 3. – С. 22–29.
3. Попов А. М., Гончаренко П. В. Алгоритми класифікації гідроакустичних сигналів в умовах завад. // Прикладні наукові дослідження в сфері обробки сигналів. – 2019. – Т. 17, № 2. – С. 12–18.
4. Brown D. R. Signal Processing Techniques for Single-Point Identification Systems. – London: Wiley, 2019. – 275 p.

Фішер О.Є, Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Барабаш О.В., д.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

МЕТОДИ ОПТИМІЗАЦІЇ МОДЕЛЕЙ КОМП'ЮТЕРНОГО ЗОРУ ДЛЯ РОЗПІЗНАВАННЯ ЗОБРАЖЕНЬ НА АВТОНОМНИХ ПРИСТРОЯХ

Комп'ютерний зір може виражатися у вигляді різних задач: детекція і класифікація об'єктів, сегментація зображень, відстеження об'єктів тощо. Кожен тип задач потребує нейронної мережі певної архітектури, натренованої під цей тип задач. Для детекції та класифікації зазвичай використовують конволюційні нейромережі, а для відстеження – рекурентні [1].

Використання нейронних мереж співіснує зі значними вимогами до обчислювальної потужності, паралельності обробки даних і енерговитрат. Ефективність архітектури, використаних моделей, кількість параметрів, якість тренувального сету – все це впливає не тільки на точність виконання задач, покладених на мережу, але на швидкість і необхідну потужність, а значить і енерговитратність мережі.

Невеликі, вузькоспеціалізовані й енергоефективні мережі для розв'язання задач комп'ютерного зору вже багато років можна зустріти на мобільних пристроях, лептопах, роботах і камерах відеоспостереження, де вони допомагають розпізнавати людей, обличчя, виконувати базові задачі розпізнавання тексту і сегментації зображень, розпізнавання об'єктів. Проте ефективна робота таких мереж в реальному часі відносно обмежена і деякі задачі потребують надсилання запитів до хмарних обчислювальних потужностей.

Фінальну ефективність системи можна описати такими параметрами [2] як:

- 1) швидкість розпізнавання: скільки часу у нейронної мережі займає виконання обробки сигналу;
- 2) точність розпізнавання: наскільки результат розпізнавання відповідає дійсності;
- 3) апаратне забезпечення: які обчислювальні потужності є в наявності для роботи нейронної мережі.

Більшість портативних пристроїв, що використовують комп'ютерний зір, мають доступ до хмарних обчислювальних потужностей, що дозволяє покладатись не лише на обчислювальну потужність самого пристрою. Проте, деякі пристрої мають у своєму розпорядженні лише обмежені ресурси бортового комп'ютера. Незалежність таких пристроїв від зовнішніх обчислювальних потужностей дозволяє використовувати системи побудовані на основі таких пристроїв в умовах обмеженого доступу до мережі інтернет [3].

Іншою проблемою є фізичні обмеження швидкості передачі даних. Комп'ютерний зір, що працює на самому пристрої не буде мати проблему зі швидкістю передачі даних на аналіз хмарній нейронній мережі. Це дозволяє пристрою швидше реагувати на зміну ситуації навколо нього.

Запропонована до розробки система буде фокусуватись на задачі розпізнавання об'єктів на зображенні при максимізації параметра швидкості розпізнавання та мінімізації використання обчислювальних потужностей. Точність розпізнавання має зберігатись в оптимальних рамках визначених експериментально.

Серед методів оптимізації розроблюваної системи пропонуються такі параметри:

- вибір оптимальної кількості параметрів нейронної мережі;
- оптимізація кількості шарів;
- вибір оптимальної функції активації;
- вибір оптимальної функції втрати;
- вибір оптимальної архітектури нейронної мережі, або вдосконалення наявних.

Використовуючи розроблювальну нейронну мережу, оптимізовану під обмежені потужності автономних пристроїв, уможливилося використання її на крайових пристроях, використовуючи лише крайові обчислення для прийняття рішень у режимі реального часу.

Використовуючи камеру пристрою, запропонована до розробки нейронна мережа дозволить використовувати розпізнану візуальну інформацію для прийняття рішень на автономних пристроях без використання додаткових обчислювальних потужностей, або без наявності оператора.

Список використаних джерел

1. Szeliski R. (2022) Computer Vision: Algorithms and Applications [2nd ed.]. (p. 291–320).
2. Bochkovskiy A., Wang C., Mark Liao H. (2020) YOLOv4: Optimal Speed and Accuracy of Object Detection, arXiv:2004.10934.
3. Douglas C. (2024) AI at the Edge vs. AI in the Cloud: A Comparative Analysis of High-End and Edge AI Systems, 10.13140/RG.2.2.18865.80488.

Гейко О.О., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Варава І.А., к.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

ІНСТРУМЕНТИ ПРЕДСТАВЛЕННЯ МАТЕМАТИЧНИХ ФОРМУЛ В ПРОГРАМНОМУ ЗАБЕЗПЕЧЕННІ

Математична модель – це спрощене представлення реального об’єкта, процесу або явища за допомогою математичних формул, рівнянь та інших математичних інструментів. Це як створити малюнок, який схожий на справжній предмет, але не має всіх дрібних деталей. В сучасному світі математичні моделі можуть лягти в основу цифрових моделей і комп’ютерного моделювання.

Комп’ютерне моделювання є універсальним інструментом, який знаходить застосування в багатьох сферах людської діяльності. Воно дозволяє нам краще розуміти складні системи, приймати обґрунтовані рішення та розробляти нові технології.

Виходячи із визначення математичної моделі можна зробити висновок, що вона містить певні припущення та спрощення, які можуть не повністю відображати всі нюанси реального процесу або явища. З огляду на це, виникає потреба в виконанні її верифікації.

Верифікація математичних моделей є невід’ємною частиною сучасного наукового процесу. Вона забезпечує довіру до отриманих результатів і дозволяє ефективніше використовувати математичні моделі для вирішення реальних задач. Але такі перевірки математичних моделей зазвичай передбачає проведення великої кількості розрахунків, порівняння результатів моделювання з експериментальними даними, а також візуалізацію результатів. Тому розробка автоматизованих систем для верифікації моделей є актуальним завданням, яке дозволить підвищити швидкість і точність наукових досліджень.

Однією із задач які можна розглянути в розрізі розробки програмного забезпечення для верифікації математичних моделей – це отримання даних від користувача. Безсумнівно, для користувача набагато легшим буде надавати математичну модель у вигляді формул та рівнянь. В зв’язку з цим і було проведено огляд існуючих нотацій та бібліотек. Це необхідно для створення зручного інтерфейсу та забезпечення точності результатів. Коли користувач вводить рівняння або математичні моделі, система повинна коректно розпізнати, зчитати та відобразити вираз, а також виконати відповідні обчислення без втрати сенсу та точності.

На даний момент існує декілька способів задання математичних формул, один з них – LaTeX. LaTeX (вимовляється «латех») — мова розмітки даних та пакет макросів TeX для високоякісного оформлення документів. Вважається стандартом де-факто для підготовки математичних і технічних текстів для публікації в наукових виданнях. Був створений Леслі Лампортом (англ. Leslie Lamport) на початку 1980-х років [1].

На рисунку 1 наведено приклад використання розмітки LaTeX та її відображення.

<code>\lim\limits_{x \to \infty} \exp(-x) = 0</code>	$\lim_{x \rightarrow \infty} \exp(-x) = 0$
--	--

Рисунок 1 – Приклад розмітки LaTeX

Не дивлячись на високу популярність розмітки LaTeX для відображення математичних формул, не існує прямих способів її інтерпретування в мовах програмування. Хоча, в мові програмування Python, ця задача вирішується за допомогою використання декількох бібліотек.

PyLaTeX – це бібліотека на Python для створення та компіляції файлів LaTeX. Мета цієї бібліотеки – стати простим, але розширюваним інтерфейсом між Python і LaTeX (рисунк 2).

$$\begin{pmatrix} 2 & 3 & 4 \\ 0 & 0 & 1 \\ 0 & 0 & 2 \end{pmatrix} \begin{pmatrix} 100 \\ 10 \\ 20 \end{pmatrix} = \begin{pmatrix} 310 \\ 20 \\ 40 \end{pmatrix}$$

Рисунок 2 – Із Python в LaTeX

Для зворотного ж перетворення необхідно буде використати бібліотеки latex2sympy та sympy (рисунок 3). latex2sympy: Ця бібліотека може безпосередньо аналізувати вирази LaTeX в об'єкти SymPy, що полегшує роботу з математичними виразами в Python.[3]

```
from latex2sympy2 import latex2sympy
from sympy import Symbol, Integral, sin, cos, exp

latex_expr = r"\int_{-\infty}^{\infty} \sin(x) \cdot e^{-x^2} \cdot dx"
sympy_expr = latex2sympy(latex_expr)
# Define symbols
x = Symbol('x')
# Create a mathematical expression
expr = Integral(sin(x)*exp(-x**2), (x, 1, 10))
# Evaluate or manipulate the expression
result = expr.evalf(10) # Numerical evaluation

print(result)
```

✓ 0.1s

0.1297382013

Рисунок 3 – Із LaTeX в Python

Отже, використання засобів інтерпретації математичних формул є важливим кроком для створення якісного, надійного та інтуїтивно зрозумілого програмного продукту, який підтримує наукові та інженерні обчислення. Це забезпечує точність, легкість у використанні та підтримку складних обчислень, дозволяючи користувачам працювати зі звичним для них форматом введення та отримувати коректні результати. Інтерпретація не лише допомагає уникнути помилок, але й значно розширює функціональність продукту, надаючи можливості для зростання та покращення відповідно до потреб користувачів.

Список використаних джерел

1. Oetiker T., Partl H., Hyna I., Schlegl E. (2002). Не надто короткий вступ до LATEX2ε (М. Поляков, Пер.). (Оригінал опубліковано 2002 р.) <https://web.archive.org/lshortuk.pdf>.
2. PyLaTeX – PyLaTeX 1.4.2 documentation. (б. д.). GitHub Pages. <https://jeltef.github.io/PyLaTeX/current/>.
3. latex2sympy2. (б. д.). PyPI. <https://pypi.org/project/latex2sympy2/>.

Романів Р.С., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Бандурка О.І., д.ф., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

ЗАСТОСУВАННЯ БЛОКЧЕЙН ТЕХНОЛОГІЇ ДЛЯ ЗАБЕЗПЕЧЕННЯ ФУНКЦІОНАЛЬНОЇ СТІЙКОСТІ РОЗПОДІЛЕНИХ ІНФОРМАЦІЙНИХ СИСТЕМ МОНІТОРИНГУ РУХУ ТРАНСПОРТНИХ ЗАСОБІВ

У минулому технологія блокчейн застосовувалася з великим успіхом у сфері безпеки, і було показано, що вона є ефективним засобом полегшення поширення інформації між кількома пов'язаними системами. Через ці особливості автори [1] провели дослідження, яке намагалося інтегрувати технологію блокчейну з існуючими системами перевірки трафіку CAV, щоб швидко й ефективно захистити інформацію про транспортні засоби та усунути оманливу інформацію, якою обмінюються шкідливі транспортні засоби [1].

Щоб протестувати запропоновану систему, команда відстежувала кількість зловмисників і здатність системи на основі блокчейну виявляти, коли користувачі генерують шкідливу інформацію, причому технологія блокчейну використовує систему на основі репутації для відстеження надійності певних користувачів на основі їхні минулі дії [2].

Отримані результати показали, що ця система була дуже ефективною у відрізненні звичайних даних від даних, створених зловмисниками [3]. Його успіх посилює сферу особливої важливості для використання технології блокчейну в CAV: пом'якшення атак і безпека користувачів. Грунтуючись на цих результатах, можна сказати, що технологія блокчейн має потенціал для значного розвитку багатьох аспектів роботи та надійності CAV, які раніше стосувалися питань.

Діаграма на рисунку 1 представляє процес взаємодії між відправниками та перевізниками в системі управління транспортними перевезеннями з використанням блокчейн технології.

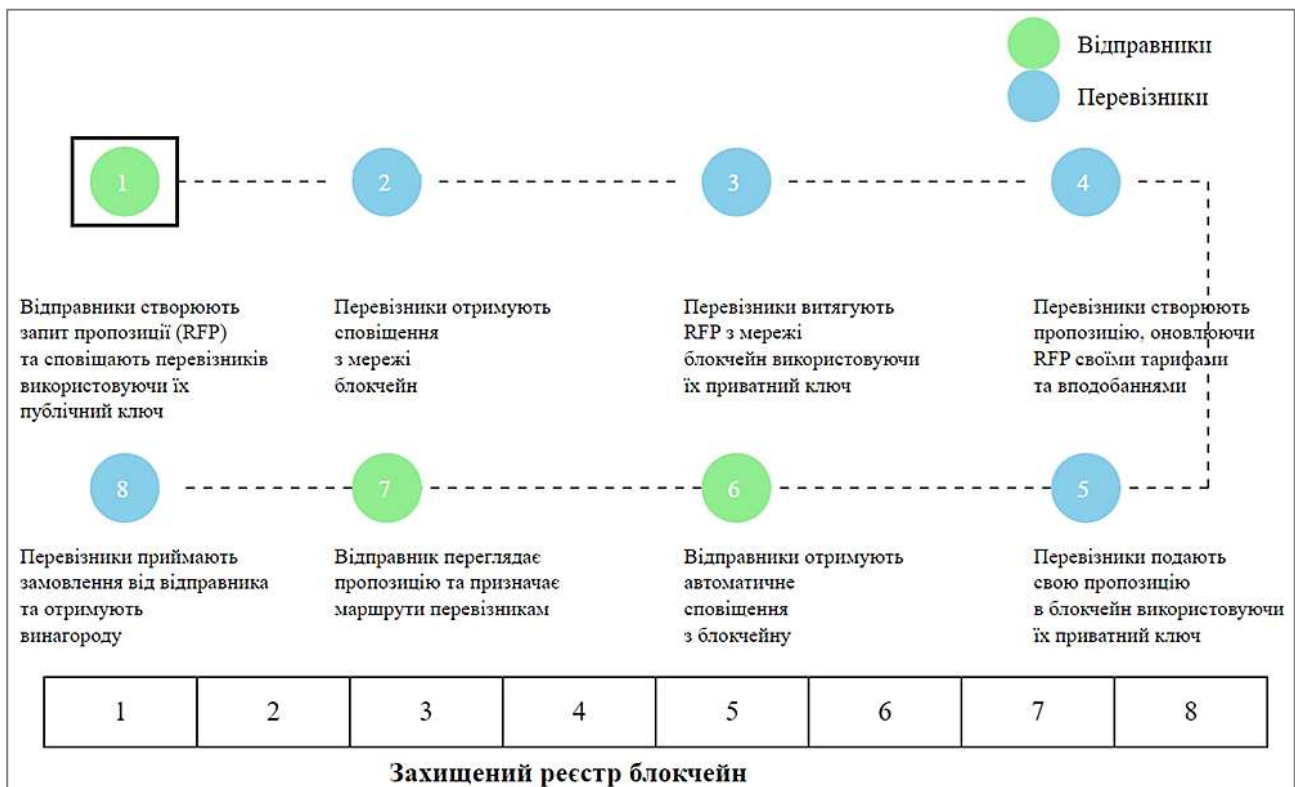


Рисунок 1 – Діаграма призначення завдань для перевізників

Процес складається з 8 послідовних кроків:

- 1) початок процесу: відправники створюють запит пропозиції (RFP – Request for Proposal) та сповіщають перевізників, використовуючи їх публічний ключ;
- 2) етап отримання: перевізники отримують сповіщення з мережі блокчейн;
- 3) етап вилучення: перевізники витягують RFP з мережі блокчейн, використовуючи свій приватний ключ;
- 4) етап створення пропозиції: перевізники створюють пропозицію, оновлюючи RFP своїми тарифами та вподобаннями;
- 5) етап подання: перевізники подають свою пропозицію в блокчейн, використовуючи їх приватний ключ;
- 6) етап сповіщення: відправники отримують автоматичне сповіщення з блокчейну;
- 7) етап перегляду: відправник переглядає пропозицію та призначає маршрути перевізникам;
- 8) завершальний етап: перевізники приймають замовлення від відправника та отримують винагороду.

Внизу діаграми зображено захищений реєстр блокчейн, який складається з 8 блоків, що відповідають кожному етапу процесу. Це забезпечує:

- незмінність та прозорість всіх транзакцій;
- безпечне зберігання даних;
- відстеження всього процесу;
- автоматизацію взаємодії між учасниками.

Ця система забезпечує надійний, прозорий та автоматизований процес організації перевезень, де всі транзакції та взаємодії захищені технологією блокчейн.

Для програмної реалізації, було написано псевдокод (рисунок 2), який описує процес реєстрації транспортних засобів в системі ITS та BFMS, а також механізм розподілу завдань і нагородження транспортних засобів.

```

1. Function RegisterInIVTP():
2.   // Obtain and encrypt vehicle identifier
3.   vehicleId ← get_vehicle_identifier()
4.   encryptedId ← encrypt(vehicleId)
5.   // Add vehicle node to ITS
6.   add_vehicle_node_to_ITS()
7.   return encryptedId
8. EndFunction
9. Function RegisterInBFMS():
10.  // Obtain identifiers
11.  vehicleId ← get_vehicle_identifier()
12.  encryptedId ← get_encrypted_vehicle_identifier()
13.  // Check identifier match in database
14.  if encryptedId == vehicleId_in_database then
15.    add_vehicle_to_BFMS()
16.    send_event("vehicle_added_to_blockchain")
17.  end_if
18. EndFunction
19. Function AssignTask():
20.  // Distribute task to all vehicles
21.  allVehicles ← distribute_task()
22.  // Process task acceptance
23.  vehicleX ← accept_task()
24.  // Check task completion and reward allocation
25.  if task == completed then
26.    allocate_reward(vehicleX)
27.  end_if
28. EndFunction

```

Рисунок 2 – Псевдокод процесу реєстрації

Алгоритм BFMS складається з трьох основних функцій:

1. RegisterInIVTP:
 - бере ID транспортного засобу;
 - шифрує його;
 - додає транспорт в систему ITS;
 - повертає зашифрований ID;
2. RegisterInBFMS:
 - перевіряє ID транспорту та його зашифровану версію;
 - якщо дані співпадають – додає транспорт в систему BFMS;
 - записує інформацію в блокчейн;
3. AssignTask:
 - розсилає завдання всім транспортним засобам;
 - один з транспортних засобів (X) приймає завдання;
 - після виконання завдання транспорт X отримує винагороду.

Це простий алгоритм реєстрації та управління транспортними засобами через блокчейн, який забезпечує безпеку даних через шифрування та автоматизує процес призначення завдань.

Запропонована система забезпечує:

- підвищення безпеки даних через криптографічне шифрування;
- прозорість всіх транзакцій завдяки розподіленому реєстру;
- автоматизацію процесів через смарт-контракти;
- зниження ризиків шахрайства та помилок.

