

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
ДЕРЖАВНИЙ УНІВЕРСИТЕТ
ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ
НАВЧАЛЬНО-НАУКОВИЙ ІНСТИТУТ ІНФОРМАЦІЙНИХ ТЕХНОЛОГІЙ
КАФЕДРА КОМП'ЮТЕРНИХ НАУК

VII МІЖНАРОДНА НАУКОВО-ПРАКТИЧНА КОНФЕРЕНЦІЯ

«СУЧАСНІ ДОСЯГНЕННЯ КОМПАНІЇ HEWLETT PACKARD ENTERPRISE В
ГАЛУЗІ ІТ ТА НОВІ МОЖЛИВОСТІ ЇХ ВИВЧЕННЯ І ЗАСТОСУВАННЯ»

ЗБІРНИК ТЕЗ

11 грудня 2025

VII МІЖНАРОДНА НАУКОВО-ПРАКТИЧНА КОНФЕРЕНЦІЯ «Сучасні досягнення компанії HEWLETT PACKARD ENTERPRISE в галузі ІТ та нові можливості їх вивчення і застосування».

ЗБІРНИК ТЕЗ. – К.: ДУІКТ, 2025р.

Збірник містить тези доповідей учасників конференції, представлених на Міжнародній науково-практичній конференції «Сучасні досягнення компанії HEWLETT PACKARD ENTERPRISE в галузі ІТ та нові можливості їх вивчення і застосування», яка проводилась 11 грудня 2025р. на кафедрі Комп'ютерних наук Навчально-наукового інституту інформаційних технологій Державного університету інформаційно-комунікаційних технологій, м. Київ.

Робочі мови – українська та англійська.

На конференції проведено апробацію результатів наукових досліджень, обговорено перспективи та різноманітні підходи до вирішення сучасних проблем в галузі інформаційних технологій та досягнення компанії HEWLETT PACKARD ENTERPRISE.

ОГАНІЗАТОРИ КОНФЕРЕНЦІЇ

Державний університет інформаційно-комунікаційних технологій
Товариство з обмеженою відповідальністю «Sophela»,
представник компанії Hewlett Packard Enterprise в Україні
Навчально-науковий інститут інформаційних технологій
кафедра Комп'ютерних наук, кафедра Штучного інтелекту

ПРОГРАМНИЙ КОМІТЕТ

Нестеренко К.С. – к.т.н., доцент, директор Навчально-наукового інституту Інформаційних технологій, Державний університет інформаційно-комунікаційних технологій.

Піщіков А.В. – директор ТОВ «Sophela».

Вишнівський В.В. – д.т.н., проф., Державний університет інформаційно-комунікаційних технологій.

Зінченко О.В. – д.т.н., доцент, Державний університет інформаційно-комунікаційних технологій.

Сторчак К.П. – д.т.н., проф., Державний університет інформаційно-комунікаційних технологій.

Шукула О.М. – д.фіз.-мат.н., проф., Державний університет інформаційно-комунікаційних технологій.

Срібна І.М. д.т.н., доцент, Державний університет інформаційно-комунікаційних технологій.

Барабаш О.В. – д.т.н., проф., Національний технічний університет України «Київський політехнічний інститут ім. Ігоря Сікорського».

Мусієнко А.П. – д.т.н., доцент, Національний технічний університет України «Київський політехнічний інститут ім. Ігоря Сікорського».

Аль-Амморі А.Н. – д.т.н., проф., Національний транспортний університет.

Поперешняк С.В. – к.ф.-м.н., доцент, Національний технічний університет України «Київський політехнічний інститут ім. Ігоря Сікорського».

Щербина І.С. – к.т.н., доцент, Державний університет інформаційно-комунікаційних технологій.

Іщоряков С.М. – к.т.н., доцент, Державний університет інформаційно-комунікаційних технологій.

Березовська Ю.В. – доктор філософії, Державний університет інформаційно-комунікаційних технологій.

Шумигора Т.І. – старший комерційний представник ТОВ «Sophela».

ЗМІСТ

1	Багажков Д.С. , Катков Ю.І. ВЕБ-СКРЕЙПІНГ У СИСТЕМІ КОЛЕКЦІОНУВАННЯ КОНТЕНТУ НА ГІБРИДНИХ РІШЕННЯХ НРЕ.....	9
2	Бондар В.О., Катков Ю.І. ЗАСТОСУВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ ОПТИМІЗАЦІЇ ЯКОСТІ МЕДІАСТРИМІНГУ В РЕАЛЬНОМУ ЧАСІ.....	11
3	Загороднюк Я.А., Катков Ю.І. МОДЕЛЬ ВЕРИФІКАЦІЇ SMART HOME ІОТ ЧЕРЕЗ ЦИФРОВІ ДВІЙНИКИ З ВИКОРИСТАННЯМ РІШЕНЬ НРЕ.....	13
4	Нечипорук А.В., Катков Ю.І. ПЯТИВИМІРНА МОДЕЛЬ ПЕРСОНАЛІЗАЦІЇ НАВЧАННЯ АНГЛІЙСЬКОЇ МОВИ НА ОСНОВІ АДАПТИВНИХ АЛГОРИТМІВ.....	15
5	Сагайдак В.А. ARTIFICIAL INTELLIGENCE FOR WIRELESS LOCAL NETWORK PERFORMANCE ANALYSIS.....	18
6	Valuiskyi H.O., Chorna V.V. THE ROLE OF THE JAVA VIRTUAL MACHINE IN ACHIEVING CROSS-PLATFORM COMPATIBILITY.....	20
7	Маслюков І.С., Чорна В.В. ХМАРНІ ТЕХНОЛОГІЇ ЯК КЛЮЧОВИЙ ІНСТРУМЕНТ ЦИФРОВОЇ ТРАНСФОРМАЦІЇ СУЧАСНИХ ОРГАНІЗАЦІЙ.....	21
8	Господинько М.А., Чорна В.В. ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ ТА ШТУЧНИЙ ІНТЕЛЕКТ У СУЧАСНІЙ ПРОМИСЛОВІЙ ІНДУСТРІЇ.....	23
9	Калашникова К.С., Чорна В.В. ІНФОРМАЦІЙНО-ПОШУКОВІ ТА ЕКСПЕРТНІ СИСТЕМИ: ІННОВАЦІЙНІ ПІДХОДИ В ОСВІТНІЙ СФЕРІ ВІД ІТ-КОМПАНІЙ.....	25
10	Нескреба М.С., Чорна В.В. ІНТЕЛЕКТУАЛЬНІ АДАПТИВНІ СИСТЕМИ ПРОГНОЗУВАННЯ ПОВЕДІНКИ КОРИСТУВАЧІВ У ЦИФРОВИХ ПЛАТФОРМАХ НА ОСНОВІ МУЛЬТИМОДАЛЬНИХ AI-МОДЕЛЕЙ.....	27
11	Смірнов Д.Є., Чорна В.В. ШТУЧНИЙ ІНТЕЛЕКТ В ІНДУСТРІЇ: ВІД АВТОМАТИЗАЦІЇ ПРОЦЕСІВ ДО ЦИФРОВОЇ ТРАНСФОРМАЦІЇ ПІДПРИЄМСТ.....	29
12	Сліпко Д.Р., Чорна В.В. ВИКОРИСТАННЯ ІНФРАСТРУКТУРНИХ ПЛАТФОРМ HEWLETT PACKARD ENTERPRISE ДЛЯ ПІДВИЩЕННЯ ЕФЕКТИВНОСТІ ІТ-СИСТЕМ.....	31
13	Куліков Є.О., Чорна В.В. РОЗРОБКА ТА ВПРОВАДЖЕННЯ ВЕБ-СИСТЕМИ РЕКОМЕНДАЦІЙ ФІЛЬМІВ MOVIEVIBE.....	32
14	Стойко М.М. ІІІ-МЕТОДИ ІНТЕГРАЦІЇ ГЕТЕРОГЕННИХ МЕРЕЖ У СПЕЦІАЛЬНИХ АСУ СП.....	33
15	Гніденко М.М. ВПЛИВ РОЗМІЩЕННЯ КОНТРОЛЕРІВ SDN НА ЕФЕКТИВНІСТЬ МЕРЕЖІ.....	35