Перспективи подальших досліджень

Планується розширення функціоналу системи через впровадження:

- механізмів машинного навчання для оптимізації маршрутів;
- додаткових рівнів верифікації учасників;
- інтеграції з IoT-пристроями для розширення можливостей моніторингу.

Список використаних джерел

1. Al-Ali M. S., Al-Mohammed H. A., Alkaeed M. Reputation Based Traffic Event Validation and Vehicle Authentication using Blockchain Technology. In Proceedings of the 2020 IEEE International Conference on Informatics, IoT, and Enabling Technologies (ICIoT), Doha, Qatar, 2–5 February 2020; pp. 451–456. [Google Scholar]
2. Ленков С. В. Алгоритм прогнозування для показників надійності і вартості експлуатації об'єктів радіоелектронних засобів озброєння / С. В. Ленков, Г. В. Банзак, В. М. Цицарєв, Я. М. Проценко // Системи обробки інформації. – 2016. – № 9 (146). – С. 28–30.
3. Haque M. A., Shetty S., Krishnappa B. ICS-CRAT: A Cyber Resilience Assessment Tool for Industrial Control Systems. IEEE 5th Intl Conference on Big Data Security on Cloud (BigDataSecurity), IEEE Intl Conference on High Performance and Smart Computing, (HPSC) and IEEE Intl Conference on Intelligent Data and Security (IDS). 2019. pp. 273–281. DOI: 10.1109/BigDataSecurity-HPSC-IDS.2019.00058. <https://doi.org/10.1109/BigDataSecurity-HPSC-IDS.2019.00058>

Ворвуль Д.М., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Безсмертний В.Ю., Київський національний університет імені Тараса Шевченка
Мусієнко А.П., д.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

ОЦІНКА ПРОДУКТИВНОСТІ ГОТОВИХ РІШЕНЬ ПОШУКОВО- ДОПОВНЕНОЇ ГЕНЕРАЦІЇ ДЛЯ НАДАННЯ РЕКОМЕНДАЦІЇ З ВИБОРУ НАЙКРАЩОГО РІШЕННЯ

У цій статті представлено детальний аналіз ефективності існуючих готових рішень пошуково-доповненої генерації (Retrieval-Augmented Generation, RAG) у фінансовій сфері, з особливим акцентом на компаніях, що займаються злиттям та поглинанням. З огляду на критичну важливість швидкого та точного доступу до релевантної інформації в цьому секторі, дослідження спрямоване на розробку комплексного програмного забезпечення для оцінювання продуктивності RAG-систем за тринадцятьма ключовими показниками. Ці показники були ретельно обрані для відображення різних аспектів роботи систем, включаючи їхню здатність надавати точні, релевантні та обґрунтовані відповіді. Оцінка включає глибоке порівняння відповідей, наданих різними RAG-системами, з еталонними експертними відповідями для більш точного визначення їхніх сильних та слабких сторін, а також потенціалу для подальшого вдосконалення.

У рамках дослідження для аналізу було зібрано обширний датасет, що містить запитання та відповіді, надані провідними експертами галузі фінансів. Продуктивність систем оцінювалась на основі автоматичної перевірки відповідей з використанням передових великих мовних моделей, а також за допомогою прямих експертних оцінок. Для наочності та полегшення порівняння було створено лідерборд, який відображав результати систем за різними показниками, такими як релевантність, точність, стійкість до шуму, обґрунтованість, узгодженість, швидкість роботи та інші важливі критерії. Оцінювання проводилось на основі вдосконаленого фреймворку RGAR (Retrieval, Generation, Additional Requirement), що дозволило врахувати додаткові вимоги та особливості фінансової сфери при аналізі роботи систем.

Додатково, для забезпечення максимальної об'єктивності та точності оцінок, було застосовано комбінований підхід, що поєднує автоматизовані методи з експертними аналізами. Це дозволило не лише оцінити технічні характеристики систем, але й врахувати практичну цінність та зручність використання з точки зору кінцевих користувачів.

Основні результати дослідження вказують на те, що система Azure OpenAI GPT-4 у поєднанні з Azure AI Search показала найкращі результати серед протестованих рішень. Вона продемонструвала найвищий рівень точності у відповідях, високу стійкість до контрфактної інформації та здатність ефективно інтегрувати різноманітні інформаційні джерела. Ці характеристики роблять її особливо придатною для використання у фінансовому секторі, де важлива кожна деталь та висока надійність даних.

Практичне застосування отриманих результатів є надзвичайно важливим для компаній у фінансовому секторі. Завдяки цим висновкам, організації зможуть більш обґрунтовано підходити до вибору RAG-рішень, обираючи ті, що найбільше відповідають їхнім специфічним потребам та забезпечують швидку і точну обробку критично важливої інформації. Це сприятиме підвищенню ефективності процесів прийняття рішень та загальної конкурентоспроможності компаній на ринку.

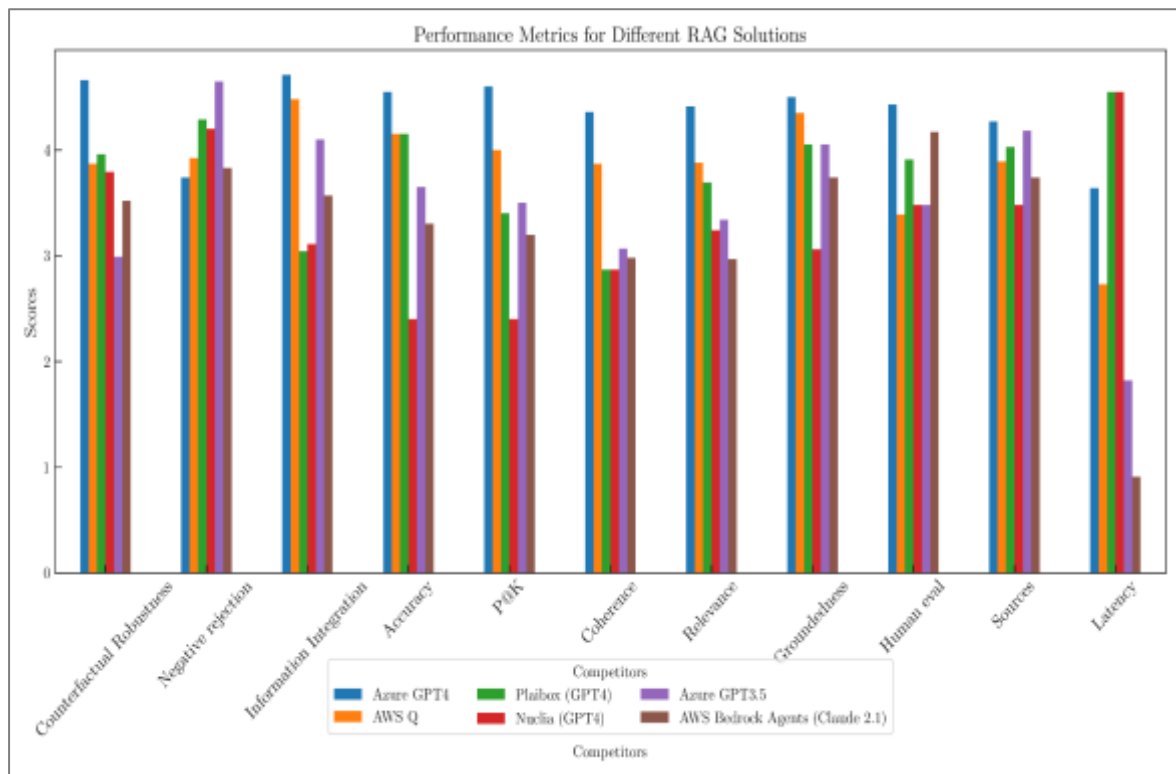


Рисунок 1 – Фінальний лідербординг

Висновки

Висновки дослідження чітко вказують на те, що RAG-системи мають значний потенціал у фінансовій сфері, особливо в галузі злиття та поглинання, де своєчасний доступ до точної та релевантної інформації може мати вирішальне значення. Використання таких систем дозволяє значно підвищити ефективність процесів прийняття рішень, оскільки вони забезпечують швидкий доступ до необхідних даних та здатні надавати обґрунтовані відповіді на складні запити. Це особливо важливо в контексті сучасного динамічного фінансового ринку, де інформація швидко змінюється, а конкуренція стає все більш жорсткою.

Таким чином, впровадження RAG-систем може стати ключовим фактором успіху для компаній, що прагнуть залишатися конкурентоспроможними та приймати обґрунтовані рішення на основі актуальних даних. Результати дослідження також вказують на необхідність подальших досліджень та розвитку в цій сфері, щоб максимально використовувати можливості, які надають сучасні технології штучного інтелекту у фінансовому секторі.

Список використаних джерел

1. Kwiatkowski T., et al. (2019). Natural Questions: A Benchmark for Question Answering Research. Transactions of the Association for Computational Linguistics, № 7, pp. 453–466.
2. Lewis P., et al. (2020). Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks. arXiv preprint arXiv : 2005. 11401.
3. Shao Z., Gong Y., Shen Y., Huang M., Duan N., Chen W. (2023). Enhancing Retrieval-Augmented Large Language Models with Iterative Retrieval-Generation Synergy (Version 2). arXiv. <https://doi.org/10.48550/ARXIV.2305.15294>.

Секція 2

ПРОГРАМНЕ ЗАБЕЗПЕЧЕННЯ КІБЕРФІЗИЧНИХ СИСТЕМ

Добровольський Р.В., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Shiwei Zhu Information Research Institute,
Qilu University of Technology (Shandong Academy of Sciences)
Сенченко В.Р., с.н.с., Інститут проблем реєстрації інформації НАН України
Коваль О.В., д.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

ПРОГРАМНЕ ЗАБЕЗПЕЧЕННЯ ПЕРЕВІРКИ СЕМАНТИЧНОЇ СУМІСНОСТІ ОНТОЛОГІЧНИХ МОДЕЛЕЙ ДЛЯ ПРЕДМЕТНОЇ ОБЛАСТІ СИСТЕМИ ОЦІНКИ РІВНЯ НАУКОВОЇ РОБОТИ ВИКЛАДАЧІВ КАФЕДРИ

Забезпечення семантичної сумісності є складним завданням у системах, які працюють з великим обсягом різномірних даних та потребують узгодження між процесами, що представлені у вигляді сценаріїв. Основні проблеми, які виникають при цьому, включають:

- різномірність даних та моделей: дані можуть надходити з різних джерел і мати різні структури, що ускладнює їх інтеграцію;
- семантичні невідповідності: різні концепції або терміни можуть бути використані для позначення однакових понять, що призводить до суперечностей у даних;
- логічна узгодженість переходів: у складних сценаріях важливо забезпечити логічну послідовність кроків, щоб уникнути ситуацій, коли переходи між етапами здійснюються без дотримання необхідних умов [1].

Семантичний медіатор [2] є критичним елементом системи, який забезпечує сумісність та логічну цілісність процесів аналітичної діяльності. Його головна функція полягає у перевірці відповідності між елементами онтологічного сценарію, що дозволяє забезпечити узгодженість процесів у межах заданої предметної області. Це досягається завдяки аналізу семантичної сумісності, що допомагає синхронізувати сценарії з базою знань предметної області та зменшує ризик семантичних невідповідностей.

Медіатор також відіграє ключову роль у забезпеченні логічної узгодженості між кроками сценарію. Під час виконання сценарію аналітичної діяльності він аналізує кожен перехід між етапами, забезпечуючи їх відповідність логічним правилам і умовам сценарію. Це включає перевірку наявності всіх необхідних умов для переходу на наступний етап, що дозволяє уникнути логічних розривів або непослідовностей у сценарії. Таким чином, медіатор сприяє підтриманню цілісності процесу та підвищенню точності виконання аналітичних завдань [3].

Важливим компонентом системи є семантичний мапер, що є онтологічною моделлю [4], побудованою на основі структури бази знань і сценарію процесу. Структура цього елементу було створено з допомогою сучасного інструменту проєктування онтологій Protégé [5]. Мапер виконує функцію узгодження та зіставлення елементів сценарію з концепціями предметної області. Основні завдання, які виконує мапер:

- формування відповідностей між елементами сценарію та базою знань, забезпечуючи відображення реальних процесів і об'єктів предметної області;
- перевірка семантичної сумісності між елементами сценарію і домену, запобігаючи розбіжностям у моделях;
- мапер використовує механізми міркування (reasoner [6]) для перевірки відповідності між сценарієм і базою знань, що дозволяє уникнути семантичних конфліктів і забезпечити логічну узгодженість.

На діаграмі (рисунку 1) зображені основні класи та зв'язки, які використовуються мапером для підтримки цілісності моделі.

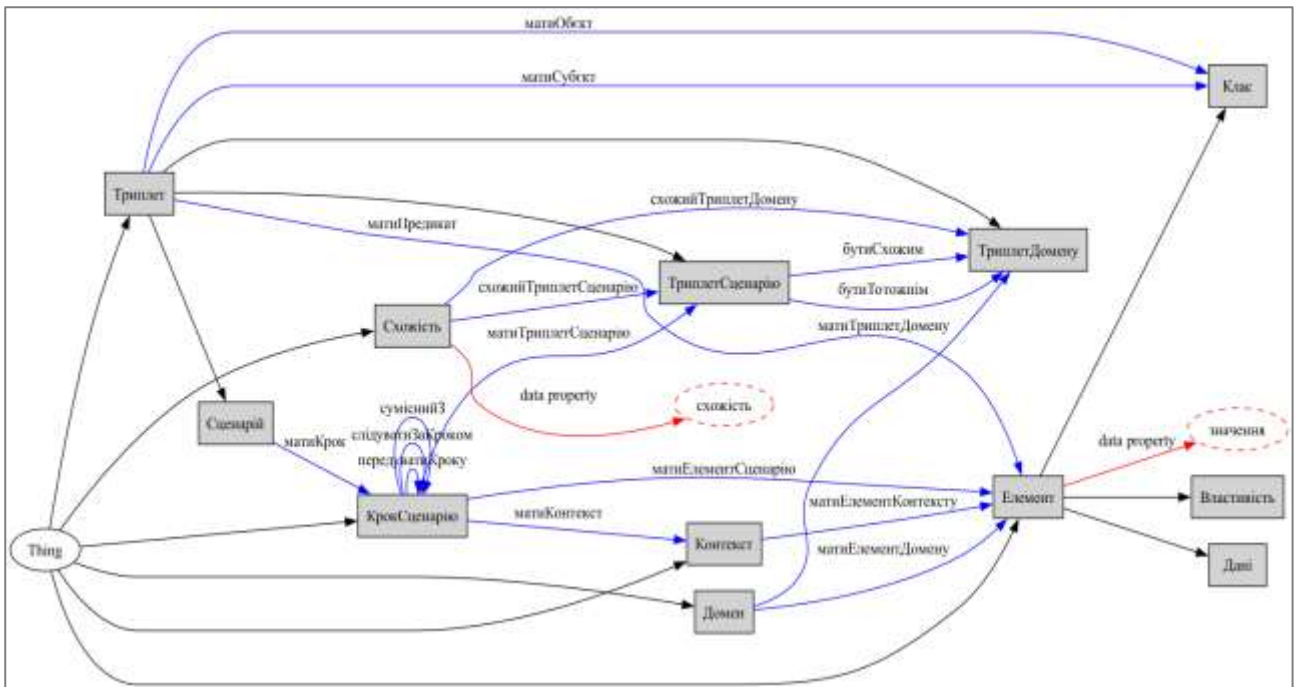


Рисунок 1 – Структура онтології семантичного медіатора

Отже, у представлений розробці семантичний медіатор є центральним інструментом для забезпечення семантичної сумісності і логічної узгодженості елементів сценарію у процесі аналітичної діяльності. Семантичний мапер, у свою чергу, сприяє узгодженню між сценарієм та базою знань, забезпечуючи відображення реальних концептів і процесів у системі. Це дозволяє зменшити ризик семантичних невідповідностей, підвищити точність і достовірність аналітичних процесів та забезпечити надійну підтримку користувачів у їхніх аналітичних завданнях. Система семантичного медіатора забезпечує узгодженість та логічну цілісність процесів, що робить її незамінною у контексті складних аналітичних сценаріїв.

Список використаних джерел

1. Сенченко В. Р. Семантична сумісність процесів складних сценаріїв аналітики. Інститут проблем реєстрації інформації НАН України, 2021. URL: <http://drsp.ipri.kiev.ua/article/view/244965> (дата звернення: 02.11.2024)
2. Moujane A., Chiadmi D., Benhlima L. (2006). Semantic Mediation: An Ontology-Based Architecture Using Metadata. Presented at the [Conference Name if available], April 2006. Available at ResearchGate. URL: <https://www.researchgate.net/publication/228643806>
3. Коваль О. В. «Методи та засоби комп'ютерного моделювання сценаріїв аналітичної діяльності: автореф. дис. ... докт. техн. наук: 01.05.02.» Київ, 2021. 134 с.
4. W3C OWL Working Group. URL: <https://www.w3.org/2001/sw/wiki/OWL>
5. Protégé (software). URL: <https://en.wikipedia.org/wiki/software>
6. Ontology Reasoning with Large Data Repositories. URL: IOS Press, <https://www.iospress.nl/book/ontology-reasoning-with-large-data-repositories/>

Єрмосін О.О., Національний технічний університет України
 «Київський політехнічний інститут імені Ігоря Сікорського»
 Gong Xiaodong, Institute of Oceanographic Instrumentation
 Qilu University of Technology (Shandong Academy of Sciences)
 Коваль О.В., д.т.н., Національний технічний університет України
 «Київський політехнічний інститут імені Ігоря Сікорського»

СИСТЕМА ПОБУДОВИ ОНТОЛОГІЧНОЇ МОДЕЛІ НА ОСНОВІ СЦЕНАРІЇВ АНД ДЛЯ ВИКОРИСТАННЯ В ПРЕДМЕТНІЙ ОБЛАСТІ «ОЦІНКА РІВНЯ НАУКОВОЇ РОБОТИ»

Інноваційна діяльність університетів та інститутів є визначною характеристикою, що прямо впливає на їх репутацію та конкурентоспроможність, а тому стає надзвичайно важливо об'єктивно оцінювати наукові здобутки підрозділів. У сучасному світі кожен навчальний заклад або підрозділ має бути зацікавлений у вимірюванні наукової активності своїх викладачів, щоб забезпечити якість освіти, підтримувати темпи професійного зростання та сприяти впровадженню інновацій. Так як традиційні методи оцінки, що засновані на експертних оглядах, потребують додаткових витрат часу та ресурсів, і водночас можуть бути дещо суб'єктивними, було розроблено комплексну систему, що дозволить автоматизувати даний процес.

Так як модулі, що описуються в даній роботі, є складовими системи «Оцінка рівня наукової роботи працівників кафедри», буде доцільно коротко описати й інші компоненти, щоб сформувавши концепцію загальної системи та призначення розроблених модулів.