16	Худік Б.О., Гашійук А.О. РОЗРОБКА МЕТОДИКИ ВИЯВЛЕННЯ ФІШИНГОВИХ ЕЛЕКТРОННИХ ЛИСТІВ ЗА ДОПОМОГОЮ АЛГОРИТМІВ МАШИННОГО НАВЧАННЯ.....	37
17	Радковський Д.О., Чичкарьов Є.А. СИСТЕМА МОНІТОРИНГУ ПСИХОЕМОЦІЙНОГО СТАНУ ЛЮДИНИ З ВИКОРИСТАННЯМ ШТУЧНОГО ІНТЕЛЕКТУ.....	39
18	Вишнівський В.В., Мороз М.В. ПОРІВНЯЛЬНИЙ АНАЛІЗ МОДЕЛЕЙ U-NET, DPT ТА MIDAS ДЛЯ ОЦІНКИ ГЛИБИНИ У SLAM-СИСТЕМАХ.....	41
19	Вишнівський В.В., Мороз М.В. ПОЄДНАННЯ ГЛИБИННИХ МАП І СЕМАНТИЧНОГО РОЗПІЗНАВАННЯ ДЛЯ ПІДВИЩЕННЯ ТОЧНОСТІ ЛОКАЛІЗАЦІЇ У ЗАДАЧАХ SLAM	43
20	Вишнівський В.В., Мороз М.В. ВИКОРИСТАННЯ НЕЙРОМЕРЕЖЕВИХ ОЗНАКОВИХ ДЕСКРИПТОРІВ У ВІЗУАЛЬНИХ ЗАДАЧАХ SLAM.....	45
21	Чичур А.І., Подгорбунський В.В. ОЦІНКА ЛІНГВІСТИЧНОЇ НОВИЗНИ ТЕКСТІВ, ЗГЕНЕРОВАНИХ ВЕЛИКИМИ МОВНИМИ МОДЕЛЯМИ	46
22	Чичур А.І., Подгорбунський В.В. АНАЛІЗ СИСТЕМ КОНТРОЛЮ ТА МОНІТОРИНГУ LLM.....	49
23	Чичур А.І., Метелюк А.В. ПОРІВНЯННЯ МЕТОДІВ ПЕРЕХОДУ ВІД МОНОЛІТНОЇ ДО МІКРОСЕРВІСНОЇ АРХІТЕКТУРИ ДЛЯ СИСТЕМИ МОНІТОРИНГУ ТА АЛЕРТИНГУ РЕКЛАМНИХ КАМПАНІЙ.....	51
24	Чичур А.І., Брагінець Т.В. ПЕРЕВАГИ ТА НЕДОЛІКИ ВПРОВАДЖЕННЯ ТЕХНОЛОГІЙ ДОПОВНЕНОЇ РЕАЛЬНОСТІ В ТУРИСТИЧНІ ГАЛУЗІ.....	54
25	Товсточуб І.С., Ярохно Д.В. ЕВОЛЮЦІЯ ПЛАСКОГО ТА СКЕВОМОРФНОГО ДИЗАЙНІВ В ІГРОВИХ UI.....	55
26	Товсточуб І.С., Ярохно Д.В. МЕТРИКИ ЕФЕКТИВНОСТІ ВІЗУАЛЬНИХ МАРКЕРІВ НА ПРИКЛАДІ ЖОВТОЇ ПІДКРАСКИ.....	58
27	Товсточуб І.С., Колесник В.Я. ІНТЕЛЕКТУАЛЬНА СИСТЕМА РАНЬОГО ВИЯВЛЕННЯ АНОМАЛІЙ ДОСТУПУ ДО КОРПОРАТИВНОГО VPN.....	60
28	Товсточуб І.С., Колесник В.Я. ІНТЕЛЕКТУАЛЬНА СИСТЕМА РАНЬОГО ВИЯВЛЕННЯ АНОМАЛІЙ ДОСТУПУ ДО КОРПОРАТИВНОГО VPN.....	62
29	Тестов Т., Чичкарьов Є., Вишневський В., Зінченко О. ВИДІЛЕННЯ КЛЮЧОВИХ СЛІВ З БАГАТОМОВНОГО ТЕКСТУ ПРИ ОРГАНІЗАЦІЇ KEYWORD DRIVEN TESTING	64
30	Гніденко М.П., Шепетун О.О. ВИКОРИСТАННЯ МОДУЛЬНИХ НЕЙРОМЕРЕЖЕВИХ МОДЕЛЕЙ ДЛЯ ПОКРАЩЕННЯ ЯКОСТІ ІДЕНТИФІКАЦІЇ ЛЮДИНИ НА ФОТОЗОБРАЖЕННЯХ.....	67
31	Іщераков С.М., Дубовицький Д.С. ВПЛИВ СЕРВЕРНОГО РЕНДЕРУ НА ПРОДУКТИВНІСТЬ ВЕБ-СТОРИНОК.....	68