Отже, призначенням загальної системи є надання користувачам інструментів для оцінки наукової діяльності викладачів за допомогою використання BPMN сценаріїв. Для реалізації цієї мети було спроектовано п'ять компонентів: модуль створення, валідації та збереження BPMN-сценаріїв [1] в форматі OWL [2], модуль аналізу семантичної сумісності конвертованих сценаріїв з онтологією предметної області та виконання сценаріїв, модуль поглибленого аналізу семантичної сумісності для покращення точності виконання сценаріїв, модуль створення нових онтологічних моделей на основі знань з веб-ресурсів, модуль злиття нових онтологій з онтологією предметної області.

Модуль створення та перетворення онтологічних моделей є першим модулем, з яким починає взаємодіяти користувач системи. Саме цей компонент надає можливість користувачам проектувати сценарії, що будуть використовуватися іншими модулями системи (рисунк 1). Окрім створення сценаріїв, даний модуль надає користувачам можливість завантажувати сценарії, які вони створили в інших додатках або власноруч.

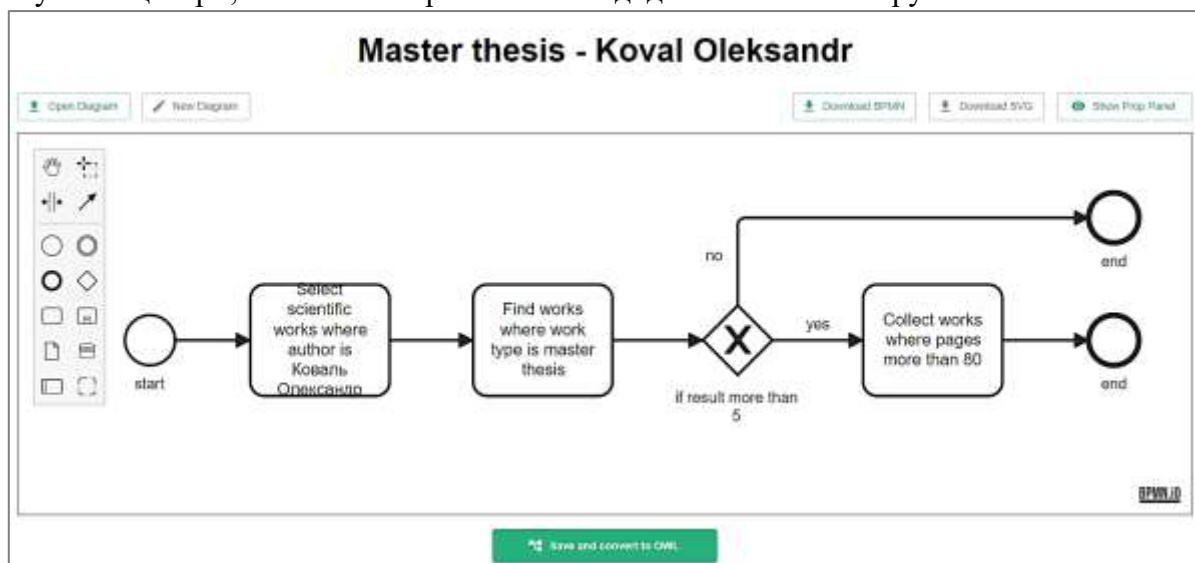


Рисунок 1 – BPMN сценарій для оцінки наукової діяльності

Перш ніж перейти до виконання сценарію, його потрібно перевірити, щоб переконатись, що даний сценарій відповідає нотації BPMN та додатковим вимогам системи. Цей процес є другим завданням даного модуля, він повинен передувати конвертації сценарію, щоб не витрачати ресурси системи на некоректно спроектовані сценарії. Останнім і найважливішим з основних завдань цього модулю є конвертація BPMN сценарію, представленого у вигляді XML файлу, в OWL формат. Це досягається завдяки зіставленню елементів BPMN відповідним концептам онтології та встановленню зв'язків між ними за допомогою властивостей об'єктів. Процес завершується збереженням сценарію в обох форматах в базі даних для його подальшого використання іншими модулями системи.

Модуль формування онтологій на основі даних з зовнішніх джерел надає користувачам можливість в автоматизованому режимі знаходити нові знання про наукові роботи для підтримки актуальності основної бази знань системи за рахунок розширення її новими знаннями. Від користувача система потребує лише URL ресурсу, на якому потенційно можуть бути знайдені відомості про наукові роботи та їх авторів. Після цього система автоматично аналізує зміст веб-сторінки з використанням попередньо навченої моделі штучного інтелекту для знаходження релевантних даних. Даний процес не обмежується знаходженням нових знань. Для кожного успішно проаналізованого ресурсу система надає короткий звіт про порівняння новоствореної онтології з онтологією предметної області. Також слід зазначити, що використання моделі штучного інтелекту для аналізу великих об'ємів даних вимагає певного часу, тому всі описані вище процеси відбуваються в асинхронному режимі, щоб задіяти всі виділені на цей компонент ресурси для зменшення часу виконання.

Отже, розроблені модулі системи забезпечують ефективне формування онтологій на базі BPMN сценаріїв та нових знань зі сторонніх ресурсів. Створені модулі мають важливе значення для функціонування комплексної системи «Оцінка рівня наукової роботи працівників кафедри», адже не тільки забезпечують створення BPMN сценаріїв, з яких починається робота користувача, а й допомагають підтримувати актуальність даних, що зберігаються в базі знань, тим самим розширюючи можливості системи.

Список використаних джерел

1. Business Process Model and Notation (BPMN), Version 2.0. URL: <https://www.omg.org/spec/BPMN/2.0.2/PDF> (дата звернення: 1.05.2024).
2. An introduction to the owl web ontology language. URL: <https://www.cse.lehigh.edu/~heflin/IntroToOWL.pdf> (дата звернення: 1.05.2024).

Єрмосіна О.О., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Gong Xiaodong Institute of Oceanographic Instrumentation
Qilu University of Technology (Shandong Academy of Sciences)
Коваль О.В., д.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

АНАЛІЗ ТА ВИКОНАННЯ СЦЕНАРІЇВ АНД З ВИКОРИСТАННЯМ ОНТОЛОГІЧНОЇ МОДЕЛІ ДЛЯ ЗАДАЧІ «ОЦІНКА РІВНЯ НАУКОВОЇ РОБОТИ»

Задача оцінки наукової роботи викладачів кафедри охоплює аналіз наукових видань, якість викладання, наукове керівництво, участь у дослідницьких проєктах та конференціях, а також інші наукові процеси. Завдяки використанню онтологічних моделей та BPMN сценаріїв можна автоматизувати оцінювання, забезпечуючи об'єктивність та ефективність даного процесу. Результати аналізу можна представити у вигляді звіту з наукової діяльності викладача за тими критеріями, що задає користувач у своєму сценарії. Таким чином, можна гнучко налаштовувати критерії оцінки відповідно до потреб користувача.

Комплексна система «Оцінка рівня наукової роботи працівників кафедри» складається з кількох окремих модулів:

- 1) створення, перевірка, зберігання BPMN-сценаріїв та їх перетворення в онтологічну модель (формат OWL [1]);
- 2) формування нових онтологічних моделей з наукометричних веб-ресурсів для розширення онтології (бази знань) предметної області;
- 3) первинний аналіз семантичної сумісності [2] з предметною областю та виконання сценарію на основі онтологічної моделі;
- 4) поглиблений аналіз семантичної сумісності онтологічних моделей сценаріїв з предметною областю («Семантичний медіатор»);
- 5) розширення та злиття онтологічних моделей для доповнення бази знань.

У даній роботі розглядається модуль аналізу та виконання сценаріїв аналітичної діяльності. Дана система була створена для первинної перевірки семантичної сумісності OWL сценаріїв з предметною областю, конвертації елементів сценарію у формат, придатний для виконання SPARQL [3] запитів, та виконання сценаріїв на основі онтології предметної області.

Аналіз семантичної сумісності з предметною областю є важливим етапом виконання OWL сценарію. Дана перевірка необхідна для пошуку розбіжностей між сценарієм аналітичної діяльності та предметною областю. Семантичний аналіз здійснюється перед виконанням сценарію, щоб уникнути помилок через невідповідність концептів. Хоча глибокий семантичний аналіз виконує модуль «Медіатор», для забезпечення певної автономності модулю «Виконання сценарію» реалізовано первинну перевірку сумісності, яка визначає, чи всі елементи сценарію відповідають предметній області. Несумісні сценарії не виконуються, а користувач отримує інформацію про розбіжності між сценарієм та предметною областю.

Виконання сценаріїв включає в себе визначення послідовності процесу, розпізнавання контексту завдання з назви процесу, формування черги виконання завдань, та побудову і виконання SPARQL запиту в базу знань предметної області на основі встановленої послідовності завдань.

В онтологічній моделі сценарію в кожного елементу процесу можна знайти інформацію про попередній та наступний елементи. На основі цих даних формується черга завдань. Послідовність завдань починається з події «Старт».

Прості задачі додаються до черги виконання. Задачі в підпроцесах обробляються послідовно, поки не буде знайдена кінцева подія підпроцесу.

Сценарій може розгалужуватись, якщо використовуються шлюзи. Система підтримує кілька типів шлюзів:

- ексклюзивні — обирається одна гілка сценарію на основі умови шлюзу;

- паралельні — дозволяється одночасне виконання кількох гілок;
- інклюзивні — виконуються декілька гілок залежно від умови шлюзу.

Після завершення формування черги задач сценарію система формує SPARQL запит для пошуку інформації в базі знань предметної області.

У результаті виконання сценарію користувач отримує звіт у вигляді таблиці, сформованої на основі результатів пошуку в онтології бази знань. Дані результати можна завантажити у вигляді Excel файлу (рисунок 1).

ScientificWork	Author	ScientificWorkType	PublicationDate
Аналіз Відеопотоку Ідентифікація Людей За Статтю Т Сігайов Андрій Олександрович		Bachelor Thesis	2020-09-22T11:34:36
Інформаційний Веб Сервіс З Розрахунку Оцінки Ризику Швайко Валерій Григорович		Bachelor Thesis	2022-10-12T11:57:12
Симуляція Потоків Даних На Сервері З Метою Пеї Федорова Наталія Володимирівна		Bachelor Thesis	2022-10-12T11:52:46
Створення Програмних Засобів Формування Баз Даних Швайко Валерій Григорович		Bachelor Thesis	2020-09-28T13:50:48
Моніторинг Рівнів Гідропостів З Використанням Техн Швайко Валерій Григорович		Bachelor Thesis	2019-09-12T15:45:16
Створення Геоінформаційної Баз Даних Генеральні Швайко Валерій Григорович		Bachelor Thesis	2019-09-12T12:49:12
Web Додаток Аналізу Та Візуалізації Результатів Еко Швайко Валерій Григорович		Bachelor Thesis	2022-10-11T15:54:29
Інформаційний Пошук За Використання Семантичної Коваль Олександр Васильович		Bachelor Thesis	2021-10-01T07:46:30
Підсистема Оплати Послуг Сігайов Андрій Олександрович		Bachelor Thesis	2020-09-28T13:43:49
Обробка Та Передавання Надвеликих Масивів Даних Федорова Наталія Володимирівна		Bachelor Thesis	2022-10-13T08:38:29
Геоінформаційна База Даних Моніторингу Екологічних Швайко Валерій Григорович		Bachelor Thesis	2021-09-29T13:41:51
Програмні Засоби Моніторингу Соціально Економічних Швайко Валерій Григорович		Bachelor Thesis	2021-09-29T09:51:04
W E B Додаток Створення Цифрових Карт Сільгоспу Швайко Валерій Григорович		Bachelor Thesis	2022-10-13T07:53:20
Моделювання Управлінських Процесів В Інженерних Швайко Валерій Григорович		Bachelor Thesis	2020-09-25T09:24:28
Створення Web Додатку Формування Баз Даних Під Швайко Валерій Григорович		Bachelor Thesis	2021-10-04T08:09:07
Моделювання Оцінки Соціально Політичних Ризиків Швайко Валерій Григорович		Bachelor Thesis	2019-09-11T16:22:40
W E B Додаток Для Ведення Геоінформаційної Баз Даних Швайко Валерій Григорович		Bachelor Thesis	2022-10-21T06:50:48
Прототип Соціальної Мережі Швайко Валерій Григорович		Bachelor Thesis	2022-10-26T08:38:32
Створення Бібліотеки Функцій Для Аналізу Гідроакусу Швайко Валерій Григорович		Bachelor Thesis	2019-09-12T12:31:16
Програмний Додаток Формування Сценарію Моделювання Швайко Валерій Григорович		Bachelor Thesis	2021-10-04T13:38:57
Проектування Та Розроблення W E B Додатків На Платформі Сігайов Андрій Олександрович		Bachelor Thesis	2020-09-25T09:32:06
Модуль Студент Системи Управління Дипломними Коваль Олександр Васильович		Bachelor Thesis	2019-08-13T11:11:12
Система Уніфікації Методів Доступу До Баз Даних Коваль Олександр Васильович		Bachelor Thesis	2019-07-17T12:30:36

Рисунок 1 – Результати виконання сценарію, збережені в Excel файл

Отже, розроблена система забезпечує ефективний та об'єктивний підхід до оцінки наукової діяльності викладачів кафедри. Система дозволяє проводити перевірку семантичної сумісності та виконувати сценарії аналізу наукової роботи викладачів. Подальшими етапами вдосконалення системи може бути інтеграція з іншими платформами або впровадження автоматизованого формування та оновлення онтологічних словників.

Список використаних джерел

1. Web ontology language (OWL) and semantic web. URL: <https://doi.org/10.1145/1295289.1295290> (дата звернення: 1.05.2024).
2. Сенченко В. Р. Семантична сумісність процесів складних сценаріїв аналітики. Реєстрація, зберігання і обробка даних. 2021. Т. 23, № 3. С. 48–61. URL: <https://doi.org/10.35681/1560-9189.2021.23.3.244965> (дата звернення: 1.05.2024).
3. What Is SPARQL? URL: <https://www.ontotext.com/knowledgehub/fundamentals/what-is-sparql/> (дата звернення: 01.05.2024).

Чжан Мінцзюнь Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Стативка Ю.І., к.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

КОНВЕРТАЦІЯ РЕЙТИНГУ НЕОДНОРІДНИХ ОБ'ЄКТІВ З РАНГОВОЇ ШКАЛИ У МЕТРИЧНУ

Впорядкування неоднорідних об'єктів передбачає, що серед них знайдуться такі, яким, за певним показником, не можуть присвоєні числові значення. Добре відомий приклад — непорівнянність таких вишів, як Лондонська школа економіки, з рештою університетів, за показником N&S Шанхайського рейтингу університетів, де значення None інтерпретується як значення 0, а вага показника розподіляється між вагами інших первинних показників [1].

Нехай є множина M неоднорідних об'єктів, і непорівнянні з іншими за індикатором X_n утворюють множину $M^u \subset M$, а решта $M^c = M \setminus M^u$, $M^c \subset M$ — порівнянних за всіма n індикаторами. Множину $n - 1$ індикаторів, за якими всі об'єкти порівнянні позначимо через $K = \{X_1, X_2, \dots, X_{n-1}\}$

Тоді на множині M за індикаторами K може бути обчислений композитний рейтинговий показник P_K , який утворює порядок за відношенням $\leq_K$ на M , а на множині M^c , за всіма n індикаторами, — агрегований показник $\hat{P}_{agr}$ з порядком за відношенням $\leq_{agr}$. Тоді на множині M можемо визначити частковий порядок з відношенням $\leq$, яке отождествляється: між об'єктами $O_1 \in M^u$ та об'єктами $O_2 \in M$ — з $\leq_K$, між об'єктами $O_1, O_2 \in M^c$ — з $\leq_{agr}$. А це означає, що на так побудованій частково впорядкованій множині M може бути побудований і лінійний порядок, див. [2], тобто може бути обчислений агрегований рейтинговий показник у ранговій шкалі P^{Rank} .

Для відображення P^{Rank} на метричну шкалу ($P^{Rank} \Rightarrow P^{Metric}$) скористаємось аналогією з пружною системою. Вважатимемо, що є $n - 1$ послідовно з'єднаних пружин, жорсткість кожної з яких k^{O_i} пропорційна різниці Δ^{O_i} між значеннями P_K чи $\hat{P}_{agr}$, в залежності від рангу об'єкта, та належності до множини M^u чи M^c , тобто $\Delta^{O_i} = |\hat{P}_{agr}^{O_i} - \hat{P}_{agr}^{O_{i+1}}|$ при $O_i, O_{i+1} \in M^c$ та, інакше, $\Delta^{O_i} = |\hat{P}_K^{O_i} - \hat{P}_K^{O_{i+1}}|$.

Тоді коефіцієнт жорсткості еквівалентної системі пружини $k = \left(\sum_O \frac{1}{k^{O_i}}\right)^{-1}$, та загальна її деформація $L = B - A$ спричинить силу пружності $F = kL$. Отже, кожна з пружин буде деформована на $x^{O_i} = F/k^{O_i}$, а об'єкти будуть розміщені на відрізку $[A; B]$ на відстанях Δ^{O_i} , пропорційних коефіцієнтам жорсткості $k^{O_{ij}}$, тобто мірі віддаленості об'єктів за P_K та $\hat{P}_{agr}$.

Список використаних джерел

1. ShanghaiRanking's Academic Ranking of World Universities. Methodology 2024. Вилучено з <https://www.shanghairanking.com/methodology/arwu/2024>
2. Нікольський Ю. В., Пасічник В. В., Щербина Ю. М. (2007) Дискретна математика. Київ: Видавнича група BVH.

Voitko A.S., National Technical University of Ukraine
«Igor Sikorsky Kyiv Polytechnic Institute»
Kuzminyh V.O., c.t.sc., National Technical University of Ukraine
«Igor Sikorsky Kyiv Polytechnic Institute»

DYNAMIC MESSAGE ROUTING STRATEGY IN MICROSERVICE ARCHITECTURES FOR DATA PROCESSING

The demand for efficient data processing architectures in modern distributed systems that ensure flexibility, scalability, and reliability is growing. Microservice architectures, in particular, have become a standard for processing large volumes of data, thanks to their ability to distribute processing among independent services that can be easily adapted to various tasks and environmental changes. A key aspect of these architectures is inter-service communication, often achieved through message brokers. This approach enables asynchronous message exchanges, reducing dependencies among services and increasing the overall system performance.

The microservices design principle with message brokers relies on asynchronous interaction through an intermediary that decouples services, simplifies scaling, and enhances system reliability. In this architecture, each microservice performs one main function and communicates with other microservices via messages relayed by a broker (e.g., RabbitMQ, Apache Kafka, or ActiveMQ).

As an example of a microservice architecture, a data processing pipeline will be considered, whose main tasks include data extraction, transformation, and loading from one or more sources into a target repository, often for subsequent analysis or processing. Communication through a message broker is an effective approach for building ETL data processing pipelines. This architecture allows the processing to be broken down into separate stages performed by independent microservices that communicate asynchronously. Each stage (or microservice) is responsible for a specific pipeline function, simplifying development, scalability, and fault recovery.

An ETL pipeline typically consists of a fixed sequence of steps that each piece of data must undergo before loading into a database. The key steps are Extract, where data is loaded from external sources; Transform, where data undergoes preliminary processing; and Load, where data is directly loaded into the database. However, when scaling the system, each step can be divided into separate microservices with specific functionalities to ease the architecture's structure (reducing coupling) or optimize resource usage (offloading GPU-demanding steps onto dedicated servers).