32	Дубовицький Д. С. Емуляція браузерного середовища для виконання скриптів JavaScript на сервері: методи та виклики.....	70
33	Шепетун О.О., Гніденко М.П. ВИКОРИСТАННЯ МОДУЛЬНИХ НЕЙРОМЕРЕЖЕВИХ МОДЕЛЕЙ ДЛЯ ПОКРАЩЕННЯ ЯКОСТІ ІДЕНТИФІКАЦІЇ ЛЮДИНИ НА ФОТОЗОБРАЖЕННЯХ	71
34	Шикула О.М., Целованський Т.Р. КОМПЛЕКСНІ МЕТОДИ ОПТИМІЗАЦІЇ ПІДСИСТЕМИ ПАМ'ЯТІ НА СЕРВЕРНИХ БАГАТОСОКЕТНИХ СИСТЕМАХ.....	72
35	Миронов Д.В., Верещак А.О. ВПРОВАДЖЕННЯ АРХІТЕКТУРИ НУЛЬОВОЇ ДОВІРИ (ZERO TRUST) ДЛЯ ПІДВИЩЕННЯ БЕЗПЕКИ СУЧАСНИХ КОМП'ЮТЕРНИХ МЕРЕЖ.....	74
36	Звенигородський О.С., Шаш М.С. МОДЕЛЬ АНАЛІЗУ ЮНІТ-ЕКОНОМІКИ SAAS ЗА ДОПОМОГОЮ ШТУЧНОГО ІНТЕЛЕКТУ.....	76
37	Серих С.О., Попов А.О. ВИКОРИСТАННЯ CROSS-LINGUAL MULTI-SPEAKER TTS ТА VOICE CONVERSION ДЛЯ АУГМЕНТАЦІЇ ДАНИХ У ЗАДАЧАХ РОЗПІЗНАВАННЯ МОВИ.....	80
38	Серих С.О., Попов А.О. ПОПЕРЕДНЄ НАВЧАННЯ БЕЗ СПОСТЕРЕЖЕННЯ ДЛЯ ЕФЕКТИВНОГО СИНТЕЗУ МОВЛЕННЯ В УМОВАХ НИЗЬКОРЕСУРСНИХ МОВ.....	81
39	Шикула О.М., Коник С.П. ТЕЛЕГРАМ ЧАТ-БОТ НА МОВІ PYTHON З ІНТЕГРАЦІЄЮ В SRM СИСТЕМУ.....	83
40	Шикула О.М., Бугайченко А.Д. ІНФОРМАЦІЙНА СИСТЕМА ПРИЙОМУ ТА АНАЛІЗУ ЗАЯВОК ТЕХНІЧНОЇ ПІДТРИМКИ (Python/Django + Bootstrap + PostgreSQL).....	85
41	Шикула О.М., Івасюк А.І. ВЕБ-ЗАСТОСУНОК ДЛЯ МОНІТОРИНГУ ПОГОДНИХ ТА ЕКОЛОГІЧНИХ ПОКАЗНИКІВ НА МОВІ JAVASCRIPT З ВИКОРИСТАННЯМ ФРЕЙМВОРКУ VUE.JS.....	86
42	Шикула О.М., Борода К.О. ЧАТ-БОТ-МАГАЗИНУ В TELEGRAM НА МОВІ PYTHON ІЗ ВИКОРИСТАННЯМ СКБД MYSQL.....	88
43	Шикула О.М., Кирилов Д.В. РОЗРОБКА САЙТУ КІННОГО КЛУБУ НА МОВАХ ПРОГРАМУВАННЯ JAVASCRIPT ТА PHP З ВИКОРИСТАННЯМ СКБД MYSQL.....	90
44	Шикула О.М., Лоскучерява С.С. РОЗРОБКА САЙТУ ФІРМИ З ОРГАНІЗАЦІЇ СВЯТКОВИХ ЗАХОДІВ ДЛЯ ДІТЕЙ НА МОВАХ ПРОГРАМУВАННЯ JAVASCRIPT ТА PHP З ВИКОРИСТАННЯМ СКБД MYSQL.....	92
45	Шикула О.М., Марохонько М.О. МЕРЕЖЕВЕ СХОВИЩЕ NAS (NETWORK ATTACHED STORAGE) З ВИКОРИСТАННЯМ REACT.JS, NODE.JS, MONGODB.....	94
46	Шикула О.М., Мазуренко А.В. ВЕБ-ЗАСТОСУНОК ДЛЯ СТВОРЕННЯ РЕЗЮМЕ НА МОВІ ПРОГРАМУВАННЯ JAVASCRIPT З ВИКОРИСТАННЯМ ФРЕЙМВОРКУ VUE.JS.....	95
47	Шикула О.М., Марохонько М.О. ДОСЛІДЖЕННЯ МЕТОДІВ ШИФРУВАННЯ ПОВІДОМЛЕНЬ ТА ПРОБЛЕМИ ЗАХИСТУ ДАНИХ	97

	У СУЧАСНИХ ЧАТ СЕРВІСАХ.....	
48	Гніденко М.П., Сітко Д.О. ІНТЕЛЕКТУАЛЬНІ МЕТОДИ ПОБУДОВИ ВЕКТОРНИХ КОНТУРІВ НА ОСНОВІ ТОПОЛОГІЧНОЇ ІНТЕРПРЕТАЦІЇ ПІКСЕЛЬНИХ СТРУКТУР.....	98
49	Кисіль Т.М., Кривий А.В. AIOPS: МЕТОДОЛОГІЯ ЗАСТОСУВАННЯ МАШИННОГО НАВЧАННЯ ДЛЯ КОРЕЛЯЦІЇ ПОДІЙ ТА ОПТИМІЗАЦІЇ МОНІТОРИНГУ.....	101
50	Кисіль Т.М., Марченко Б.С., Сущенко В.В. АРХІТЕКТУРНИЙ ПОРІВНЯЛЬНИЙ АНАЛІЗ АСИНХРОННИХ МОДЕЛЕЙ ОБРОБКИ ПОДІЙ ПРИ ПРОЕКТУВАННІ ЧАТ-БОТІВ.....	103
51	Кисіль Т.М., Товстенко М.В. MODEL CONTEXT PROTOCOL (MCP): АРХІТЕКТУРА ТА ІНТЕГРАЦІЯ LLM З КОРПОРАТИВНИМИ БАЗАМИ ДАНИХ.....	106
52	Худік Б.О., Гайшук А.О. РОЗРОБКА МЕТОДИКИ ВИЯВЛЕННЯ ФІШИНГОВИХ ЕЛЕКТРОННИХ ЛИСТІВ ЗА ДОПОМОГОЮ АЛГОРИТМІВ МАШИННОГО НАВЧАННЯ.....	108
53	Кисіль Т.М. CNN-LSTM ПІДХІД ДО РОЗПІЗНАВАННЯ ЕМОЦІЙНИХ СТАНІВ ЗА ПОЗОЮ ОСІБ В РЕЖИМІ РЕАЛЬНОГО ЧАСУ.....	110
54	Кисіль Т.М., Шимко В.О. АДАПТИВНЕ ПРОГНОЗУВАННЯ ЕНЕРГОСПОЖИВАННЯ МІСЬКИХ МЕРЕЖ НА ОСНОВІ ГІБРИДНИХ МОДЕЛЕЙ ГЛИБИННОГО НАВЧАННЯ.....	113
55	Стежко С.О., Кондратенко Н.Ю. МОВНА ОСВІТА В СИСТЕМІ ПІДГОТОВКИ МАЙБУТНІХ ПЕДАГОГІВ: ІНТЕГРАЦІЙНИЙ ПІДХІД.	116

Катков Ю.І., д.т.н., професор ДУІКТ
Багажков Д.С., студент ДУІКТ

ВЕБ-СКРЕЙПІНГ У СИСТЕМІ КОЛЕКЦІОНУВАННЯ КОНТЕНТУ НА ГІБРИДНИХ РІШЕННЯХ HPE

Постановка завдання. Проблема полягає в тому, що більшість систем каталогізації контенту, наприклад, інформації про мангу, не мають автоматизованого оновлення даних. Завданням дослідження є розроблення веб-системи, здатної збирати та структурувати інформацію з відкритих джерел, а також інтегруватися з гібридними інфраструктурними рішеннями Hewlett Packard Enterprise (HPE) для надійного зберігання даних і масштабування. Необхідно проаналізувати методи веб-скрейпінгу, визначити вимоги до системи, спроектувати архітектуру, реалізувати модулі збору та відображення даних, протестувати працездатність платформи.

Мета роботи. Метою дослідження є створення веб-системи автоматизованого збору та каталогізації інформації про мангу з використанням веб-скрейпінгу та підтримкою розгортання в гібридних обчислювальних середовищах HPE для надійної обробки й зберігання даних

Результати дослідження. Стрімке зростання веб-даних потребує інструментів автоматизованого збору та структурування інформації. Гібридні обчислювальні рішення HPE забезпечують масштабовану інфраструктуру для таких задач, поєднуючи локальні ресурси та хмарні сервіси [1-4]. Використання веб-скрейпінгу дозволяє створювати системи з актуальними каталогами контенту для колекціонування манги без ручного оновлення.

Було проведено аналіз сучасних методів веб-скрейпінгу та наявних систем для каталогізації контенту про мангу. Здійснено проєктування веб-системи, архітектура якої орієнтована на можливість розгортання в гібридних ІТ-середовищах HPE. Гібридне ІТ-середовище HPE — це інфраструктура, що поєднує локальні ресурси, хмарні сервіси та edge-обчислення в єдину керовану систему. Воно створюється завдяки технологіям HPE GreenLake (хмара як сервіс), HPE ProLiant (сервери), HPE Alletra Storage (масиви зберігання), HPE Aruba ESP (мережа й edge), а також HPE Ezmeral (керування даними та аналітика). Таке середовище забезпечує масштабованість, високу доступність і гнучке керування ресурсами.