In this scenario, message routing becomes a critical issue. A fixed set of sequential steps leads to a high load on the message broker and individual services, as every service must process each message. One way to address this issue is using the content-based routing (CBR) approach, which directs messages to different recipients based on content analysis. In CBR, the message broker examines specific attributes or data within each message and uses established rules to determine which service or group of services should receive it.

However, applying the CBR approach to an ETL pipeline has certain limitations:

- limited or unavailable content analysis functionality in the broker;
- increased load on the message broker due to additional logic;
- complexity in configuring and maintaining routing rules.

The following adaptations of CBR principles for ETL pipeline design are proposed to address these challenges:

- shift the message analysis logic from the broker to the services;
 - routing can be centralized in one service or distributed among all services;
- establish a set of requirements and dependencies for each service;
 - requirements represent a set of characteristics that a message must meet for routing to a particular service;
 - dependencies represent the requirement of subsequent processing steps on the results of prior steps;

- create a unique queue stack for each message, which outlines its processing route;
 - each service can read and modify the stack;
 - the first element in the stack points to the next service for message delivery;
 - the last element in the stack always designates the final ETL pipeline step;
 - a step counter aids in debugging infinite loops.

The proposed system reduces the load on the broker by minimizing the computational burden due to the absence of additional message analysis logic and the reduction in message traffic through optimized routing. Transferring the strategy choice to the services facilitates development by allowing full programming language usage in writing rules, and binding rules to services simplifies organization by incorporating rules into the overall version control system. Network topology is generated automatically based on specified requirements and dependencies (e.g., dependency – generating embeddings requires translating text from any language into English, rule – the text field for translation must not be empty, and the language must be supported).

Thus, the proposed methodology reduces computational resource usage by optimizing data processing paths and facilitates software development by automating the generation of microservice communication topology and integrating rules with existing development tools.

Список використаних джерел

1. Oliveira B., Oliveira O., Belo O. (2023). Overcoming traditional ETL systems architectural problems using a service-oriented approach. *Expert Systems*, 40(10), e13442. <https://doi.org/10.1111/exsy.13442>
2. Carzaniga A., Rutherford M. J., Wolf A. L. (2004). A routing scheme for content-based networking. In *IEEE INFOCOM 2004* (T. 2, pp. 918–928). Hong Kong, China. <https://doi.org/10.1109/INFCOM.2004.1356979>

Tararenko R.A., National Technical University of Ukraine
«Igor Sikorsky Kyiv Polytechnic Institute»
Tararenko E.R.,
Taras Shevchenko National University of Kyiv

KNOWLEDGE MODEL OF DATA IN THE REPRESENTATION OF INFORMATION OF INTELLIGENT ANALYTICAL SYSTEMS

Achieving the goal, quality, efficiency and competitiveness through decision-making is complicated by the requirements of the need for comprehensive analytics [1], which has already become the norm for the analysis of complex systems. But the requirements of comprehensive analysis present qualitatively new criteria that cannot be solved by simple approaches. The multitude of problems that need to be solved in order to enter the «solution space» of complex systems significantly increases the «threshold of the costs of successful solutions» but at the same time opens up «niches of new opportunities» and opportunities to further reduce costs and implementation time. First of all, this is related to the complex system of relationships and the dynamics of the behavior of complex systems. The practical implementation of tasks related to complex systems gives rise to a number of theoretical, methodological, technical, applied and other problems, both individually and in their combined application, which impose significant limitations on the implementation of tasks and achievement of the goal. We are observing a crisis in each of the components of the decision-making complex, which requires the search for qualitatively new solutions. For example: the famous economist V.M. Polterovych [2] presented 8 main factors of the economic crisis: improvement of mathematical tools, in-depth research and generalization of basic models, coverage of the theory of new spheres of economic life, accumulation of empirical data, changing the «stringency standard», collective nature of summarizing work, the principle of coexistence, «behavioral» revolution in theoretical macroeconomics, organizational growth.

The development of new technical capabilities was described by Viktor Meier-Schönberger and Kenneth Cukier as «a revolution that will change the way we live, work and think...» [3] and in the language of presenting complex systems in big data $N = \text{«ALL»}$. In tasks with large data, the process of data preparation takes 45-80 % of the time, which states the fact of low efficiency and a low level of automation. The need to increase the level of automation and use of information analytical systems – decision support systems and directly expert systems is indisputable. For complex systems, classical and simplified approaches to automation are unacceptable, and approaches based on knowledge that characterize the uniqueness and scale of systems are used. Therefore, databases built on knowledge are a key element of building analytical chains of conclusions. For the tasks of complex systems, modern databases are still not fast enough [4], so the quality of information analytical systems is influenced by the database model and its technical implementation.

We will distinguish three main categories of building a knowledge model of databases – methodology, architecture and organization.

We consider the methodology based on the need to use adequate universal models [5] – we consider the cognitive data model as a transforming projection system of the real-world system. Practically – the formation of digital doubles.

The problem that information technologies and systems «don't know how to process information – they don't understand it» is solved by comparing analytical systems for generating conclusions with defined control knowledge. At the same time, we use the approach proposed and implemented in [6, 7]. We evaluate the information in the database model by quality, expanding the representation of indicators as a non-negative component of their states. Each indicator of the database is considered from six main categories of its assessment X (X_0 , X_I , X_{II} , X_T , X_N , X_m), where X_0 is the true value; X_I – statistical significance; X_{II} – estimated value; X_T – technological value; X_N – normative value; X_m is the consumer value. In the process of building logical chains of

conclusions, we introduce a new concept of «ESSENCE», which will be taken into account in the data model.

The model of the architecture of the knowledge database is built on representations of 4D space, within which the search for solutions is carried out: 1st dimension – Data; 2nd dimension – Volume; 3rd dimension – Velocity (speed); The 4th dimension – Knowledge (knowledge) – ordered representations obtained on the basis of a holistic approach.

The organization has a multitude of possible solutions and depends on the «software and technical» means and subject area.

Conclusions. It must be recognized that the development of artificial intelligence tools is considered as a search for unified solutions and approaches that erase the boundary between individual specialized ones, and definitely require the use of software and technical complexes of a qualitatively new level of complexity. Accordingly, digital doubles and other new technologies indicate the need for qualitatively new knowledge databases that overcome the critical limitations of existing ones. The development of technical capabilities indicates the tendency of the possible fusion of knowledge databases with software complexes, despite the fact that this is not yet obvious, but the need to increase the efficiency of complex software and technical complexes, especially in robotics, indicates this.

References

1. Chernikov A. (2008). Business analytics should become «all-encompassing». Computer Review, (40), October 21.
2. Polterovich V. M. (1997, March 18). The Crisis of Economic Theory. Report at the Scientific Seminar of the Department of Economics and CEMI RAS «Unknown Economy» Retrieved from <https://ts1.cemi.rssi.ru/rus/publicat/e-pubs/d9702t/d9702t.htm>
3. Mayer-Schönberger V., Cukier K. (2013). Big Data: A Revolution That Will Transform How We Live, Work, and Think. Eamon Dolan/Houghton Mifflin Harcourt.
4. Alizar A. (2024, October 21). Why are DBMS so slow. Habr. <https://habr.com/ru/companies/ruvds/articles/851330/>
5. Klir G. J. (1985). Architecture of systems problem solving. Plenum Press.
6. Taranenko R. A. (2005). Information quality: A review of representations in the context of integrity paradigm. USiM, T. 199 (5), pp. 3–12.
2. Katerinchik R. (2020, October 1). Will business analysts replace programmers? Corezoid cases from the Artjoker team. VC.ru. <https://vc.ru/services/163238-zamenyat-li-biznes-analitiki-programmistov-keisy-po-corezoid-ot-komandy-artjoker1>
3. Runkler T. A. (2020). Data analytics: Models and algorithms for intelligent data analysis. Springer Fachmedien Wiesbaden GmbH, part of Springer Nature.
4. Rivera G., Cruz-Reyes L., Dorronsoro B., Rosete A. (Eds.). (2023). Data analytics and computational intelligence: Novel models, algorithms and applications. Springer Nature.

Савко В.Я., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Гаврилко Є.В., д.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

РОЗРОБКА ІНТЕЛЕКТУАЛЬНОГО СПЕЦІАЛЬНОГО ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ УПРАВЛІННЯ ВЕНТИЛЯЦІЙНИХ СИСТЕМ КОНФАЙНМЕНТУ НА ОСНОВІ УРАХУВАННЯ ІНЕРЦІЙНОСТІ СИСТЕМИ

Забезпечення ядерної, радіаційної та, комплексно, екологічно безпечної експлуатації НБК над Об'єктом «Укриття» (ОУ) на Чорнобильській АЕС є складною і багатофакторною задачею. На рисунку 1 наведений вигляд НБК. НБК – надвелика спеціалізована споруда з великою кількістю спеціалізованого обладнання та потужного устаткування. Крім всього ключовою функцією НБК ОУ є забезпечення задачі ЯБ та РБ.



Рисунок 1 – Новий Безпечний Конфайнмент об'єкту «Укриття» Чорнобильської атомної електростанції. Вигляд з середини

Завданням, що ставиться в цій публікації є розробка інтелектуального СПМЗ на основі машинного навчання з підкріпленням на основі генетичних алгоритмів, що дозволить під час експлуатації ВС НБК врахувати інерційність повітряних мас під НБК та оптимізувати по критерію мінімізації витрати електроенергії та викиди РА.

Метою роботи є забезпечення оптимізації витрати електроенергії ВС при мінімумі неконтрольованих викидів РА за межі НБК.

Данні з зовнішніх та внутрішніх датчиків, збираються в базу даних НБК. Ці дані, отримані з широкого спектру зовнішніх та внутрішніх датчиків, є фундаментом для всіх наступних етапів аналізу та управління. Зібрані в базу даних НБК, вони проходять через процес обробки з використанням програмних засобів, що дозволяє точно визначати параметри роботи ВС.

Основні зібрані гідрометеорологічні показники включають:

$$M_i = F(M_1; M_2; M_3; M_4; M_5 \dots), \quad (1)$$

де M_1 швидкість вітру – визначає інтенсивність та потенційний вплив вітру на НБК, впливаючи на розрахунки тиску та вентиляційні потреби; M_2 – вологість повітря – зовнішня та внутрішня вологість визначають умови працездатності ВС, а також впливають на рівні конденсації та загальний комфорт у НБК; M_3 – опади (дощ/сніг) – наявність та тип опадів впливають на робочі параметри НБК, зокрема на управління зовнішніми впливами та

забезпечення стабільності систем; M_4 – напрям вітру – важливий для аналізу впливу вітрового навантаження на конструкції НБК та для оптимізації розміщення вентиляційних виходів.

Тренування ГАл (генеративного алгоритму) для прогнозування параметрів тиску в ОО та КП є ключовим елементом в управлінні ВС НБК. Використання даних, що містять n вимірів за часом і k параметрів експлуатаційних даних на кожен момент часу, дозволяє створити високодеталізовану картину робочих процесів.

Забезпечення інерційності ВС є критично важливим аспектом управління середовищем в НБК на ВУ ЧАЕС. Це дозволяє ефективно контролювати обсяги повітряного потоку, що переміщується між внутрішнім простором НБК і зовнішнім середовищем, з метою мінімізації витоків та оптимізації енергоспоживання. Систематичний підхід до налаштувань агрегатів забезпечує високу адаптивність системи до зовнішніх змін та експлуатаційних потреб.

Для перевірки обраної методики використано модельну ситуацію з приміщенням об'ємом 50 м^3 , де через систему вентиляції та нещільності здійснювався контроль тиску. Сумарна площа нещільностей становила 0.02 м^2 , що відповідає реалістичним умовам для середньостатистичного приміщення з мінімальними втратами через недоліки у герметизації. Задача полягала у забезпеченні стабільного позитивного тиску в $+10 \text{ Па}$ всередині приміщення відносно атмосферного тиску зовні, що мінусує на 10 Па .

На рисунку 2 представлено графік результату роботи системи прогнозування швидкості вітру, що використовується для моніторингу та адаптації вітрових умов. Графік показує прогнозовану швидкість вітру з кроком у одну годину, надаючи уявлення про зміни вітрових умов.

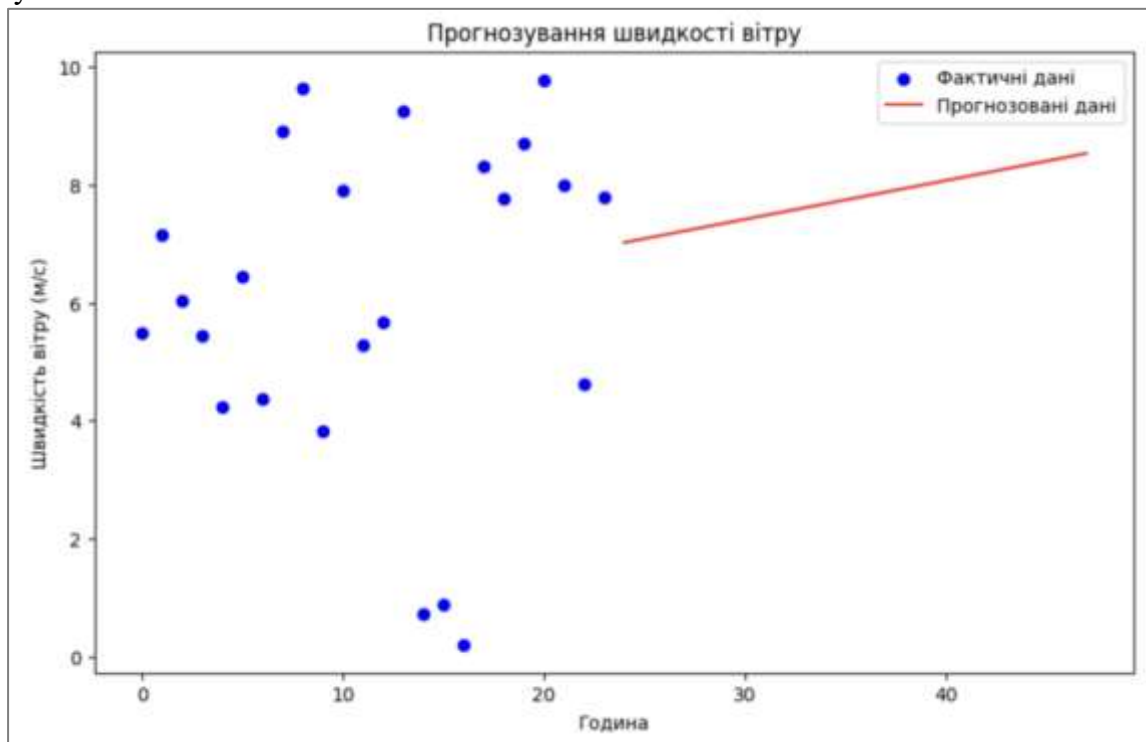


Рисунок 2 – Порівняння показів датчика КП-ОС та прогнозованого значення для тестової вибірки

На осі абсцис відображено час доби в форматі годин, що дозволяє легко ідентифікувати, у який час доби очікується зміна швидкості вітру. Вісь ординат показує швидкість вітру в метрах за секунду (м/с).

Ця модель дозволяє аналізувати, як швидко ВС зможе досягти потрібного рівня тиску в НБК і як цей тиск буде підтримуватися або змінюватися в часі залежно від змін умов M .

Описані підходи та результати розрахунків демонструють ефективність використання розробленої методики для зміни підходів до управління ВС НБК контролю та управління станами.

Отримані дані можуть бути використані для оптимізації роботи ВС з метою економії енергії та забезпечення необхідних умов комфорту та безпеки в приміщеннях.

Список використаних джерел

1. Pysmennyy Y., Havrylko Y., Krukovskyi P., Starovit I., Diadiushko Y. Development of special mathematical software for management of ventilation units of the New Safe Confinement of the Chernobyl Nuclear Power Plant. Nuclear and radiation safety. 2022, T. 2 (94), С. 35–43.
2. Loboda P., Starovit I., Shushura O., Havrylko Y., Ventilation control of the New Safe Confinement of the CHNPP based on neuro-fussy networks. Informatyka, Automatyka, Pomiary W Gospodarce I Ochronie Środowiska. 2023. T. 13 (4), pp. 114–118.
3. Krukovskyi P. G., Dyadyushko E. V., Sklyarenko D. I., Starovit I. S. Unorganized emissions of air with radioactive aerosols from the new safe confinement of the Chernobyl Nuclear Power Plant into the surrounding environment. Issues of atomic science and technology, № 6, 2017, pp. 181–186.

Максименко П.О., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Гусева І.І., к.е.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

ОБМЕЖЕННЯ І ВИКЛИКИ КОМП'ЮТЕРНОГО ЗОРУ В УМОВАХ НИЗЬКОГО ОСВІТЛЕННЯ ПРИ СТВОРЕННІ КАРТ ДЛЯ ПОТРЕБ ТРАНСПОРТНОЇ СФЕРИ

Комп'ютерний зір (КЗ) є важливим інструментом для автоматизації процесів розпізнавання та аналізу зображень у різних галузях. Особливо у транспортній сфері КЗ використовується для створення детальних карт та навігаційних систем, які покращують безпеку та ефективність перевезень. Проте його застосування в умовах низького освітлення залишається складним завданням через ряд обмежень і викликів. Одним з основних обмежень є зниження якості зображень через шум, низький контраст та втрату деталей, що ускладнює точне розпізнавання дорожньої інфраструктури та об'єктів [1]. Це може призвести до помилок у картографуванні. Це критично, адже в багатьох реальних сценаріях, таких як нічні перевезення або погані погодні умови, якість зображень суттєво погіршується [1], що впливає на точність і надійність систем КЗ у транспортній сфері.

Основним компонентом при розв'язанні проблеми КЗ в умовах низького освітлення є розробка методів покращення якості зображень та адаптації моделей до складних умов.

Основна ідея застосування цих методів полягає в підвищенні видимості та деталізації зображень. При цьому зазвичай використовуються три основні підходи, описані далі.

1. Методи підвищення яскравості та контрастності (white box) – ці методи базуються на традиційних алгоритмах обробки зображень, які детально змінюють параметри зображення для покращення його візуальних характеристик. Основні техніки включають:

- гістограмна рівномірність (Histogram Equalization): Цей метод розподіляє інтенсивності пікселів більш рівномірно по всьому діапазону яскравості, що дозволяє покращити контрастність зображення;

- адаптивна гістограмна рівномірність (Adaptive Histogram Equalization): Розширює стандартну гістограмну рівномірність, застосовуючи її до малих областей зображення для збереження локальних деталей;

- фільтри для зменшення шуму: Умови низького освітлення часто призводять до збільшення шуму в зображенні. Використовуються фільтри, такі як медіанний або Гауссів;

- методи масштабування яскравості: Регулювання кривих яскравості та контрастності для підвищення видимості темних областей.