Створено модуль веб-скрейпінгу, що виконує автоматичне отримання та аналіз HTML-документів, вилучає релевантну інформацію та зберігає її у реляційній базі даних. Асинхронні механізми дозволили оптимізувати швидкість обробки запитів. Логування та обробка помилок забезпечують стабільну роботу навіть при зміні структури джерел.

База даних спроектована з урахуванням можливості перенесення або масштабування на високопродуктивні масиви зберігання даних HPE, що дозволяє обробляти збільшені обсяги контенту та забезпечує надійність і швидкодію. Серверна частина реалізована через REST API та готова до розгортання як у локальному середовищі, так і в гібридних рішеннях HPE GreenLake чи HPE Alletra Storage.

Розроблено інтерфейс користувача з адаптивним дизайном, що дозволяє формувати особисті колекції, здійснювати пошук, сортування та перегляд деталізованої інформації. Проведені модульні, інтеграційні та навантажувальні тести показали стабільність системи, її працездатність і сумісність з ІТ-інфраструктурою HPE.

Висновки та перспективи подальших досліджень

Створено веб-систему автоматизованого збору та обробки контенту з використанням веб-скрейпінгу. Інтеграція з гібридними обчислювальними рішеннями HPE підвищує масштабованість і надійність системи, забезпечує ефективне зберігання даних і дозволяє легко розширювати її функціонал.

Перспективи подальших досліджень. Подальший розвиток передбачає впровадження рекомендаційних моделей, розширену аналітику, підтримку додаткових джерел даних та повну інтеграцію з платформами HPE GreenLake для керування навантаженнями, резервного копіювання й оптимізації витрат на зберігання.

Список використаних джерел

1. Hewlett Packard Enterprise. *HPE GreenLake Cloud Platform: Technical Overview*, 2025. URL: <https://www.hpe.com/us/en/greenlake.html> (date of access: 24.11.2025).
2. PostgreSQL Global Development Group. *PostgreSQL 16 Documentation*, 2025. URL: <https://www.postgresql.org/docs/> (date of access: 24.11.2025).
3. HPE Documentation. *HPE Alletra Storage — Unified Data Infrastructure*, 2024. URL: <https://www.hpe.com/us/en/storage/alletra.html> (date of access: 24.11.2025).
4. Mitchell R. *Web Scraping with Python: Collecting Data from the Modern Web*. O'Reilly, 2023. URL: <https://www.oreilly.com/library/view/web-scraping-with/9781491985564> (date of access: 24.11.2025).

Катков Ю.І., д.т.н., професор ДУІКТ
Бондар В.О., студент ДУІКТ

ЗАСТОСУВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ ОПТИМІЗАЦІЇ ЯКОСТІ МЕДІАСТРИМІНГУ В РЕАЛЬНОМУ ЧАСІ

Постановка задачі. Сучасні медіаплатформи стикаються з постійно зростаючим попитом на відеостримінг високої якості в умовах нестабільних мережевих ресурсів. Користувачі очікують плавного відтворення, мінімальних затримок та відсутності буферизації, що висуває високі вимоги до механізмів управління передачею даних [1,2]. Традиційні алгоритми адаптивного бітрейту часто не здатні оперативно реагувати на різкі коливання пропускну здатності мережі [3,4]. Тому основною проблемою сучасних медіаплатформ є забезпечення максимальної якості стримінгу за умов динамічних змін мережевих параметрів, зокрема пропускну здатності, затримок та втрат пакетів. Для вирішення проблеми потрібно розробити механізм, який зможе передбачати зміни мережі та адаптивно регулювати бітрейт у реальному часі. Традиційні моделі адаптивного бітрейту (ABR) працюють за заздалегідь визначеними правилами, що призводить до необ'єктивних рішень і погіршення досвіду користувача. У цьому контексті застосування методів ШІ для створення інтелектуальної системи адаптивного бітрейту набуває особливої актуальності, даючи можливість значно покращити стабільність, якість та ефективність медіастримінгу в реальному часі.

Мета роботи. Метою роботи є дослідження та обґрунтування використання методів штучного інтелекту (ШІ) для створення інтелектуальної системи адаптивного бітрейту, здатної динамічно оптимізувати якість відеостримінгу. Передбачається розроблення моделі, яка здатна прогнозувати стан мережі, обирати оптимальний рівень бітрейту та забезпечувати стабільну якість перегляду незалежно від зовнішніх умов.

Результати дослідження. Адаптивний бітрейт (ABR) на основі ШІ — це технологія, що використовує алгоритми машинного навчання для динамічного вибору оптимального бітрейту під час медіастримінгу [5,6]. На відміну від класичних ABR-рішень, які орієнтуються переважно на миттєві параметри мережі, інтелектуальні системи здатні аналізувати історичні дані, моделювати майбутні стани каналу та прогнозувати потенційні коливання пропускну здатності. Створення такої системи включає збір телеметрії мережі, формування наборів даних, тренування моделей (наприклад, нейронних мереж, моделей підкріпленого навчання), а також інтеграцію їх у систему доставки контенту.

Інтелектуальна система адаптивного бітрейту здатна враховувати широкий спектр параметрів: буфер клієнта, ймовірність втрат пакетів, затримки, поточний бітрейт, історію поведінки користувача та специфіку пристрою. Моделі підкріпленого навчання навчаються на основі винагороди, що максимізує якість досвіду користувача, мінімізуючи буферизацію й різкі стрибки якості. Це дозволяє системі приймати значно зваженіші рішення у порівнянні з традиційними евристичними підходами. Використання ШІ у ABR суттєво зменшує кількість переключень між якісними рівнями відео, забезпечує більш плавний стримінг та

покращує загальне QoE (Quality of Experience). [6,7]. Такі моделі також здатні адаптуватися до нових умов без необхідності повного перепроєктування алгоритмів, що робить їх ефективними у масштабованих стримінгових сервісах та CDN-інфраструктурах. CDN (Content Delivery Network) — це географічно розподілена мережа серверів, що зберігає копії веб-контенту, серед яких зображення, CSS, JS та файли HTML. Коли користувач заходить на сайт, що використовує CDN, мережа доставляє контент із найближчого сервера, що значно покращує швидкість завантаження та продуктивність.

Висновки та перспективи. Застосування штучного інтелекту в системах адаптивного бітрейту демонструє значне підвищення якості медіастримінгу в умовах реального часу. Інтелектуальні ABR-моделі забезпечують гнучке реагування на мережеві коливання, покращують стабільність та знижують ризик буферизації. Перспективи подальших досліджень пов'язані з удосконаленням моделей прогнозування, інтеграцією мультиагентних систем, застосуванням генеративних методів для симуляції мережевих сценаріїв та розвитком гібридних ABR-підходів для високонавантажених стримінгових сервісів.

Список використаних джерел.