2. Глибокі нейронні мережі для покращення зображень (black box) – цей підхід використовує моделі глибокого навчання, зокрема згорткові нейронні мережі (CNN). Основні методи включають:

- автокодери (Autoencoders): Навчаються відтворювати вхідне зображення, але з покращеною якістю. Наприклад, модель LLNet [1] використовує глибокий автокодер для покращення зображень в умовах низького освітлення;

- Generative Adversarial Networks (GANs): Використовують дві мережі (генератор і дискримінатор) для створення зображень, які максимально наближені до реальних;

- моделі на основі перенесення стилю: Використовують попередньо навчені моделі для перенесення характеристик високоякісних зображень на зображення з низьким освітленням.

3. Гібридні методи (grey box) – ці методи комбінують традиційні алгоритми обробки зображень з методами глибокого навчання, використовуючи переваги обох підходів. Основні приклади включають:

- комбіновані моделі: Використання попередньої обробки зображення традиційними методами перед подачею його на вхід нейронної мережі;

- моделі з врахуванням фізичних характеристик: Інтеграція фізичних моделей формування зображення в структуру нейронної мережі для покращення її продуктивності;
- адаптивні методи: Застосування алгоритмів, які динамічно обирають між традиційними та нейронними методами залежно від умов зображення.

У розробці рішень для комп'ютерного зору в умовах низького освітлення часто використовуються моделі, що поєднують класичні алгоритми та глибоке навчання. Наприклад, метод LIME [2] спочатку оцінює карту освітленості зображення, використовуючи традиційні методи, а потім застосовує алгоритми для підвищення яскравості та збереження деталей. Іншим прикладом є модель RetinexNet [2], яка базується на теорії Retinex і використовує глибоку нейронну мережу для декомпозиції зображення на компоненти відображення та освітлення, що дозволяє ефективно покращувати якість зображення. При розробці рішень для КЗ в умовах низького освітлення створено моделі, що містять як класичні алгоритми обробки зображень [2,3], так і сучасні моделі з використанням глибокого навчання, які дозволяють відновлювати деталі та покращувати якість зображень.

Однією з сфер застосування цих алгоритмів є інструментальні засоби розміщення візуальних об'єктів в шарах мапи, де вони можуть відігравати ключову роль, оскільки такі засоби часто можуть зіткнутися з проблемами під час використання в темний час доби, що погіршує їх ефективність. Покращені алгоритми дозволяють точніше ідентифікувати та геоприв'язувати об'єкти, такі як дорожні знаки, транспортні засоби або вибоїни навіть в умовах низького освітлення, що значно покращує ефективність такого засобу. Це сприяє розвитку навігаційних систем, підвищенню безпеки дорожнього руху та ефективності автономних транспортних засобів, які покладаються на точні та актуальні дані мапи.

Таким чином, подолання обмежень комп'ютерного зору в умовах низького освітлення є важливим для підвищення точності та надійності систем у реальних сценаріях. Використання традиційних методів обробки зображень, глибоких нейронних мереж та їхніх гібридів дозволяє ефективно покращувати якість зображень і адаптувати моделі до складних умов.

Список використаних джерел

1. Guo X., Li Y., Ling H. (2017). LIME: Low-Light Image Enhancement via Illumination Map Estimation. *IEEE Transactions on Image Processing*, Т. 26 (2), pp. 982–993. <https://doi.org/10.1109/TIP.2016.2639450>
2. Chen W., Wang W., Yang W., Liu J. (2018). Deep Retinex Decomposition for Low-Light Enhancement. *У Proceedings of the British Machine Vision Conference (BMVC)*. https://www.researchgate.net/publication/327033239_Deep_Retinex_Decomposition_for_Low-Light_Enhancement
3. Li C., Guo J., Han J., Cong R. (2018). LightenNet: A Convolutional Neural Network for Weakly Illuminated Image Enhancement. *Pattern Recognition Letters*, 104, pp. 15–22. <https://www.sciencedirect.com/science/article/abs/pii/S0167865518300163>

Ярмак Д.О., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Варава І.А., к.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

ЗАСОБИ ФОРМУВАННЯ СЦЕНАРІЇВ ДИНАМІЧНИХ ГІДРОАКУСТИЧНИХ ЕКСПЕРИМЕНТІВ

Динамічні гідроакустичні експерименти є важливими для вивчення підводного поширення звуку, яке має застосування в морській навігації, покращенні навігації під водою для підводних човнів та підводної картографії.

Експеримент є ключовим елементом наукового дослідження, що є методом перевірки гіпотез та збору даних. У гідроакустиці експерименти зосереджені на тому, як звукові хвилі поширюються крізь воду, взаємодіють із різними середовищами та піддаються впливу умов навколишнього середовища. Ці експерименти часто передбачають надсилення звукових хвиль від джерела та запис того, як вони відбиваються від об'єктів, морського дна або як вони поширюються на відстані.

Експерименти можна розділити на два типи: статичні та динамічні. Статичні експерименти – це ті, де умови навколишнього середовища залишаються відносно постійними з часом. На відміну від статичних експериментів, ці динамічні установки відстежують рухомий об'єкт (наприклад, судно або морську тварину), який генерує звук уздовж заздалегідь визначеної траєкторії. Цей рух дозволяє дослідникам досліджувати, як звук поводить себе за мінливих умов навколишнього середовища, таких як різна глибина, солоність і температура [1].

Формування динамічних гідроакустичних експериментів передбачає планування та конфігурацію руху об'єкта, що випромінює шум, по заданій траєкторії. Це рухоме джерело діє як динамічна змінна, дозволяючи дослідникам спостерігати поведінку звуку під час його взаємодії з підводними об'єктами, такими як траншеї, хребти. Основні вхідні дані включають батиметрію, умови навколишнього середовища (наприклад, солоність, температуру) і траєкторію об'єкта. Правильна інтеграція цих джерел даних має вирішальне значення для точного представлення реальних умов.

Географічні інформаційні системи (ГІС) пропонують важливі можливості для керування та візуалізації просторових і часових аспектів динамічних гідроакустичних експериментів. ГІС дозволяє дослідникам накладати кілька шарів даних, таких як батиметричні карти та змінні навколишнього середовища, на траєкторію об'єкта, що випромінює звук [2]. Без ГІС ручна інтеграція та керування цими даними були б трудомісткими та схильними до помилок [3].

Ефективний ГІС-інструмент для формування динамічних гідроакустичних експериментів має бути модульним, забезпечуючи гнучку настройку та керування. Нижче наведено детальну структуру інструменту, розділеного на основні модулі:

Модуль імпорту/експорту – цей модуль обробляє інтеграцію різних джерел даних, таких як батиметрія, умови навколишнього середовища чи вже готові експерименти для їх модифікації. Експорт створених в застосунку експериментів, що включає в себе сцену, на якій відбувається експеримент, батиметрія, гідроакустичні датчики, морські об'єкти та їх траєкторія руху по сцені, для їх подальшого використання, наприклад для моделювання поширення звуку.

Модуль візуалізації, в який входить створення комплексної експериментальної карти, що відображає дані, розбиті на відповідні логічні шари такі як: гідроакустичні датчики, сцена проведення експерименту, морські об'єкти та їх траєкторія. Візуалізація навколишнього середовища допомагає користувачам досліджувати поширення звуку в середовищах із змінною глибиною, пропонуючи розуміння просторових зв'язків між рухом об'єкта та підводними особливостями

Модуль побудови сценаріїв експериментів. В цьому модулі описуються функції для створення чи оновлення компонентів експерименту. При цьому створення експерименту відбувається за певним сценарієм (в певній послідовності): необхідним і ключовим елементом експерименту є сцена, на якій будуть розміщуватися решта складових. На сцені будуть розміщуватися гідроакустичні датчики та морські об'єкти, що будуть використовуватися в якості джерела шуму. Критичним в цьому модулі є функція для побудову траєкторії руху морського об'єкта, що дозволить дослідникам перевірити, як різні шляхи впливають на поширення звуку.

Модуль аналізу узбережжя та перешкод. Для експериментів поблизу прибережних регіонів або складних підводних рельєфів має бути змога аналізувати взаємодію звуку з прибережними об'єктами та перешкодами, такими як рифи, скелі та підводні утворення. Такий аналіз необхідний для оцінки впливу відбиття звуку, заломлення та перешкоди вздовж траєкторії джерела звуку.

Кожен модуль сприяє спрощеному робочому процесу, перетворюючи складні дані в доступний формат для точного налаштування експерименту та адаптивного моніторингу в реальному часі. Ця структура гарантує, що інструмент на основі ГІС може ефективно підтримувати динамічні гідроакустичні експерименти, допомагаючи подолати розрив між експериментальним формуванням і надійним аналізом даних.

Отже, формування динамічних гідроакустичних експериментів потребує складних інструментів, які керують і візуалізують обширні просторові набори даних, забезпечуючи точний контроль рухомих джерел звуку уздовж заздалегідь визначених траєкторій. Технологія ГІС значно покращує ці завдання шляхом інтеграції даних про навколишнє середовище, забезпечує точне батиметричне картографування. Організувавши ці функції в модульну структуру, інструмент на основі ГІС може спростити процес формування експерименту, підвищити точність даних і, зрештою, розширити дослідницькі можливості в галузі морської акустики та наук про навколишнє середовище.

Список використаних джерел

1. Urick R. J. (1983). Principles of Underwater Sound (3rd ed.). Peninsula Publishing.
2. Jensen F. B., Kuperman W. A., Porter M. B., Schmidt H. (2011). Computational Ocean Acoustics (2nd ed.). Springer.
3. Burrough P. A., McDonnell R. A. (1998). Principles of Geographical Information Systems. Oxford University Press.

Лебедєв А.В., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Федорова Н.В., д.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

АКТУАЛЬНІСТЬ ТА ПРОБЛЕМИ ВИКОРИСТАННЯ МУТАЦІЙНОГО ТЕСТУВАННЯ ДЛЯ АНАЛІЗУ ЯКОСТІ КОДУ СИСТЕМ З МАШИННИМ НАВЧАННЯМ

Системи з машинним навчанням (МН) стали невіддільною частиною багатьох критичних сфер, таких як медицина, фінанси, енергетика та військова промисловість. Кількість систем з МН непинно зростає з кожним роком (рисунк 1) [1]. Враховуючи їх важливість та кількість, якість та надійність моделей МН є критично важливими. Традиційні методи тестування програмного забезпечення часто не враховують специфіку моделей МН, що може призвести до невиявлених помилок та вразливостей.

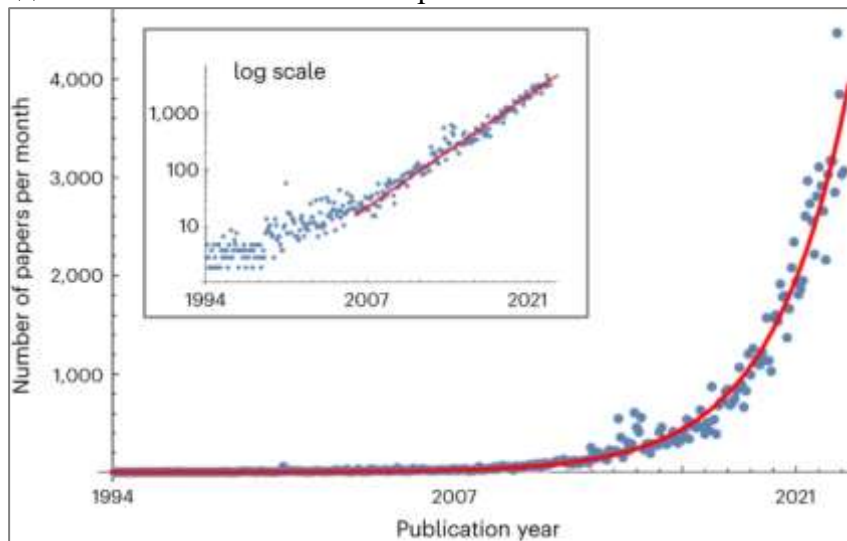


Рисунок 1 – Кількість наукових робіт, що публікуються в категоріях AI та ML

Більшість моделей МН є складними та непрозорими, що ускладнює їх аналіз та тестування. Кількість параметрів та гіперпараметрів може суттєво впливати на результат моделі, а їх повна перевірка є ресурсомісткою [2]. Тому виникає потреба у застосуванні нестандартних підходів для перевірки якості коду моделей МН, таких як мутаційне тестування.

Мутаційне тестування – це метод, який передбачає внесення невеликих змін у вихідний код або моделі з метою оцінки ефективності наявних тестів [3]. Однак, застосування цього методу до систем з МН стикається з кількома викликами:

- 1) відсутність спеціалізованих мутаційних операторів для МН, адже більшість існуючих операторів призначені для традиційного програмного забезпечення і не враховують специфіку моделей МН;
- 2) генерація мутантів для складних моделей може призвести до експоненційного зростання кількості тестів, що вимагає значних обчислювальних ресурсів [4];
- 3) відсутність інструментів, які дозволяють легко інтегрувати мутаційне тестування в процес розробки та CI/CD пайплайни.

Для якісного аналізу систем з моделями МН необхідно мати унікальні оператори, які будуть змінювати структуру моделей, гіперпараметри, функції активації, а також виконувати петрубацію вхідних даних. До цих операторів можна віднести зміну кількості нейронів у шарі, модифікацію функції втрат або застосування шуму до навчальних даних.

Питання великої кількості мутантів згенерованих в результаті виконання аналізу може бути вирішене за допомогою використання алгоритмів машинного навчання для пріоритезації

мутантів за ступенем їх впливу на модель. За допомогою використання даного підходу велика кількість неякісних чи продубльованих мутантів буде виключена з кінцевої вибірки й не матимуть впливу на результат тестування.

Інтеграція аналізу в життєвий цикл розробки разом з візуалізацією та інтерпретацією результатів дозволяють автоматизувати процес мутаційного тестування та спрощують процес аналізу. Тому розв'язання цих проблем є критично важливими частинами системи. До того ж метрики, які враховують зміни в продуктивності моделі, такі як точність, повнота, F1-міра та інші дозволяють кількісно оцінити вплив кожного мутанта на модель та систему загалом.

Окрім відомих викликів з мутаційним тестуванням, існують також ряд практичних проблем з імплементацією системи, які потребують вирішення:

- 1) використання розподілених обчислень для мутаційного тестування великих моделей, адже цей процес може вимагати значних ресурсів;
- 2) перевірка практичності мутаційних операторів, тобто вони мають дійсно моделювати можливі помилки розробників, які мають вплив на якість моделі;
- 3) забезпечення сумісності з різними популярними бібліотеками та фреймворками МН, таких як TensorFlow, PyTorch та інші.

Отже, запропонований підхід до використання мутаційного тестування для систем з МН може значно підвищити якість та надійність моделей. Розробка спеціалізованих мутаційних операторів та інтеграція в життєвий цикл розробки сприятиме більш ефективному виявленню помилок та вразливостей.

Список використаних джерел

1. Krenn M., Buffoni L., Coutinho B., Eppel S., Foster J. G., Gritsevskiy A., Lee H., Lu Y., Moutinho J. P., Sanjabi N., Sonthalia R., Tran N. M., Valente F., Xie Y., Yu R., Kopp M. (2023). Forecasting the future of artificial intelligence with machine learning-based link prediction in an exponentially growing knowledge network. *Nature Machine Intelligence*. <https://doi.org/10.1038/s42256-023-00735-0>
2. Hutter F., Kotthoff L., Vanschoren J. (Ред.). (2019). *Automated machine learning*. Springer International Publishing. <https://doi.org/10.1007/978-3-030-05318-5>
3. Jia Y., Harman M. (2011). An analysis and survey of the development of mutation testing. *IEEE Transactions on Software Engineering*, T. 37 (5), pp. 649–678. <https://doi.org/10.1109/tse.2010.62>
4. Papadakis M., Kintis M., Zhang J., Jia Y., Traon Y. L., Harman M. (2019). Mutation testing advances: An analysis and survey. *У Advances in computers* (pp. 275–378). Elsevier. <https://doi.org/10.1016/bs.adcom.2018.03.015>

Хоменко О.М., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Коваль О.В., д.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

ПІДХОДИ ДО АНАЛІЗУ КАСКАДНИХ ЕФЕКТІВ В РОБОТІ ЕЛЕКТРОМЕРЕЖ

Безперебійна робота критичної інфраструктури є важливою умовою для забезпечення функціонування необхідних послуг (наприклад, комунальні послуги), відсутність яких може негативно вплинути на життя людей. З часом складність критичної інфраструктури збільшується, зв'язки стають складнішими зростає ризик помилки. Надійність роботи енергомережі, адаптація до різних типів збоїв є важливим завданням для забезпечення роботи залежних систем. Небезпечним явищем в енергосистемах, яке виникає зрідка, але спричиняє значні економічні та соціальні наслідки є блекаут. Збій в роботі компонента системи породжує каскад збоїв, що може стати початком масштабних перебоїв і в результаті виникає блекаут. Наприклад, один із генераторів виходить з ладу, потужність в мережі стає незбалансованою, і частота починає зменшуватися. Якщо резервної потужності генерації для забезпечення потреб навантаження недостатньо, то частота буде продовжувати зменшуватися, що може призвести до відключення інших генераторів. Через певний проміжок часу це може призвести до виведення з ладу всієї мережі, відновлення якої може тривати тривалий період часу, що призводить до економічних втрат. Одним із можливих сценаріїв розвитку каскадних збоїв є вихід із ладу лінії електропередачі, в результаті потужність перерозподіляється по мережі і потенційно може призвести до перевантаження інших ліній.

Для розробки стратегій дій щодо уникнення або зменшення наслідків каскадних збоїв в енергомережах аналізуються історичні дані (рисунок 1).

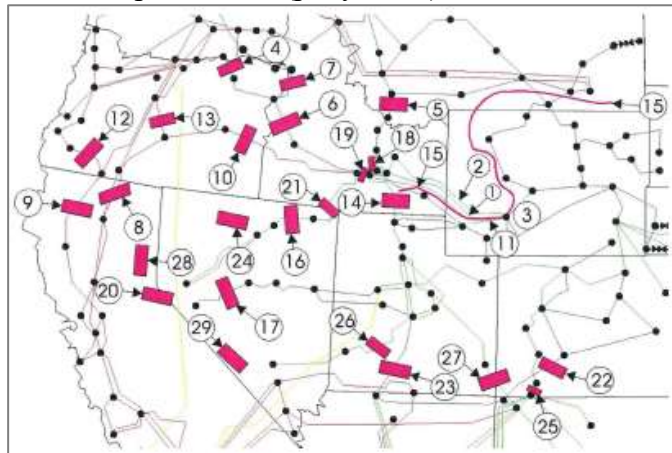


Рисунок 1 – Розвиток каскаду, послідовності подій на прикладі блекаута 2 липня 1996 року на заході США [4]

Приклади виникнення та розвитку блекаутів в електромережах різних країн та їх наслідки:

1) 14 серпня 2003 року в електромережах Сполучених Штатів Америки і Канади був блекаут, від якого постраждали приблизно 50 мільйонів людей, генерація зменшилася приблизно на 61800 мегават (МВт), а вартість знеструмлення оцінюється в 7-14 мільярдів доларів. Електроенергію для більшості споживачів було успішно відновлено протягом кількох годин, у деяких районах Сполучених Штатів електроенергії не було протягом двох днів, а в деяких частинах Онтаріо зазнали постійних відключень електроенергії на два тижні через брак генеруючих потужностей [1];

2) 30 і 31 липня 2012 року в індійській електромережі сталося два блекаута (приблизні втрати генерації 36000 МВт та 48000 МВт), що вплинуло на функціонування інфраструктури, постраждали приблизно 620 мільйонів людей [2];

3) 23 січня 2023 року в Пакистані стався блекаут, що стався через дефіцит реактивної потужності, перевантаження лінії електропередачі та тривав більше 12 годин. Блекаут вплинув на роботу Інтернету та мобільному зв'язку, водяні насоси, які працюють від електрики, не працювали, багато банкоматів перестали працювати, текстильна промисловість зазнала збитків на 70 мільйонів доларів [3].

Нестабільність в роботі електромережі може трансформуватися у катастрофічні наслідки, тому запобігання виникнення збоїв, що потенційно можуть перетворитися в неконтрольований ланцюг подій – каскадний ефект є важливим аспектом в системах електроенергії. Оперативна обробка даних та їхній аналіз є важливим критерієм для прийняття рішень при виникненні збоїв в критичній інфраструктурі. Для запобігання розвитку каскадів в мережах використовують різні підходи. Розглянемо їх далі.

1. Теорія графів. Об'єкти критичної інфраструктури та взаємозв'язки визначаються за допомогою графу, використовуючи матрицю суміжності, матрицю Лапласа. Концепції з теорії графів використовуються для визначення вразливих компонентів в мережі та оцінки її стійкості. Наприклад, в аналізі електромереж використовують теорію складних мереж, використовуючи топологічний та гібридний підходи. Топологічний підхід базується на структурі та її топологічних концепціях, наприклад: average path length, node degree distribution, betweenness, size of the largest component, network efficiency [6]. В гібридному підході використовується поєднання концепцій складних мереж та електротехніки (наприклад, імпеданс лінії) для покращення топологічного підходу, тому метрики із топологічного підходу змінюються на відповідні аналоги: net-ability, electrical degree centrality, electrical betweenness centrality, entropic degree, effective graph resistance [6]. При оцінці міцності графів використовують robustness metrics, які згруповані в шість категорій: clustering, connectivity, distance, throughput, spectral methods, geographical metrics [7]. На вибір метрик для оцінки характеристик графу впливає тип графу, співвідношення між точністю та обмеженнями при обчисленні (деякі алгоритми є NP-складним, але мають апроксимаційні рішення).

2. Мережа Баєса – імовірнісна графова модель, яка є одним із інструментів для аналізу, отримання знань із даних, може використовуватися для моделювання складних недетермінованих систем. Мережа Баєса та її модифікації використовуються для оцінки ризиків, аналізу стійкості критичної інфраструктури [8, 9]. Існують різновиди мережі Баєса, які використовуються у випадках, коли потрібна розширена структура: Causal interaction models, Dynamic Bayesian networks, Influence diagrams [10]. Існують гібридні алгоритми, які поєднують підходи на основі оцінок і обмежень. Навчання мережі Баєса є NP-складною проблемою частково тому, що множина розв'язків графів зростає надекспоненціально (super-exponentially) з кількістю вершин (випадкових величин), тому для побудови мережі Баєса великої розмірності використовують евристичний метод [11]. При збільшенні кількості випадкових величин та їх можливих станів розмір таблиць умовної ймовірності може зростати експоненціально, що ускладнює їх визначення та підтримку. Відсутність даних може спричинити упереджені або неточні результати, а складність структури мережі може ускладнити її інтерпретацію.

3. Мережа Петрі – це орієнтований, зважений дводольний граф, який складається з двох типів вершин, які називаються позиціями та переходами, де дуги йдуть або від місця до переходу, або від переходу до місця. Стан або маркування в мережах Петрі змінюється відповідно до правила переходу [12]. Існують модифікації мереж Петрі, які використовуються для моделювання та аналізу складної динаміки в розумних електромережах (Smart Grid): management, reliability, networking, security (запобігання можливим каскадним ефектам) [13]. Сучасні електромережі складаються з великою кількістю компонентів, а складність зв'язків між ними збільшується, що може ускладнити побудову та підтримку мережі Петрі, ризик виникнення помилки при оновленні моделі.

Каскадні збої поширюються нелокально в енергосистемах: наступний компонент, який вийде з ладу після збою в лінії електропередачі, може бути віддаленим, як географічно, так і

топологічно, що ускладнює процес моделювання та розробку алгоритмів для запобігання та зменшення наслідків збоїв [4, 5].

Машинне навчання та глибоке навчання динамічно розвиваються, розроблюються нові підходи, які можуть бути використані для аналізу каскадних збоїв в енергосистемах [14]. Нейронні мережі демонструють високу ефективність в розпізнаванні, узагальненні патернів. Деякі з моделей можуть бути масштабованими для обробки великих наборів даних, що дозволяє виявляти та прогнозувати потенційні каскадні збої в режимі реального часу. Розробка програмного забезпечення, використовуючи сучасні підходи, архітектури, інструменти, дозволяє підвищити ефективність прийняття рішень експертом для стабілізації роботи системи після виникнення збоїв.

Список використаних джерел

1. August 14, 2003, Northeast Blackout Impacts and Actions and the Energy Policy Act of 2005 David W., Hilt P. E. North American Electric Reliability Council, URL: <https://www.nerc.com/pa/rrm/ea/AugustInvestigation.pdf>.
2. Lai L. L., Zhang H. T., Mishra S., Ramasubramanian D., Lai C. S., Xu F. Y., «Lessons learned from July 2012 Indian blackout», 9th IET International Conference on Advances in Power System Control, Operation and Management (APSCOM 2012), Hong Kong, 2012, pp. 1–6, doi: 10.1049/cp.2012.2173.
3. 2023 Pakistan blackout, URL: https://en.wikipedia.org/wiki/2023_Pakistan_blackout.
4. Hines P. D. H., Dobson I. Rezaei P., «Cascading Power Outages Propagate Locally in an Influence Graph That is Not the Actual Grid Topology», in IEEE Transactions on Power Systems, T. 32, № 2, pp. 958–967, March 2017, doi: 10.1109/TPWRS.2016.2578259.
5. Dobson I., Carreras B. A., Newman D. E., Reynolds-Barredo J. M., «Obtaining Statistics of Cascading Line Outages Spreading in an Electric Transmission Network From Standard Utility Data», in IEEE Transactions on Power Systems, T. 31, № 6, pp. 4831–4841, Nov. 2016, doi: 10.1109/TPWRS.2016.2523884.
6. Cuadra L., Salcedo-Sanz S., Del Ser J., Jiménez-Fernández S., Geem Z. W., «A Critical Review of Robustness in Power Grids Using Complex Networks Concepts», Energies 2015, 8, 9211–9265. <https://doi.org/10.3390/en8099211>.
7. Oehlers M., Fabian B., «Graph Metrics for Network Robustness – A Survey», Mathematics 2021, 9, 895. <https://doi.org/10.3390/math9080895>.
8. Wright M., Chizari H., Viana T., «A Systematic Review of Smart City Infrastructure Threat Modelling Methodologies: A Bayesian Focused Review», Sustainability. 2022, 14, 10368. <https://doi.org/10.3390/su141610368>.
9. Eldosouky A., Saad W., Mandayam N., «Resilient critical infrastructure: Bayesian network analysis and contract-Based optimization», Reliability Engineering & System Safety, T. 205, January 2021, 107243. <https://doi.org/10.1016/j.ress.2020.107243>.
10. Daly R., Qiang S., Aitken S., «Learning Bayesian Networks: Approaches and Issues», Knowledge Engineering Review, T. 26, № 2, pp. 99–127, 2011. <https://doi.org/10.1017/S0269888910000251>.
11. Kitson N. K., Constantinou A. C., Zhigao G., Liu Y., Chobtham K., «A survey of Bayesian Network structure learning», Artificial Intelligence Review T. 56, pp. 8721–8814, 2023. <https://doi.org/10.1007/s10462-022-10351-w>.
12. Murata T., «Petri nets: Properties, analysis and applications», in Proceedings of the IEEE, T. 77, № 4, pp. 541–580, April 1989. <https://doi.org/10.1109/5.24143>.
13. Ge M., Rossi B., Chren S., Blanco J. M., «Petri Nets for Smart Grids: The Story So Far», SAC '24: Proceedings of the 39th ACM/SIGAPP Symposium on Applied Computing, pp. 661–670. <https://doi.org/10.1145/3605098.3635989>.
14. Md Sami N., Naeini M., «Machine Learning Applications in Cascading Failure Analysis in Power Systems: A Review», Electric Power Systems Research, T. 232, July 2024, 110415. <https://doi.org/10.1016/j.epsr.2024.110415>.

Ляшко І.І., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Залевська О.В., к.т.н, Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

КОНЦЕПТУАЛЬНІ ОСНОВИ ТА ПЕРСПЕКТИВИ РОЗРОБКИ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ ДЛЯ РОЗПІЗНАВАННЯ МІМІКИ У ВІРТУАЛЬНИХ СЕРЕДОВИЩАХ НА ОСНОВІ АНТРОПОМОРФНИХ ДАНИХ

Реалістичне відтворення міміки у віртуальних середовищах є ключовим для покращення якості взаємодії та створення ефекту присутності. Однак існуючі методи, що базуються лише на камерах, часто не забезпечують достатньої точності. Інтеграція даних з камер і LiDAR-сенсорів, у поєднанні з нейронними мережами, дозволяє створити точну модель міміки, адаптовану до індивідуальних виразів користувача. Це дослідження пропонує підхід, який дозволяє в реальному часі проєктувати унікальну міміку користувача на аватарах, забезпечуючи природну та інтуїтивну взаємодію.

Віртуальні середовища активно впроваджуються у різних сферах, таких як освіта, віддалена співпраця, терапія та медична реабілітація, де точне відтворення міміки користувачів допомагає підвищити ефективність комунікації та створює відчуття присутності. Реалістична передача емоцій у VR дозволяє віддаленим учасникам відчувати природнішу взаємодію, що є особливо важливим у дистанційних командах і навчальних платформах, де точне вираження емоцій та міміки впливає на загальне сприйняття інформації та емоційний зв'язок між учасниками. Особливо в соціальних мережах і відеоіграх реалістичні аватари дозволяють користувачам занурюватися в нові віртуальні досвіди, підтримуючи природну емоційну експресію. Важливість цього аспекту підтверджується дослідженнями у галузі телеприсутності, які показують, що аватари, що відтворюють унікальні вирази обличчя, значно покращують соціальну взаємодію та відчуття зв'язку в цифровому просторі.

Згідно з прогнозами, глобальний ринок віртуальної реальності зростає з 15,9 мільярда доларів США у 2024 році до 38,0 мільярда доларів США у 2029 році, демонструючи середньорічний темп зростання (CAGR) 19,1 % протягом цього періоду [1]. Це зростання підкреслює зростаючу потребу в реалістичних та ефективних засобах комунікації у віртуальних середовищах. Крім того, дослідження показують, що точне відтворення міміки у VR може підвищити ефективність навчання на 30 %, покращуючи розуміння та запам'ятовування інформації [2]. У сфері віддаленої співпраці реалістичні аватари сприяють підвищенню продуктивності команд на 25 %, завдяки покращенню невербальної комунікації та зменшенню непорозумінь [3]. Ці показники підкреслюють важливість розробки технологій, що забезпечують точне відтворення міміки у віртуальних середовищах, для задоволення зростаючих ринкових потреб та підвищення якості взаємодії користувачів. Віртуальні середовища набувають широкого застосування в різних галузях, зокрема в освіті, віддаленій співпраці, терапії та медичній реабілітації. Згідно з дослідженням, опублікованим у журналі *Frontiers in Digital Health* [4], використання віртуальної та доповненої реальності в медичній освіті сприяє покращенню навчальних результатів і навичок медичних працівників.

Розробка систем розпізнавання та візуалізації міміки у віртуальних середовищах стикається з низкою викликів. По-перше, використання камер для захоплення виразів обличчя може бути обмеженим через варіації освітлення, положення голови та часткове перекриття обличчя, що впливає на точність розпізнавання емоцій. По-друге, інтеграція даних з різних сенсорів, таких як камери та LiDAR, потребує складних алгоритмів фузії для забезпечення узгодженості та точності тривимірних моделей обличчя. Крім того, обробка великих обсягів даних у реальному часі вимагає значних обчислювальних ресурсів, що може призвести до затримок, неприйнятних для інтерактивних застосувань. Також, існують етичні та приватні питання, пов'язані з використанням біометричних даних, що потребують ретельного розгляду та дотримання відповідних норм і стандартів. Нарешті, забезпечення універсальності системи для різних користувачів з унікальними мімічними особливостями є складним завданням,

оскільки необхідно враховувати індивідуальні відмінності у виразах обличчя та емоційних реакціях.

Запропонована концепція системи розпізнавання та візуалізації міміки для віртуальних середовищ передбачає інтеграцію даних з камер та LiDAR-сенсорів, використання нейронних мереж для розпізнавання унікальних емоційних виразів користувача та трансляцію цих даних на стандартизовану «основу» для аватарів. Така архітектура дозволяє адаптувати систему до різних типів аватарів, зберігаючи при цьому високу точність та швидкодію. Основні компоненти системи включають:

- система збору даних та попередня обробка: використання камер та LiDAR для захоплення зображень та тривимірної структури обличчя;
- модуль попередньої обробки та фільтрації даних: синхронізація та очищення даних для забезпечення їхньої узгодженості;
- фузія даних і створення універсальної основи для аватара: об'єднання даних з різних сенсорів для формування стандартизованої 3D-сітки обличчя;
- розпізнавання міміки за допомогою нейронних мереж: аналіз даних для визначення емоційних виразів користувача;
- трансляція та рендеринг міміки на аватарі: передача оброблених даних на віртуальний аватар для відображення міміки;
- адаптивне навчання та вдосконалення моделі: постійне оновлення системи для підвищення точності та персоналізації.

Система збору даних та попередня обробка. На першому етапі використовуються камери високої роздільної здатності для зчитування зображень обличчя та LiDAR-сенсори для захоплення тривимірної структури. Використання двох джерел дозволяє компенсувати недоліки кожного сенсора: камери забезпечують деталізацію кольорових зображень, тоді як LiDAR надає точні дані глибини та форми обличчя, незалежні від освітлення та позиції користувача. Це дає змогу отримати більш комплексне уявлення про об'єкт, покращуючи точність відтворення міміки [5].

Попередня обробка та фільтрація даних. Дані з камер і LiDAR синхронізуються, після чого проходять через фільтр Калмана або інші методи згладжування для усунення шуму і підготовки до фузії. Така попередня обробка є необхідною для забезпечення узгодженості даних між двома сенсорами та підвищення стабільності системи, особливо у реальному часі.

Фузія даних і створення універсальної основи для аватара. Фузія даних з камер та LiDAR реалізується за допомогою алгоритмів згладжування та спільного вирівнювання тривимірних точок обличчя, що дозволяє створити точну модель обличчя користувача. Універсальна «основа» являє собою стандартизовану 3D-сітку обличчя з контрольними точками, яка може бути адаптована до будь-якого аватара. Це дозволяє передавати міміку користувача без потреби у попередньо заданих позах, забезпечуючи точну трансляцію унікальних емоцій [6].

Розпізнавання міміки за допомогою нейронних мереж. Для розпізнавання міміки система використовує згорткові нейронні мережі (CNN) для аналізу зображень обличчя з камер, а також 3D-CNN або рекурентні нейронні мережі (RNN) для обробки LiDAR-даних, що дозволяє визначати позиції та зміну контрольних точок обличчя у тривимірному просторі. Завдяки поєднанню двовимірних і тривимірних даних система отримує повнішу картину міміки, що дозволяє максимально точно передавати унікальні вирази користувача.

Адаптивне навчання та вдосконалення моделі. На завершальному етапі система використовує механізми адаптивного навчання для автоматичного покращення точності. Модель, що базується на машинному навчанні, постійно оновлюється та коригується, враховуючи індивідуальні особливості міміки конкретного користувача, що підвищує рівень персоналізації та реалістичності.

Розробка універсальної системи для розпізнавання та візуалізації міміки користувачів у віртуальних середовищах є ключовою для створення реалістичних взаємодій та покращення ефекту присутності. Застосування інтеграції даних з камер і LiDAR-сенсорів у поєднанні з

нейронними мережами дозволяє точно передавати індивідуальні емоційні вирази користувача на віртуальних аватарах. Ця концепція забезпечує комплексний підхід до аналізу міміки, що включає збір даних, попередню обробку, фузію інформації, адаптацію під різні аватари та постійне навчання. Враховуючи зростаючу популярність віртуальних середовищ у сферах освіти, віддаленої співпраці та розваг, така система є актуальною та перспективною для подальших досліджень і впровадження.

Список використаних джерел

1. Markets and Markets, «Virtual Reality Applications Market by Technology, Component, Device Type, Application, and Geography – Global Forecast to 2029», Markets and Markets Research Report, 2023.
2. Johnson A., Lee T., «Effectiveness of Facial Expression Recognition in Virtual Learning», Frontiers in Psychology, 2023.
3. Miller S., Wang Y., «Realistic Avatars and Non-Verbal Communication in Virtual Collaboration», Computers & Graphics, 2022, T. 00371-022-02700-1.
4. Smith J., Brown R., «The Use of Virtual Reality in Medical Education: Enhancing Learning Outcomes», Frontiers in Digital Health, 2024.
5. Li Yingwei, Wei Yu Adams, Meng Tianjian, Caine Ben, Ngiam Jiquan, Peng Daiyi, Shen Junyang, Wu Bo, Lu Yifeng, Zhou Denny, Quoc V. Le, Yuille Alan, Tan Mingxing, «DeepFusion: Lidar-Camera Deep Fusion for Multi-Modal 3D Object Detection», 2022.
6. Khandaker Rahad Hossain, Reza Malekian, Oludare Isaac Abayomi, «Multi-Sensor Data Fusion for Real-Time Multi-Object Tracking» Processes, T. 11, № 2, p. 501, 2023.

Садовська І.І., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Коваль О. В., д.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

СИСТЕМА РОЗШИРЕННЯ АБО ЗЛИТТЯ БАЗ ЗНАНЬ

Розроблена система призначена для полегшення процесу злиття онтологій, які, в свою чергу, є якісним інструментом для структуризації інформації певної предметної області, кількість якої постійно збільшується [1]. Процес злиття онтологій у більшості існуючих систем реалізований за рахунок надання можливості користувачу об'єднувати дані онтологій вручну, що потребує значних зусиль та часу. Це підкреслює важливість автоматизованих систем, здатних здійснювати злиття онтологій з мінімальною участю людини [2]. Розроблена система спрямована на те, щоб мінімізувати участь людини в процесі злиття та максимально автоматизувати цей процес.