1. Zijiang Yan, Jianhua Pei, Hongda Wu, Hina Tabassum, Ping Wang. (2025). *Semantic-Aware Adaptive Video Streaming Using Latent Diffusion Models for Wireless Networks*. 2025. [електронний ресурс] — режим доступу [arXiv https://arxiv.org/abs/2502.05695?utm_source=chatgpt.com](https://arxiv.org/abs/2502.05695?utm_source=chatgpt.com) (date of access: 24.11.2025).
2. Ziyu Zhong, Mufan Liu, Le Yang, Yifan Wang, Yiling Xu, Jenq-Neng Hwang. (2025). *Video Streaming with Kairos: An MPC-Based ABR with Streaming-Aware Throughput Prediction*. 2025. [електронний ресурс] — режим доступу [arXiv https://arxiv.org/abs/2503.14271?utm_source=chatgpt.com](https://arxiv.org/abs/2503.14271?utm_source=chatgpt.com) (date of access: 24.11.2025).
3. Bekir O. Turkkan, Ting Dai, Adithya Raman та інші. (2022). «GreenABR: Energy-Aware Adaptive Bitrate Streaming with Deep Reinforcement Learning». MMSys 2022 https://research.ibm.com/publications/greenabr-energy-aware-adaptive-bitrate-streaming-with-deep-reinforcement-learning?utm_source=chatgpt.com (date of access: 24.11.2025).
4. Waqas ur Rahman та інші. (2025). *Evaluating HAS and Low-Latency Streaming Algorithms for Enhanced QoE*. Electronics, 2025. [електронний ресурс] — режим доступу [MDPI+1](#)(date of access: 24.11.2025).
5. Adithya Raman, Bekir Turkkan, Tevfik Kosar. (2024). «LL-GABR: Energy Efficient Live Video Streaming Using Reinforcement Learning». arXiv, 2024. [електронний ресурс] — режим доступу: https://arxiv.org/abs/2402.09392?utm_source=chatgpt.com (date of access: 24.11.2025).
6. Jesús Aguilar-Armijo, Ekrem Çetinkaya, Christian Timmerer, Hermann Hellwagner. (2022). «ECAS-ML: Edge Computing Assisted Adaptation Scheme with Machine Learning for HTTP Adaptive Streaming». arXiv, 2022. [електронний ресурс] — режим доступу: [arXiv https://arxiv.org/abs/2201.04488?utm_source=chatgpt.com](https://arxiv.org/abs/2201.04488?utm_source=chatgpt.com) (date of access: 24.11.2025).

7. J.W. & D. A. (інші автори). (2022). «Improving the Quality of Experience of Video Streaming Through a Buffer-Based Adaptive Bitrate Algorithm and Gated Recurrent Unit-Based Network Bandwidth Prediction». Applied Sciences, 2022. [електронний ресурс] — режим доступу [MDPI https://www.mdpi.com/2076-3417/14/22/10490?utm_source=chatgpt.com](https://www.mdpi.com/2076-3417/14/22/10490?utm_source=chatgpt.com) (date of access: 24.11.2025).

Катков Ю.І., д.т.н., професор ДУІКТ
Загороднюк Я.А., студентка ДУІКТ

МОДЕЛЬ ВЕРИФІКАЦІЇ SMART HOME ІОТ ЧЕРЕЗ ЦИФРОВІ ДВІЙНИКИ З ВИКОРИСТАННЯМ РІШЕНЬ HPE

Постановка завдання. Проблема полягає у відсутності систематичного механізму контролю змін у складних Smart Home-системах, де взаємодіють десятки IoT-сенсорів, актуаторів та модулів на основі штучного інтелекту (ШІ-модулів). Необхідно створити рішення, яке дозволяє тестувати кожне оновлення у безпечному середовищі без ризику для реальної інфраструктури. Завдання полягає у побудові моделі цифрового двійника, здатного проводити серію симуляцій та стрес-тестів в гібридному середовищі HPE.

Мета роботи. Метою є створення моделі регулярного тестування та верифікації змін у Smart Home за допомогою цифрового двійника, що функціонує в гібридному середовищі HPE та забезпечує точне симулювання поведінки пристроїв, аналіз ризиків та ухвалення рішень щодо безпечного впровадження оновлень.

Результати дослідження. Стрімке зростання кількості IoT-пристроїв у Smart Home створює потребу у безпечному та прогнозованому оновленні їхньої логіки. Навіть дрібні зміни можуть спричинити відмови або загрозу безпеці. Для усунення цього використовується цифровий двійник — це віртуальна модель Smart Home, що точно відтворює роботу IoT-пристроїв, сценаріїв автоматизації та реакцій системи, дозволяючи безпечно тестувати зміни та прогнозувати наслідки без впливу на реальне середовище. Використання цифрових двійників у поєднанні з високопродуктивним гібридним середовищем Hewlett Packard Enterprise (HPE) забезпечує точне моделювання та верифікацію таких змін до їх розгортання. Гібридне середовище HPE — це поєднання локальних обчислень, хмарних сервісів і edge-інфраструктури, об'єднаних у єдину платформу HPE, наприклад, ProLiant, HPE GreenLake, HPE Ezmeral. Воно забезпечує масштабованість, керованість, безпечне зберігання та обробку IoT-даних, підтримує моделі «as-a-service» для ефективної роботи цифрових двійників та забезпечують швидку обробку сценаріїв, масштабування та надійне зберігання великих потоків телеметрії.

Запропоновано інтелектуальну модель цифрового двійника Smart Home, яка включає симуляцію фізичного середовища, поведінки сенсорів, сценаріїв автоматизації та реакцій користувача. Для обробки симуляцій застосовано

інфраструктурні рішення HPE: HPE ProLiant для виконання складних симуляцій в режимі високого навантаження; HPE Ezmeral для аналітики великих даних та оркестрації ІІІ-моделей; HPE GreenLake для гібридного керування ІоТ-даними та масштабованої обчислювальної моделі «as-a-service».

Механізм симуляції виконує багатоступеневий аналіз:

1. Попередня перевірка конфігурацій — визначення несумісностей між новими правилами та поточними параметрами системи;
2. Запуск сценарних симуляцій — моделювання добових циклів роботи, пікових навантажень, відмов сенсорів, поведінкових аномалій;
3. Стрес-тестування — оцінка реакцій системи при зміні зовнішніх умов;
4. Аналітична обробка — за допомогою інструментів HPE Ezmeral формується аналіз конфліктів, прогноз споживання енергії, визначення ризиків;
5. Рекомендації — система генерує висновки щодо доцільності розгортання оновлення.

Гібридне середовище HPE дозволяє забезпечити високу точність моделювання, швидке масштабування симуляцій та безпечне зберігання великих масивів сенсорних даних. Запропонована модель підвищує надійність Smart Home, зменшує ризик помилок, оптимізує енергоспоживання та забезпечує безпечну еволюцію інтелектуальних систем керування.

Висновки. Розроблена модель цифрового двійника підтвердила ефективність регулярного тестування оновлень Smart Home. Інтеграція ІІІ, ІоТ та рішень HPE забезпечує передбачуваність роботи, підвищує безпеку системи та зменшує ризик збоїв при впровадженні нових конфігурацій.

Перспективи подальших досліджень. Подальший розвиток моделі передбачає впровадження самонавчальних цифрових двійників, удосконалення прогностичної аналітики, розширення сценарного моделювання та інтеграцію з хмарно-краєвими платформами HPE для створення автономних адаптивних Smart Home-систем нового покоління.

Список використаних джерел

1. Farsi M., Daneshfar F., Marquez B. Digital Twin–Driven Predictive Management in Smart Home IoT Systems. *IEEE Internet of Things Journal*, 2025. DOI: <https://doi.org/10.1109/JIOT.2024.3511023> (date of access: 24.11.2025).
2. Rabah H., Karim M., Bouabdallah A. AI-Enhanced Digital Twin Architectures for Autonomous IoT Environments. *Future Generation Computer Systems*, 2025. DOI: <https://doi.org/10.1016/j.future.2024.106789> (date of access: 24.11.2025).
3. Zhao Y., Chen L., Wang Q. Simulation-Driven Validation of Cyber-Physical Smart Home Systems Using Digital Twins. *Simulation Modelling Practice and Theory*, 2025. DOI: <https://doi.org/10.1016/j.simpat.2024.102899> (date of access: 24.11.2025).
4. Kritzinger W., Sihn W. Next-Generation Digital Twin Methodologies for Cyber-Physical Environments. *Journal of Manufacturing Systems*, 2025. DOI: <https://doi.org/10.1016/j.jmsy.2024.101561> (date of access: 24.11.2025).