В основі розробленої системи є алгоритм злиття онтологій. Він включає наступні етапи:

- 1) попередню обробку онтологій (визначення їх сумісності та групування компонентів онтологій);
- 2) попарне злиття класів та властивостей (перевірка їх семантичної сумісності, вирішення конфліктів, що виникли в результаті, та відповідне додавання елементів завантаженої онтології в базову);
- 3) злиття індивідів (порівняння індивідів основної онтології з індивідами завантаженої, у випадку їх сумісності – додавання нових фактів (statements), що розширюють онтологію; в іншому випадку – повне копіювання індивіда в базову онтологію);
- 4) валідація та завершення роботи алгоритму (перевірка узгодженості результуючої онтології за допомогою різонера Pellet [3], збереження онтології в разі валідності результуючої онтології).

Варто зазначити, що для забезпечення максимальної точності на етапі перевірки семантичної сумісності онтологій та їх компонентів, було використано нейронну мережу, яка здатна аналізувати семантичну сумісність в контексті застосування компонентів онтології.

Розроблена система складається з кількох ключових компонентів:

- 1) серверна частина застосунку: включає інструменти для взаємодії з клієнтською частиною додатку, а також модулі для обробки та злиття онтологій;
- 2) база даних: застосовується для зберігання файлів онтологій та подальшої взаємодії з ними;
- 3) веб-інтерфейс користувача: забезпечує зручність процесу злиття та перегляду його результатів.

Графічний інтерфейс користувача системи складається з наступних компонентів:

- 1) форма для завантаження онтології для злиття – власної або ж обрання з наявних автозгенерованих онтологій;
- 2) діалогові вікна для вирішення конфліктів;
- 3) таблиця для відображення результату злиття, яка містить дані про всі додані в базову онтологію компоненти (рисунок 1);
- 4) кнопка для завантаження результуючої онтології для подальшої взаємодії з нею.

У контексті системи для оцінки наукової роботи викладачів кафедри, розроблена система дає можливість збільшити знання системи та дати користувачу можливість отримати більшу кількість знань. Важливо вказати, що API розробленої системи є універсальним та може бути використаним для злиття будь-яких двох семантично сумісних онтологій, включно з тими що не стосуються предметної області наукової роботи працівників кафедри.

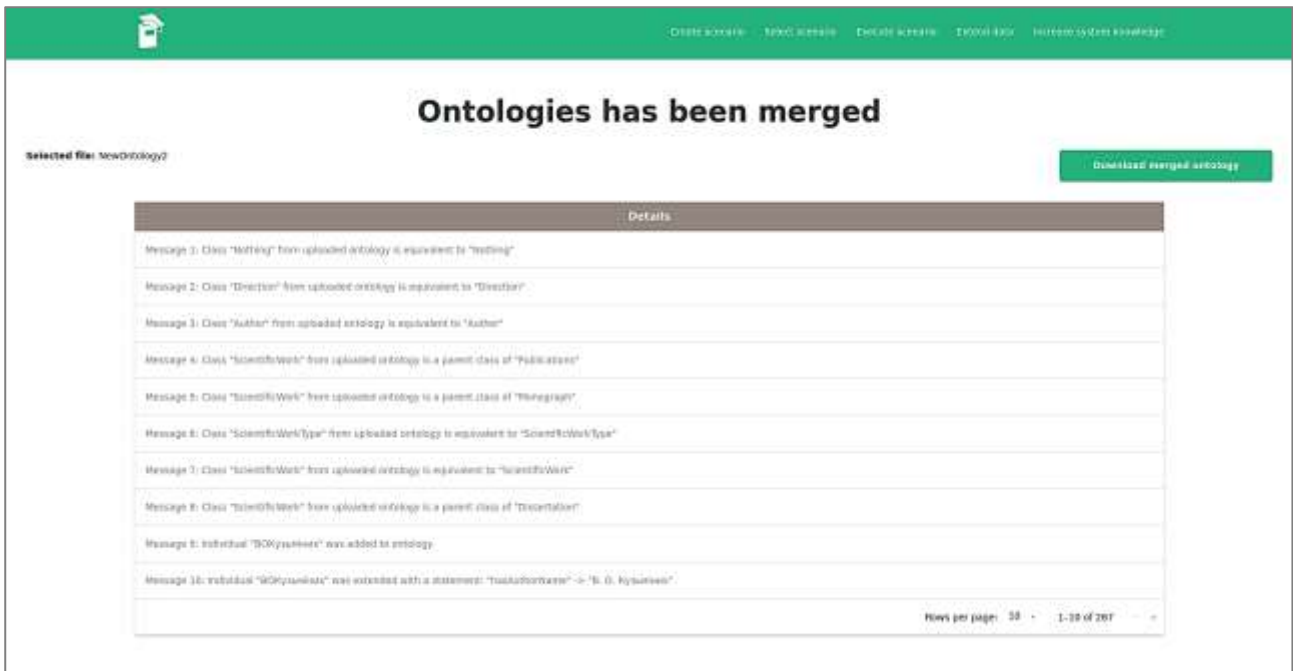


Рисунок 1 – Приклад відображення результату злиття онтологій

Отже, розроблена система автоматизує процес злиття онтологій, знижуючи потребу в ручній обробці та забезпечуючи актуальність баз знань. Вона є важливим інструментом для розширення онтологічних даних і інтеграції знань у більші системи, полегшуючи їхню підтримку та модернізацію.

Список використаних джерел

1. Gruber T. Knowledge Acquisition. 2nd ed. Palo Alto : Knowledge System Laboratory, Stanford University, 1993. Т. 5.
2. Euzenat J., Shvaiko P. Ontology Matching. 2nd ed. Berlin : Springer, 2013. 511 p. URL: <http://book.ontologymatching.org/>.
3. Pellet: A practical OWL-DL reasoner / E. Sirin et al. Journal of Web Semantics. 2007. Т. 5, № 2. pp. 51–53. URL: <https://doi.org/10.1016/j.websem.2007.03.004>.

Передера В.Р., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Недашківський О.Л., д.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

АКТУАЛЬНІСТЬ, ПРОБЛЕМИ ТА НЕТОЧНОСТІ ІСНУЮЧИХ МЕТОДІВ КЛАСИФІКАЦІЇ ТА ВИЯВЛЕННЯ ДЖЕРЕЛ ІНФОРМАЦІЇ В ІНТЕРНЕТІ НА ОСНОВІ НЕЙРОННИХ МЕРЕЖ

Актуальність дослідження методів класифікації та виявлення джерел інформації в Інтернеті зумовлена стрімким зростанням обсягу даних, які постійно генеруються в цифровому середовищі. За даними досліджень за 2023-й рік обсяг інформації в мережі Інтернет збільшився на 27 ЗБ і сягне позначки в 147 ЗБ до кінця року, а до кінця 2025-го цей показник становитиме вже 181 ЗБ [1].

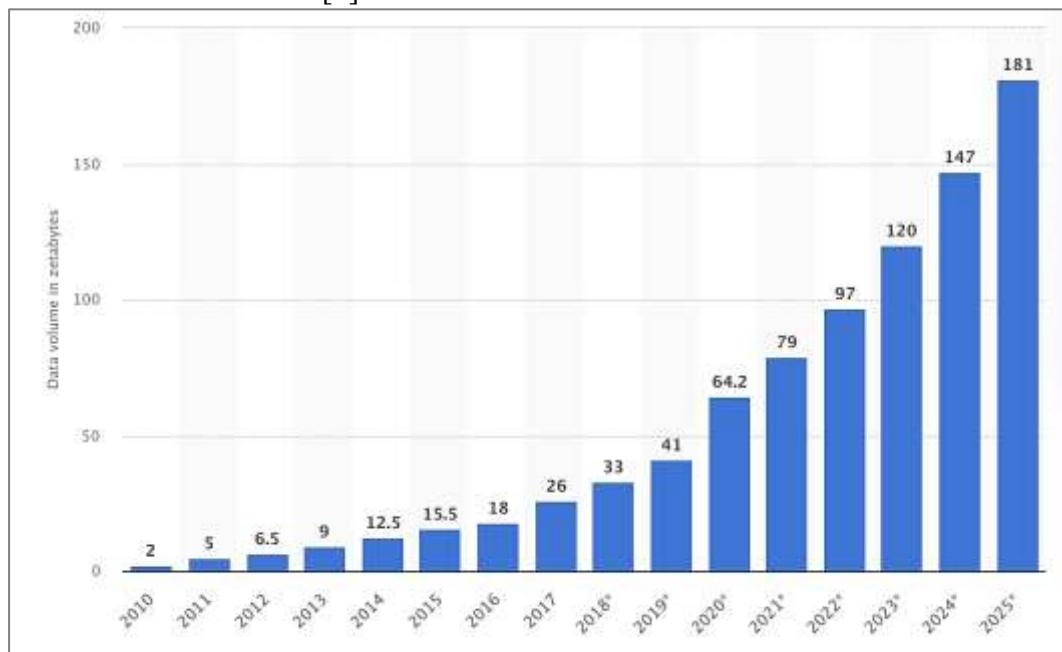


Рисунок 1 – Графік щорічного зростання обсягу інформації в мережі інтернет

Такий обсяг нової інформації одночасно збільшує кількість фейкового контенту і ускладнює його класифікацію та виявлення джерел. Водночас з цим зростає роль фактчекінгу як важливого інструменту для забезпечення достовірності інформації. Цей напрям значно розвинувся: якщо у 2008 році в Інтернеті було лише 11 сайтів, що займалися перевіркою фактів, то у 2022 році їхня кількість сягнула 424. Попри це, з 2019 року кількість нових сайтів для перевірки інформації поступово знижується. Станом на 2023 рік існує 417 активних платформ, що охоплюють понад 100 країн і працюють 69 мовами [2]. Це свідчить про значний інтерес до питання достовірної інформації, хоча тенденція до зменшення кількості нових сайтів підкреслює потребу в удосконаленні та розширенні існуючих ініціатив. У цьому контексті нейронні мережі, завдяки своїй здатності навчатися на великих обсягах даних, демонструють значний потенціал для вирішення цих завдань.

Водночас різноманітність видів неправдивої інформації ускладнює завдання її виявлення та класифікації. Зокрема, можна виділити кілька типів недостовірного контенту, кожен з яких вимагає спеціального підходу: помилкова інформація, що поширюється через недостатню обізнаність; оманлива, яка навмисно подається із спотворенням фактів; упереджена, яка відображає суб'єктивну точку зору; маніпулятивна, що використовується для досягнення певних цілей; а також гумористична, яка може вводити в оману аудиторію під виглядом жартів [3]. Розуміння цих категорій є критично важливим для створення ефективних

алгоритмів нейронних мереж, що допоможуть вчасно виявляти та класифікувати різні види дезінформації.

На сьогодні методи виявлення фейкових новин, з урахуванням основних характеристик, що використовуються в класифікаційних моделях, поділяються на три основні категорії: методи, засновані на аналізі контенту, методи на основі соціальних мереж та методи, що використовують бази знань [4].

Методи виявлення фейкових новин на основі аналізу контенту передбачають виокремлення різних семантичних ознак тексту новин, щоб оцінити їхню достовірність. Ці методи базуються на припущенні, що фейкові новини мають специфічні лінгвістичні особливості, які відрізняють їх від правдивих матеріалів. Зокрема, дослідження показують, що лінгвістичні особливості, такі як займенники першої та другої особи, слова для перебільшення (підмети, перехідні слова, модальні прислівники), частіше зустрічаються у фейкових новинах, тоді як правдиві новини характеризуються більш об'єктивною мовою, використанням чисел, займенників третьої особи та конкретних слів [5].

Метод виявлення фейкових новин на основі знань (Knowledge-based, KB) перевіряє достовірність новин, зіставляючи їх з вже відомими фактами, тому цей процес також називають фактчекінгом. Фактчекінг поділяється на два типи: ручна та автоматизована перевірка [6]. Ручна перевірка передбачає використання знань експертів у певній галузі або краудсорсинг. Цей метод забезпечує високу точність, однак має низьку ефективність і не відповідає вимогам обробки даних у сучасну епоху великих даних.

Контент-орієнтований підхід дозволяє виявляти лінгвістичні особливості, що відрізняють правдиві новини від фейкових. Однак фейкові новини іноді вводять читачів в оману, імітуючи стиль справжніх новин, що ускладнює їхнє розпізнавання за допомогою лише аналізу контенту. Щоб подолати це обмеження, можна використовувати додаткову інформацію, таку як дані про соціальний контекст і шляхи поширення новин у соціальних мережах. Крім аналізу контенту та соціального контексту, ключову роль у визначенні достовірності новин відіграє ідентифікація джерела інформації. Джерело може виступати важливим маркером правдивості, оскільки надійні джерела, як правило, мають репутацію, що базується на точності та об'єктивності поширюваної інформації. Перевірені новинні агентства, офіційні сайти та сторінки авторитетних організацій мають меншу ймовірність поширювати фейкову інформацію. У той же час ненадійні або маловідомі джерела часто можуть бути пов'язані з пропагандою або поширенням неправдивих даних. Виявлення джерела новини дозволяє не лише оцінити її достовірність, але й виявити потенційні наміри авторів новини, наприклад, якщо джерело має відомі політичні чи комерційні інтереси. Таким чином, інтеграція аналізу джерел інформації як додаткового компонента до контент-орієнтованих та соціально-орієнтованих методів може суттєво підвищити точність у виявленні фейкових новин. Цей підхід дозволяє не лише визначати лінгвістичні чи поведінкові особливості, але й враховувати репутацію і специфіку джерел, що в свою чергу додає багатовимірності процесу верифікації новинного контенту. Критичну роль в цьому процесі відіграє час та швидкість поширення інформації в мережі. За результатами дослідження, що використовує модель незалежного каскаду як метод виявлення джерел інформації в мережах виявлено, що якщо час зараження менший за певний показник, алгоритм точно визначає джерело з ймовірністю 1, особливо у великих мережах, але щойно час зараження перевищує поріг $\lceil \log n / \log \mu \rceil + 2$ (де n – кількість вузлів мережі, μ – середній ступінь зв'язку вузлів у мережі.) ймовірність виявлення джерела стає близькою до нуля [7].

Проблема підходу виявлення джерел інформації на основі моделі незалежного каскаду (IC) полягає в тому, що він має обмеження в умовах тривалої дії інформаційного «зараження». Як показують результати, після певного порогу тривалості зараження ймовірність успішної ідентифікації джерела інформації зменшується до нуля, що ускладнює виявлення фейкових новин або недостовірних джерел. Це пов'язано з тим, що інформація може швидко розповсюджуватися в мережі, і важко відстежити початкове джерело. Для покращення цього

підходу можна застосувати кілька стратегій, такі як включення додаткових параметрів, аналіз структурних характеристик мережі, факти перевірки та модернізація алгоритмів.

У результаті проведеного аналізу актуальності, проблем та перспектив розвитку методів класифікації та виявлення джерел інформації на основі нейронних мереж, можна зробити кілька важливих висновків. По-перше, зростаючі обсяги даних в Інтернеті підкреслюють необхідність у вдосконаленні існуючих підходів, оскільки традиційні методи не здатні впоратися з інформаційними викликами сучасності. По-друге, існуючі моделі, незважаючи на їхній потенціал, мають суттєві недоліки, зокрема зменшення ефективності ідентифікації джерела із збільшенням часу розповсюдження інформації. Високі показники помилок у класифікації свідчать про необхідність комплексного підходу до розробки нових алгоритмів.

Список використаних джерел

1. Chen, Zeqiu & Zhou, Jianghui & Sun, Ruizhi. (2023). An efficient content extraction method for webpage based on tag-line-block analysis. *Soft Computing*. T. 27. pp. 1–15. 10.1007/s00500-023-09076-x.
2. Гнатишин М. С., Недашківський О. Л. Методи і програмне забезпечення для виявлення неправдивої інформації в мережі Інтернет на основі нейронних мереж. Перша міжнародна науково-практична конференція «Сучасні аспекти інженерії програмного забезпечення (Modern aspects of Software Engineering)»: Науковий збірник (м. Київ, 14 грудня 2023 р.). Київ: КПІ ім.Ігоря Сікорського, 2023. С. 27–31.
3. Гнатишин М. С., Недашківський О. Л. Аналіз методів та програмного забезпечення для виявлення неправдивої інформації у мережі інтернет на основі нейронних мереж. *Журнал «Зв'язок»*. Випуск 4. 2024. С. 52-57. DOI: 10.31673/2412-9070.2024.045257.
4. Lu Y., Hangshun J., Hao S., Lei S., Nanchang C. Sustainable Development of Information Dissemination: A Review of Current Fake News Detection Research and Practice, *Systems* 2023, T. 11 (9), p. 458.
5. Rashkin H., Choi E., Jang, J. Y., Volkova S., Choi Y. Truth of varying shades: Analyzing language in fake news and political fact-checking. In *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*, Copenhagen, Denmark, 9–11 September 2017; pp. 2931–2937.
6. Zhou X., Zafarani R. A survey of fake news: Fundamental theories, detection methods, and opportunities. *ACM Comput. Surv. (CSUR)* 2020, T. 53, pp. 1–40.
7. Kai Z., Lei Y. Information source detection in networks: Possibility and impossibility results. *IEEE INFOCOM 2016 – The 35th Annual IEEE International Conference on Computer Communications*. DOI: 10.1109/INFOCOM.2016.7524363

Свтушенко Д.М. Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Донець А.Г., к.ф.-м.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Колумбет В.П., д.ф., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

ЗАСТОСУВАННЯ РОЗШИРЕНЬ БРАУЗЕРІВ GOOGLE CHROME ТА MOZILLA FIREFOX ДЛЯ АНАЛІЗУ ЗМІСТУ ВЕБ-СТОРИНОК ТА ФОРМУВАННЯ СТАТИСТИКИ В ІНФОРМАЦІЙНІЙ СИСТЕМІ ТЕСТУВАННЯ ТА ОЦІНЮВАННЯ ЗНАНЬ

Постановка проблеми. Дистанційне навчання – це навчальний процес, під час якого абітурієнти, студенти та викладачі знаходяться на відстані один від одного та взаємодіють за допомогою інтернету. Спільна робота може бути організована у різних формах. В наш час інтернет-технології дозволили навчатися дистанційно більшості бажаючих, створивши величезну мережу з безпрецедентною кількістю інформації та залучених до навчання студентів та викладачів. Онлайн-навчання стрімко з'явилося, заявило про себе та продовжує впевнено займати свою частку на ринку освітніх послуг планети [1]. На сьогоднішній день, щоб досягти успіху в будь-якій справі, необхідно постійно розвиватися, вчитися та опановувати нові знання та інформацію, підвищуючи свою професійну планку – все це дозволяє зробити можливим дистанційну освіту. На цей час переважна більшість людей, що здійснює дистанційне навчання у мережі Інтернет або працює зі спеціалізованими веб-програмами, використовують веб-браузери. Велика кількість користувачів для розширення функціональних можливостей веб-браузера використовують розширення веб-браузера. Створення розширення для веб-браузерів Google Chrome та Mozilla Firefox для обробки змісту веб-сторінок і формування статистики, яке в подальшому може бути використано як основа для створення циклу лабораторних робіт, є важливим і актуальним технічним завданням, що є невід'ємною частиною сучасних веб-технологій та відповідає світовим тенденціям [2]. В перспективі це дасть можливість створення уніфікованій бази освіти (ЄДЕБО, бази окремих навчальних закладів, освітніх підрозділів тощо) з сформованими картками споживачів будь-яких освітніх послуг.