5. Gupta A., Singh R., Mohammed F. AI-Driven Adaptive Control Models for Smart Home Automation with Digital Twins. *IEEE Access*, 2025. DOI: <https://doi.org/10.1109/ACCESS.2025.1234567> (date of access: 24.11.2025).
6. Hewlett Packard Enterprise. *HPE GreenLake Edge-to-Cloud Platform Technical Overview*, 2025. URL: <https://www.hpe.com/greenlake> (date of access: 24.11.2025).
7. Hewlett Packard Enterprise. *HPE Ezmeral Data Fabric and AI Platform Architecture Guide*, 2025. URL: <https://www.hpe.com/ezmeral> (date of access: 24.11.2025).

Катков Ю.І., д.т.н., професор ДУІКТ
 Нечипорук А.В., студент ДУІКТ

П'ЯТИВИМІРНА МОДЕЛЬ ПЕРСОНАЛІЗАЦІЇ НАВЧАННЯ АНГЛІЙСЬКОЇ МОВИ НА ОСНОВІ АДАПТИВНИХ АЛГОРИТМІВ

Постановка задачі. Попри значний прогрес у сфері освітніх технологій, наявні платформи не формують комплексної моделі студента, що включає когнітивні, мотиваційні, поведінкові, стильові та контекстні параметри. Це обмежує якість прогнозування навчальних результатів та знижує ефективність персоналізації. Відсутність гнучкої архітектури та масштабованих інфраструктур також ускладнює обробку великої кількості користувачів і реального часу адаптації траєкторій. Отже, постає завдання створити багатовимірну адаптивну систему навчання англійської мови, яка інтегрує сучасні алгоритми штучного інтелекту, моделі IRT та інструменти хмарних обчислень, зокрема інфраструктурні рішення Hewlett Packard Enterprise, для забезпечення високої продуктивності та стійкості.

Мета дослідження. Полягає в концептуальному обґрунтуванні Intelligent Adaptive Personalized Learning System (IAPLS - Інтелектуальна система персоналізованого навчання), та реалізації принципу «живої траєкторії навчання», що динамічно адаптується до когнітивних, поведінкових та емоційних параметрів користувача, за допомогою моделі інтелектуальної платформи для персоналізованого вивчення англійської мови, побудованої на основі адаптивних алгоритмів та редукованих моделей першого наближення.

Результати дослідження. Вважається, що інтелектуальна платформа для вивчення мови є практичною реалізацією певної моделі мовного навчання, тобто технологічним продуктом або середовищем, яке відтворює відповідну теоретичну концепцію й забезпечує повний цикл освітньої взаємодії. Така платформа характеризується наявністю інтерфейсу користувача, розвиненої інфраструктури даних (сервер, база знань, аналітичні підсистеми), адаптивних механізмів добору навчального матеріалу, модулів опису й подання контенту та інших компонентів, необхідних для організації персоналізованого навчального процесу. Прикладами

таких систем є Grammarly, Duolingo, ELSA Speak, Google Assistant та VoiceThread AI [1, 2, 3, 4].

Кожна з цих платформ ґрунтується на власній теоретичній моделі, що визначає принципи її функціонування. У такому контексті модель виступає «мозком» системи, тобто формальною структурою, яка описує, як система аналізує знання, приймає рішення та адаптує навчальні матеріали. Натомість платформа є її «тілом» — технічною інфраструктурою, яка реалізує зазначені механізми на практиці. Модель визначає архітектуру платформи, структуру модулів, зв'язки між ними, а також алгоритми аналізу знань, прогнозування прогресу та формування рекомендацій. Наприклад, Grammarly є інтелектуальним асистентом для перевірки письмових текстів, що використовує поєднання граматичних правил і машинного навчання [1]. Duolingo впроваджує принципи гейміфікації, ELSA Speak застосовує нейронні мережі для аналізу вимови, а Google Assistant і VoiceThread AI забезпечують діалогові та мультимодальні форми взаємодії [4, 5].

Водночас наявні моделі, реалізовані у сучасних платформах, поступово втрачають актуальність, оскільки не завжди здатні відобразити складну динаміку змін у знаннях, мотивації та когнітивних характеристиках користувача. Це зумовлює необхідність пошуку нових підходів до побудови персоналізованих систем, які можуть швидко адаптуватися до індивідуальних особливостей учня та змін навчального контексту.

Одним із таких підходів є створення моделі “живої траєкторії навчання” на основі застосування адаптивних алгоритмів для персоналізованого вивчення англійської мови. Така модель формує основу інтелектуальної платформи, що забезпечує динамічну персоналізацію навчання та створення гнучкої індивідуальної траєкторії, яка враховує когнітивні особливості кожного учня і безперервно адаптується до його прогресу.

У цьому дослідженні для концептуального опису роботи інтелектуальної платформи, яка реалізує моделі “живої траєкторії навчання” запропоновано використовувати математичні формулювання у вигляді редукованих моделей першого наближення. Такий підхід дозволяє поєднати формальну чіткість моделювання з практичністю їх застосування. Редуковані моделі забезпечують достатній рівень точності для відображення ключових закономірностей навчального процесу, не потребуючи надмірно складного математичного апарату, що є характерним для повномасштабних стохастичних або багатовимірних систем. Це дає змогу створити узагальнений, але водночас інформативний формальний опис роботи модулів платформи, підтримати побудову адаптивних алгоритмів та сформулювати основу для реалізації концепції «живої траєкторії навчання».

Для досягнення поставленої мети були вирішено такі завдання дослідження:

- 1) проведено порівняльний аналіз існуючих платформ у галузі навчання мов (Grammarly, Duolingo, ELSA Speak, Google Assistant, VoiceThread AI) для виявлення їхніх переваг, обмежень та рівня персоналізації навчального процесу;
- 2) визначені ключові компоненти інтелектуальної платформи, які забезпечують створення моделі “живої траєкторії навчання”;

3) розроблено математичні формулювання у вигляді редукованих моделей першого наближення для кожного модуля платформи, які адекватно описують динаміку навчального процесу у дискретному часі та забезпечують можливість побудови адаптивних алгоритмів.

П'ятивимірна модель IAPLS персоналізації навчання англійської мови зроблена на основі адаптивних алгоритмів. Ядро системи IAPLS – це три адаптивні алгоритми: швидке тестування рівня через Item Response Theory (математична модель, що оцінює рівень знань учня не за кількістю правильних відповідей, а за складністю завдань, які він розв'язує; генерація персоналізованих уроків за допомогою великих мовних моделей штучного інтелекту; динамічна перебудова маршруту учня після кожного виконаного завдання. Ці адаптивні алгоритми будують навчальну траєкторію з п'яти вимірів профілю користувача: когнітивного (знання, помилки), мотиваційного (цілі, інтереси), поведінкового (темп, регулярність), стильового (формати контенту), контекстного (де застосовується мова). Траєкторія навчання будується під нього, а не під абстрактного «середнього учня».

Система виконує адаптивне тестування на основі двопараметричної моделі Item Response Theory. Метод Ньютона-Рафсона оцінює параметри завдань та визначає рівень CEFR. Персоналізовані уроки генеруються через API великих мовних моделей – п'ятивимірний профіль конвертується у промпт для ШІ, що створює контент під конкретні потреби. Особливістю такого підходу у тому, що якщо студент готується до IELTS - система дає академічну лексику та формальне письмо, якщо він програміст – технічну термінологію та документацію, якщо учень готується до подорожей – розмовні фрази та діалоги з аеропорту чи готелю. Динамічна корекція траєкторії спирається на функцію корисності теми, що враховує готовність учня, релевантність його цілям та історію взаємодії. Після кожного завдання траєкторія навчання корегується з урахуванням засвоєння навчального матеріалу. Алгоритм розбиває тему на дрібніші частини та додає додаткові пояснення.