Основна частина. Розширення веб-браузера можуть виконувати різні завдання, найпопулярнішими серед яких є блокування реклами, адаптаційні зміни інтерфейсу, переклад веб-сторінок, збір інформації під час перегляду веб-сторінок про їхній вміст. На даний час використання розширень доступне майже у всіх найпопулярніших веб-браузерах, а саме в Microsoft Edge, Google Chrome, Safari, Mozilla Firefox та інших. Використання розширень, як і інших механізмів розширення функціоналу браузерів, має ряд переваг та недоліків, які обумовлені їх технічними особливостями. Серед переваг можна відокремити те, що ці додатки здатні працювати на різних платформах, де використовується веб-браузер.

На сьогоднішній день, лідируючі позиції серед браузерів посідають такі веб-браузери, як: Microsoft Edge, Google Chrome, Safari, Mozilla Firefox. Для визначення рейтингу веб-браузера важливо враховувати такі критерії, як: швидкодія, безпека, підтримка розширень, доступність на різних платформах, популярність використання серед користувачів, функціональні можливості. Серед браузерів найбільшу популярність має Google Chrome, це пояснюється багатьма критеріями [3].

Зробивши аналіз популярних браузерів для створення розширень доцільно вибрати Mozilla FireFox та Google Chrome. Це можна пояснити розповсюдженістю цих веб-браузерів, оскільки вибір цих браузерів забезпечить максимальну аудиторію для використання розширення. Також ці браузери надають інтерфейси для створення розширень, які є задокументованими, що сприяє полегшенню роботи з ними. API веб-браузера мають високий рівень функціональності, що сприяє створенню складних та функціонально широких розширень. Mozilla FireFox та Google Chrome мають досить розвинену комунікацію між

розробниками і також ці групи розробників налічують велику кількість, це сприяє у разі потреби можливості отримати допомогу від інших розробників.

Серверна частина веб-застосунку розуміється як операції, які виконуються сервером у відносинах клієнт-сервер у комп'ютерній мережі. Зазвичай сервер – це комп'ютерна програма, наприклад веб-сервер, що працює на віддаленому сервері, доступному з локального комп'ютера або робочої станції користувача. Операції можуть виконуватися на стороні сервера, оскільки вони вимагають доступу до інформації або функцій, які недоступні на клієнті, або вимагають типової поведінки, яка є ненадійною, коли це виконується на стороні клієнта. Операції на стороні сервера також включають обробку та зберігання даних від клієнта до сервера, які може переглядати група клієнтів (рисунк 1).



Рисунок 1 – Клієнт-серверна частина

Відтворення на стороні сервера або сценарії на стороні сервера – це техніка, яка використовується при розробці веб-сайтів із динамічними елементами та веб-додатків. Він заснований на використанні сценаріїв, які виконуються веб-сервером за допомогою відповідних мов сценаріїв, коли ви запитуєте відповідний вміст у браузері. Після початкового запиту всі інструкції HTML, CSS і JavaScript повністю завантажуються.

Рендеринг на стороні сервера здійснюється майже виключно розробниками, що писали програми на C і Perl, а також і в сценаріях командного рядка. Ці програми були виконані та інтерпретовані операційною системою сервера, після чого результат міг передаватися з веб-сервера до браузера, що здійснює доступ через загальний інтерфейс шлюзу (CGI).

Серверний сценарій обробляється на веб-сервері, коли користувач запитує інформацію. Такі сценарії можуть запускатися до завантаження веб-сторінки. Вони потрібні для всього, що вимагає динамічних даних, наприклад для зберігання даних для входу користувача. Деякі поширені серверні мови включають PHP, Python, Ruby і Java, які працюють як мови програмування на сервері [4].

Коли серверний сценарій обробляється, запит надсилається на сервер, а результат повертається клієнту. Це корисно для веб-сайтів, які зберігають великі обсяги даних, таких як пошукові системи чи соціальні мережі – клієнтський браузер завантажуватиме всі дані дуже повільно.

Висновки. Розширення веб-браузера є одним із ключовим інструментів для забезпечення індивідуальних потреб і уподобань користувача при роботі з Інтернет-ресурсами та веб-програмами. Розробка розширень вимагає від розробника вільного володіння наявною технічною документацією, де наведена інформація про загальні етапи створення розширень та використання програмного інтерфейсу API відповідного веб-браузера. Продемонстровано, що розширення з однаковим функціоналом може працювати водночас у двох найпоширеніших веб-браузерах, а саме Google Chrome та Mozilla Firefox, але процедура його встановлення може бути різною. Розширення здатне аналізувати вміст веб-сторінок та проводити маніпуляції з

текстовими елементами активної вкладки веб-браузера, на якій працює це розширення, наприклад виконувати заміну слів або фраз.

Список використаних джерел

1. Веллінг Л., Томсон Л. Розробка веб-додатків за допомогою PHP і MySQL. К.: Діалектика. 2019. 768 с.
2. Documentation for the development of extensions for the Mozilla FireFox. URL: <https://developer.mozilla.org/en-US/docs/Mozilla/Add-ons>
3. Publish in the Chrome Web Store. URL: <https://developer.chrome.com/docs/webstore/publish>
4. Татро К., Макинтайр П. Програмування PHP: Створюємо динамічні веб-сторінки. Львів, Сяйво. 2020. 719 с.

Клименко Я.В., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Свинчук О.В., к.ф.-м.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

МЕТОДИ І ПРОГРАМНЕ ЗАБЕЗПЕЧЕННЯ ДЛЯ МОДЕЛЮВАННЯ ПРОЦЕСІВ ТЕПЛООБМІНУ І ГІДРОДИНАМІКИ В ТУРБІННИХ СИСТЕМАХ

У сучасній енергетиці турбінні системи відіграють ключову роль у виробництві електроенергії. Точне моделювання процесів теплообміну та гідродинаміки мастила є критичним для забезпечення надійної та ефективної роботи турбін.

Існуючі програмні рішення не завжди можуть врахувати всі специфічні властивості турбінного мастила, такі як нелінійна в'язкість, залежність властивостей від температури та тиску. Це може призводити до неточностей у моделюванні та, як наслідок, до неефективної роботи обладнання.

Метою дослідження є розробка програмне забезпечення для моделювання процесів теплообміну та гідродинаміки перебігу турбінного мастила, яке забезпечить високу адекватність математичних моделей.

Для досягнення мети були поставлені наступні задачі:

- проаналізувати існуючі методи та програмні засоби моделювання процесів теплообміну та гідродинаміки;
- розробити власний програмний пакет з урахуванням специфічних властивостей мастила.
- провести валідацію розробленої моделі на основі експериментальних даних.

Проведено детальний огляд наукової літератури з питань моделювання теплообміну та гідродинаміки в системах з турбінним мастилом [1]. Вивчено можливості та обмеження існуючих програмних пакетів, таких як ANSYS Fluent, COMSOL Multiphysics та OpenFOAM.

При розробці програмного забезпечення були використані наступні математичні моделі:

- метод скінченних елементів (FEM) буде використано для розв'язання задачі теплообміну, оскільки він дозволяє точно моделювати складні геометрії та нестационарні процеси. Перевагою FEM є висока точність та можливість врахування нелінійних властивостей матеріалів, а недоліком – значні обчислювальні витрати при великій кількості елементів;
- метод скінченних об'ємів (FVM) буде застосовано для моделювання гідродинаміки потоку мастила. FVM добре підходить для розв'язання рівнянь Нав'є-Стокса та моделювання турбулентних потоків. Перевагою є збереження маси та енергії на локальному рівні, що важливо для гідродинамічних розрахунків, а недоліком може бути складність у реалізації на нерегулярних сітках;
- адаптивні сітки – застосування адаптивних сіток дозволяє зосередити обчислювальні ресурси в областях з високими градієнтами параметрів (наприклад, поблизу стінок труб або в зонах турбулентності). Перевага полягає в зменшенні загальної кількості елементів при збереженні високої точності, а недоліком є складність алгоритмів сіткогенерації та необхідність їхньої оптимізації.

На основі розроблених математичних моделей створено програмний пакет, який реалізує комбінований підхід з використанням FEM та FVM. Використано мову програмування C++ для забезпечення високої продуктивності та гнучкості. Додатково інтегровано можливість паралельних обчислень з використанням OpenMP та MPI, що дозволяє значно скоротити час розрахунків.

Розроблений програмний пакет враховує:

- нелінійну в'язкість мастила, залежність від температури та тиску;
- турбулентність потоку з використанням моделей турбулентності;

- теплопровідність матеріалів труб та обмін теплом з оточуючим середовищем;
- адаптивні сітки, які дозволяють детально моделювати області з високими градієнтами параметрів.

Проведено серію чисельних експериментів з використанням розробленого програмного забезпечення. Результати моделювання порівняно з експериментальними даними, отриманими в лабораторних умовах та з реальних промислових систем.

Розроблений підхід, що поєднує переваги методів скінченних елементів та об'ємів з використанням адаптивних сіток, дозволив створити ефективне та точне програмне забезпечення для моделювання процесів теплообміну та гідродинаміки турбінного мастила. Порівняння з альтернативними методами показало, що власне програмне забезпечення має ряд переваг у специфічних застосуваннях, хоча і вимагає подальшого розвитку та вдосконалення.

Список використаних джерел

1. Свинчук О. В. & Клименко Я. В. (2024). Аналіз сучасних засобів моделювання процесів теплообміну та гідродинаміки перебігу рідин в трубах. Інфокомунікаційні та комп'ютерні технології, Т.1 (7), С. 38–54.
2. Smith J. (2020), Heat Transfer in Turbine Systems, Energy Press.

Дашик А.С., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Залевська О.В., к.т.н., Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

ОГЛЯД ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ ДЛЯ ТОПОЛОГІЧНОЇ ОПТИМІЗАЦІЇ

Сучасне інженерне проектування все частіше потребує застосування оптимізаційних методів для створення конструкцій, які поєднують високу міцність, легкість та ефективне використання матеріалів. Топологічна оптимізація є одним з підходів у цьому напрямі, вона дозволяє знаходити оптимальні розподіли матеріалу в межах заданих геометричних та фізичних обмежень. Застосування топологічної оптимізації особливо актуальне у таких галузях, як авіабудування, автомобільна промисловість, біомедична інженерія та машинобудування, де підвищення ефективності конструкцій має вирішальне значення.

Розвиток алгоритмів топологічної оптимізації сприяв появі численних програмних рішень, кожне з яких пропонує свої методи та підходи до оптимізації. Програмне забезпечення для топологічної оптимізації сьогодні охоплює широкий спектр методів — від градієнтних методів до еволюційних, що забезпечують гнучкість і точність проектних рішень. Таблиця 1 демонструє основні методи оптимізації, доступні в популярних програмних комплексах для топологічного проектування. Кожне програмне рішення підтримує різні підходи до оптимізації, що обумовлює їх вибір у залежності від специфічних інженерних потреб

Таблиця 1 – Доступні методи оптимізації

Метод оптимізації	ANSYS	SIMULIA Abaqus FEA	Altair OptiStruct	Solidworks Generative Design	Siemens NX	COMSOL Multiphysics
Оснований на густині, SIMP	+	+	+	+	+	+
Рівня граней, Level Set	+	+			+	
Еволюційний, BESO		+	+		+	
Оснований на градієнті	+		+			+
Вільного розміру			+			
Обмеження виробництва	+	+	+	+	+	

Таблиця 1 демонструє, що більшість програмних продуктів для топологічної оптимізації підтримують метод оснований на густині, що свідчить про його універсальність та ефективність для вирішення широкого спектра інженерних задач. Методи рівня граней та еволюційної оптимізації реалізовані у вибіркових програмних платформах, таких як SIMULIA Abaqus FEA, Altair OptiStruct та COMSOL Multiphysics, що робить їх актуальними для специфічних задач, які потребують детальнішої топології. Градієнтні методи присутні в ANSYS, Altair OptiStruct та COMSOL Multiphysics, забезпечуючи гнучкість у налаштуванні конфігурацій матеріалів. Лише Altair OptiStruct підтримує метод вільного розміру, що надає більше можливостей для оптимізації нестандартних геометрій. Практично всі розглянуті платформи дозволяють враховувати обмеження виробництва, що робить їх придатними для промислових застосувань, де необхідно враховувати технологічні аспекти виготовлення конструкцій.

З огляду на різноманіття програмних продуктів для топологічної оптимізації, важливим є оцінити, наскільки широко використовуються ці інструменти науковою спільнотою. Одним із показників популярності є кількість публікацій, у яких обговорюються певні платформи для топологічної оптимізації. Рисунок 1 демонструє кількість наукових публікацій, присвячених використанню різних програмних продуктів для топологічної оптимізації, у провідних наукових базах даних.

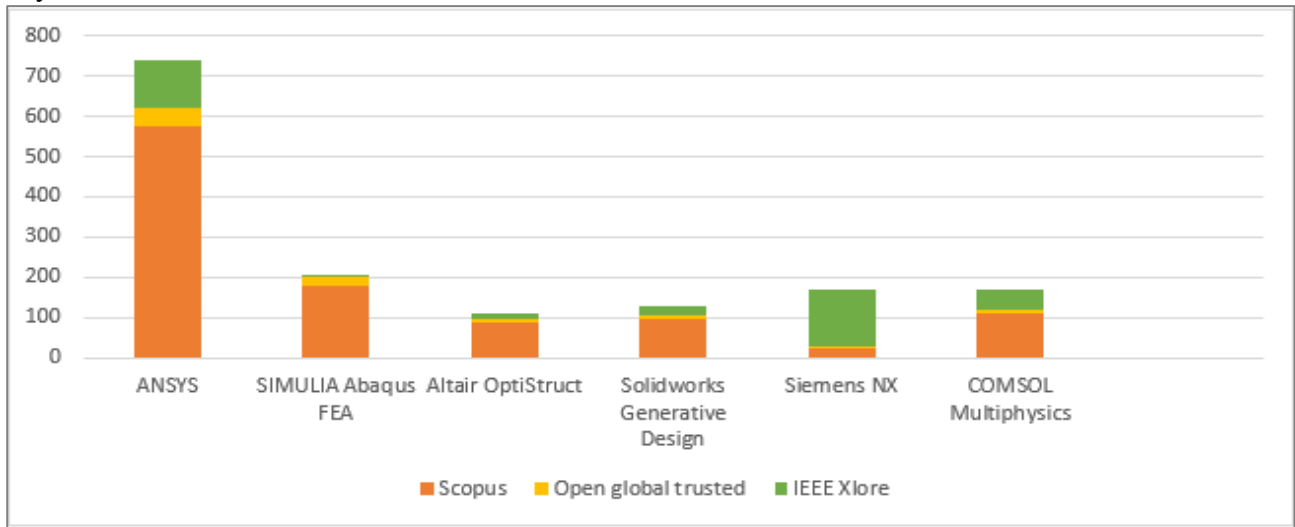


Рисунок 1 – Кількість публікацій в наукових базах даних

Результати огляду вказують на те, що ANSYS та SIMULIA Abaqus FEA є лідерами у сфері наукових досліджень, що підкреслює їхню універсальність і надійність для вирішення інженерних задач. Важливість таких платформ, як COMSOL Multiphysics, Altair OptiStruct і Solidworks Generative Design, також підтверджується їхньою активною участю в наукових публікаціях, що свідчить про їхнє широке застосування у спеціалізованих областях. Siemens NX, зі значною кількістю публікацій в індустріальних дослідженнях, демонструє свою привабливість для інженерів практиків, які працюють у галузі автоматизації проектування.

Таким чином, результати дослідження підтверджують важливість програмного забезпечення для топологічної оптимізації в сучасних інженерних науках, а також вказують на тенденції у наукових дослідженнях, пов'язаних із оптимізацією конструкцій.

Список використаних джерел

1. Alonso C. Topology Design Methods for Structural Optimization / Cristina Alonso, Osvaldo M. Querin, Mariano Victoria та ін.. – London, Great Britain: Academic Press Inc., 2017. – pp. 1–23.
2. Bendsøe M. Topology Optimization: Theory, Methods and Applications / M. Bendsøe, O. Sigmund. – Heidelberg, Germany: Springer Berlin, 2023. – p. 71–84.
3. Huei-Huang Lee. Finite Element Simulations with ANSYS Workbench 2023 / Huei-Huang Lee. – Mission, USA: SDC Publications, 2023. – pp. 4–23.

ПОКАЖЧИК ЗА АВТОРАМИ

Gong Xiaodong (Гонг Сяодун)	51, 53	Колумбет В.П.	37, 80
Kuzminyh V.O., (Кузьмініх В.О.)	56	Лебедєв А.В.	67
Shiwei Zhu (Чжу Шівей)	49	Логвиненко Т.С.	35
Taranenko E.R. (Тараненко Є.Р.)	58	Ляшко І.І.	72
Taranenko R.A. (Тараненко Р.А.)	58	Макарчук А.В.	23
Voitko A.S. (Войтко А.С.)	56	Максименко П.О.	63
Zhang Mingjun (Чжан Мінцзюнь)	55	Мусієнко А.П.	13, 47
Барабаш О.В.	11, 23, 41	Недашківський О.Л.	19, 31, 77
Бандурка О.І.	27, 45	Олексій А.О.	16
Безсмертний В.Ю.	47	Остапенко І.П.	29
Варава І.А.	6, 43, 65	Передера В.Р.	77
Верлань А.А.	16	Пироговська Т.В.	13
Ворвуль Д.М.	47	Романів Р.С.	45
Гаврилко Є.В.	60	Руй Д.І.	37
Гагарін О.О.	35	Савко В.Я.	60
Гейко О.О.	43	Садовська І.І.	75
Гнатишин М.С.	31	Свинчук О.В.	11, 83
Гусєва І.І.	63	Сенченко В.Р.	49
Дащик А.С.	85	Ситнік І.О.	33
Дембіцький В.В.	39	Статівка Ю.І.	55
Добровольський Р.В.	49	Федорова Н.В.	67
Донець А.Г.	37, 80	Фішер О.Є.	41
Євтушенко Д.М.	80	Хоменко О.М.	69
Єрмосін О.О.	51	Царева О.С.	25
Єрмосіна О.О.	53	Черноусов Д.І.	27
Залевська О.В.	72, 85	Шалденко О.В.	19
Здор К.А.	19	Шпуриєв В.В.	33
Клименко Я.В.	83	Шуклін Г.В.	8
Коваль О.В.	39, 49, 51, 53, 69, 75	Ярмак Д.О.	65



Україна 03057, м. Київ, Берестейський проспект, 37 Тел +38 (044) 204 80 82
37, Prospect Beresteiskyi, Kyiv, 03057, Ukraine Tel+38 (044) 204 80 82
www: <https://ipze.kpi.ua>, E-mail: ipzeconf.kpi.ua@lil.kpi.ua