Розгортання системи IAPLS може бути реалізоване на базі сучасних хмарних та інфраструктурних рішень Hewlett Packard Enterprise. Використання HPE GreenLake забезпечує гнучке масштабування обчислювальних ресурсів для роботи адаптивних алгоритмів у режимі реального часу. Інструменти HPE Ezmeral ML Ops дозволяють організувати повний цикл керування моделями штучного інтелекту — від тренування до промислового розгортання та моніторингу. Застосування серверів HPE ProLiant та сховищ HPE Alletra забезпечує високу продуктивність при роботі з великими масивами освітніх даних та логів користувачів. Для забезпечення масштабованого й надійного виконання алгоритмів адаптивного навчання система може бути розгорнута на платформі **HPE GreenLake**, що забезпечує гнучке хмарне керування ресурсами. Компоненти для тренування моделей ППН та обробки даних інтегруються через **HPE Ezmeral ML Ops**, яка підтримує повний цикл роботи з AI-моделями. Архітектура може бути розгорнута на серверних рішеннях HPE ProLiant DL360, що забезпечують високу продуктивність для обробки запитів і навчання моделей. Сховище даних може бути реалізоване на базі HPE Alletra з підтримкою швидкого доступу до

навчальних логів та профілів користувачів. Система працює у воєнний час – доступна з будь-якого пристрою, не прив'язана до місця, зберігає прогрес у хмарі.

Висновки. П'ятивимірне профілювання працює краще лінійного. Item Response Theory скорочує тестування втричі, великі мовні моделі генерують контент під кожного, динамічна корекція перебудовує маршрут щохвилини. Це не «адаптивність» у стилі Duolingo, що змінює складність наступного завдання. Це перебудова всієї траєкторії на основі п'яти параметрів замість одного. Система вже відповідає вимогам дистанційної освіти у воєнний час та стандартам CEFR для євроінтеграції. Подальший розвиток системи може включати застосування рішень НРЕ для апаратного прискорення моделей розпізнавання мовлення та масштабування сервісів через НРЕ GreenLake для міжнародних освітніх платформ. Потім – адаптація під інші мови з відстеженням міжмовної інтерференції.

Список використаних джерел

1. Postgraduate students' voices on leveraging Grammarly as an AI-powered tool in academic writing. [URL]: https://www.scielo.org.za/scielo.php?script=sci_arttext&pid=S2520-98682025000100007&lng=en&nrm=iso&tlng=en (date of access: 02.09.2025).
2. Mobile language app learners' self-efficacy increases after using generative AI [URL]: <https://www.frontiersin.org/journals/education/articles/10.3389/feduc.2025.1499497/full> (date of access: 02.09.2025).
3. Drachsler, H., Hummel, H. G. K., & Koper, R. (2015). Navigation support for learners in informal learning networks. *Journal of Computer Assisted Learning*, 31(3), 292–308.
4. Dillenbourg, P., & Tchounikine, P. (2007). Flexibility in macro-scripts for computer-supported collaborative learning. *Journal of Computer Assisted Learning*, 23(1), 1–13.
5. AI Chatbots in English Language Teaching and Learning: A Critical Review Lam Thi Quynh Anh [URL]: <https://jklst.org/index.php/home/article/view/208/181> (date of access: 02.09.2025).

Сагайдак В.А., доцент ДУІКТ

ARTIFICIAL INTELLIGENCE FOR WIRELESS LOCAL NETWORK PERFORMANCE ANALYSIS

Problem statement. WLAN radio is a most common access type to LAN for users, but it is also easy to make interventions in such network by other Wi-Fi networks, microwaves, solid steel walls or other materials with high attenuation for low frequency signal such as 5 GHz, distance coverage, access point location in room. For network administrator it is difficult to define root cause for each network user especially if

specific services are affected only for one device while others don't even recognize any issues due to extensive bandwidth usage in LAN without proper QoS management.

Research objective. In order to meet described above WLAN challenges AI agent is required that could give insights regarding most common network services (Web browsing, Instant messaging, streaming, e-learning, VR/AR, etc.), AP coverage and configuration within network, historical analysis for users with poor experience, network blind zones, interference from other WLAN networks or between several APs in one network.

Research results. Two WLAN solutions are capable of described – Huawei iMaster NCE CampusInsight and Juniper Mist. CampusInsight use AI algorithms for such feature as intelligent radio calibration and AI roaming. Intelligent radio calibration feature compares network metrics such as average physical-layer bandwidth of terminals, average interference rate, average channel utilization before and after calibration each time when feature is running. In addition, this feature displays differences in channel, power, and frequency bandwidth for each AP before and after the calibration. AI roaming feature is used to increase roaming success rate? reduce packet loss, delay during roaming. Juniper Mist have integrated virtual network assistant Marvis. Marvis Actions component can automatically fix issues in self-driving mode or recommend actions in driver-assist mode. Following component can check issues across wired, WAN, wireless network depending on layer (MSP layer, organization level layer, site level layer). Also, user can check SW version compliance on APs, identify bad cables, locate network loops, detect WAN link outages.

Conclusions. AI realization varies from vendor to vendor. Some solutions provide small features based on user experience, while others provide more complex solutions that analyze wider range and parts of network. Still, most of it are located inside of vendors ecosystem that makes it more difficult to retrain or develop additional features depending on network situation that will create a challenge for opensource solution or capability to bring your own AI algorithms or solution for integration.

References

1. Huawei iMaster NCE-CampusInsight DataSheet. Huawei Enterprise. URL: <https://e.huawei.com/ua/documents/products/enterprise-network/2f55644fbe3f4628a061ed26cf8bf541>.
2. Juniper Mist AI-Native Operations Guide | Mist | Juniper Networks. Juniper Networks, Now Part of HPE – Leading the Convergence of AI & Networking. URL: <https://www.juniper.net/documentation/us/en/software/mist/mist-aiops/index.html>.

Chorna V.V., assistant USUST
Valuiskyi H.O., student USUST

THE ROLE OF THE JAVA VIRTUAL MACHINE IN ACHIEVING CROSS-PLATFORM COMPATIBILITY

Problem Statement. Modern software development faces significant challenges in achieving cross-platform compatibility while maintaining performance, security, and reliability. Traditional programming languages like C and C++ allow fine-grained control over hardware, but this flexibility often leads to complex code, memory management errors, and platform-dependent behavior. As computing environments diversify—from handheld devices to large-scale servers—the need for a language and runtime system that ensures both portability and consistent behavior has become increasingly critical. Java was designed to address these challenges, but understanding how its design principles translate into practical software portability and secure execution requires systematic investigation.

Research Aim. This study aims to examine how Java and the Java Virtual Machine (JVM) enable cross-platform compatibility, simplify software development, and ensure secure execution. It focuses on key features such as bytecode, just-in-time (JIT) compilation, and runtime memory management, as well as Java's object-oriented and concurrency capabilities, to understand their role in creating reliable and portable applications.

Research Results. Java is a general-purpose programming language designed to support concurrency, object-oriented design, and a wide range of application domains. Although its syntax resembles that of C and C++, Java intentionally omits several features that complicate software development—for example, explicit pointer manipulation—so that programs are easier to reason about and less prone to certain classes of errors. The language was initially conceived as a way to bring consistency and interoperability to consumer electronics. Early prototypes targeted interactive television systems, but the industry wasn't adapted for such advanced abilities. The timing aligned far better with the rapid growth of the internet, where the idea of running the same program across many types of machines became especially valuable. This goal of seamless portability ultimately became a central design principle. To make this possible, Java programs execute on the Java Virtual Machine (JVM), an abstract execution environment created at Sun Microsystems. The first versions ran on devices comparable to early handheld organizers, but modern JVMs now operate on everything from mobile devices to large server systems [1]. Conceptually, the JVM behaves like a virtual processor: it defines its own instruction set, manages memory regions during program execution, and provides a uniform runtime environment regardless of the underlying hardware or operating system. Its design does not rely on any particular implementation approach—it may be interpreted, compiled to native instructions, or even embedded directly into hardware [2]. Programs that run on the JVM are written in Java or another compatible language and are compiled into Java bytecode. The JVM itself is unaware of the Java language; instead, it works exclusively with class files, a compact binary format containing bytecode instructions, symbolic information, and

structural metadata. Bytecode's fixed, one-byte instruction encoding keeps class files relatively small while also enabling the JVM to enforce strict rules about code structure and safety. Once compiled, a bytecode application can run on any system with a suitable JVM. During execution, the bytecode may be interpreted line by line or translated into native machine code using just-in-time (JIT) compilation [3]. Because the bytecode format and JVM behavior are standardized, applications typically run in a consistent manner across different platforms, providing both portability and a controlled, secure execution environment.

Conclusions. Java, in conjunction with the Java Virtual Machine, provides a robust solution to the longstanding challenge of software portability and security. Its architecture—centered on bytecode execution, abstraction from hardware, and strict runtime checks—allows developers to write programs that run consistently across heterogeneous systems. The language's object-oriented design, combined with built-in concurrency support, simplifies the development of complex applications while reducing the likelihood of common programming errors. Overall, Java's principles of portability, reliability, and safety make it a versatile tool for modern computing, from mobile devices to enterprise-level servers, demonstrating the continued relevance of its design in contemporary software engineering.

References

1. Lindholm T. The java virtual machine specification, java SE 8 edition. Addison-Wesley Professional, 2014. 600 p.
2. Holmes D., Gosling J., Arnold K. Java programming language. Pearson Education, Limited, 2025. 992 p.
3. Contributors to Wikimedia projects. Java bytecode - Wikipedia. Wikipedia, the free encyclopedia. URL: https://en.wikipedia.org/wiki/Java_bytecode (date of access: 08.12.2025).

Чорна В.В., асистент ПДАБА
Маслюков І.С., студент ПДАБА

ХМАРНІ ТЕХНОЛОГІЇ ЯК КЛЮЧОВИЙ ІНСТРУМЕНТ ЦИФРОВОЇ ТРАНСФОРМАЦІЇ СУЧАСНИХ ОРГАНІЗАЦІЙ

Постановка задачі. Хмарні технології сьогодні розглядаються як фундамент цифрової трансформації організацій різних секторів економіки. Вони забезпечують гнучкість, масштабованість та високу доступність інформаційних ресурсів, що дозволяє підприємствам оперативно адаптуватися до змін ринкового середовища. Наявність моделей IaaS, PaaS і SaaS відкриває можливості для гнучкого управління робочими процесами, прискорює розробку нових ІТ-рішень і забезпечує безперервність бізнес-процесів.

Мета дослідження. Розвиток хмарних обчислень тісно пов'язаний із впровадженням контейнеризації та мікросервісної архітектури, що дозволяють

створювати стійкі, незалежні та легко масштабовані програмні системи. Використання Kubernetes, Docker та інших оркестраційних інструментів забезпечує ефективний розподіл навантаження, автоматичне відновлення сервісів і можливість гнучкого розгортання застосунків у гібридних та мультихмарних середовищах. Ці підходи стають стандартом у сучасній розробці програмного забезпечення і визначають нові тенденції в управлінні ІТ-інфраструктурою.

Окрему увагу дослідників та практиків привертає інтеграція технологій штучного інтелекту та машинного навчання з хмарними рішеннями. Хмарні платформи пропонують інструменти для обробки великих даних, створення та навчання моделей ШІ, а також автоматизації інтелектуальних процесів у бізнесі. Завдяки цьому компанії отримують можливість здійснювати глибоку аналітику, прогнозувати ринкові тенденції, покращувати взаємодію з клієнтами та оптимізувати операційні процеси. Хмарні рішення для ШІ суттєво знижують бар'єр входу для організацій, які не мають власних потужних обчислювальних ресурсів.

Результати дослідження. Особливої актуальності набувають мультихмарні підходи, які забезпечують гнучке поєднання сервісів від різних постачальників. Це дозволяє мінімізувати ризики залежності від одного провайдера, підвищувати відмовостійкість та забезпечувати оптимальне використання ресурсів. Мультихмарні архітектури створюють нові можливості для організацій, які прагнуть збалансувати вартість, продуктивність та безпеку.

У контексті безпеки хмарні технології залишаються одним із найважливіших напрямів дослідження та розвитку. Зі зростанням обсягів переміщених у хмару даних загострюються ризики кібератак, витоків інформації та несанкціонованого доступу. Вирішення цих проблем вимагає комплексного підходу, що включає впровадження сучасних засобів шифрування, політик багатофакторної автентифікації, систем моніторингу активності та використання Zero Trust-архітектури.

Важливим аспектом сучасних хмарних технологій є роль DevOps та інфраструктури як коду (IaC). Ці підходи дають змогу автоматизувати розгортання, тестування та супровід програмних рішень у хмарі, що суттєво скорочує витрати часу та підвищує стабільність систем. IaC-технології, такі як Terraform, Ansible, Pulumi, дозволяють формалізувати інфраструктуру у вигляді конфігураційних файлів, які легко масштабувати, переносити та контролювати.

Економічний ефект від використання хмарних рішень проявляється у можливості переходу підприємств від капітальних інвестицій (CapEx) до операційних витрат (OpEx). Організації отримують можливість оптимізувати бюджети, сплачуючи лише за реально спожиті ресурси, що є важливою перевагою для малих і середніх підприємств. Хмара також забезпечує швидше впровадження інновацій, оскільки дозволяє скоротити час, необхідний для розгортання нових ІТ-продуктів та сервісів. Це дає компаніям змогу залишатися конкурентоспроможними в умовах динамічних ринкових змін.

Підвищений інтерес викликають також екологічні аспекти хмарних технологій. Центри обробки даних, які використовуються великими хмарними провайдерами, активно переходять на енергоефективні моделі та відновлювані

джерела енергії. Такі рішення сприяють зменшенню вуглецевого сліду ІТ-інфраструктури та підтримують глобальні тенденції до сталого розвитку.

Хмарні технології відіграють ключову роль у розвитку сучасних цифрових сервісів — від онлайн-освіти до електронної медицини, від електронної комерції до інтелектуальних систем управління містами. Масштабованість і доступність хмарних платформ роблять можливими проекти, які раніше вимагали значних інвестицій і складних технічних рішень. Сучасні користувачі очікують високої швидкості роботи застосунків, стабільності сервісів та безперервної доступності.

Висновки. Таким чином, хмарні технології створюють стратегічні переваги для розвитку бізнесу та суспільства, сприяють діджиталізації та впровадженню інновацій, формують нову модель побудови інформаційної інфраструктури. Системи на базі хмарних рішень стають основою цифрової економіки, визначають темпи технологічних змін і служать ключовим інструментом у розбудові конкурентоспроможних організацій майбутнього.

Список використаних джерел

1. Бойко, О. В. *Хмарні технології в інформаційних системах підприємств*. Київ: Наукова думка, 2020.
2. Гриценко, Л. М. *Сучасні тенденції розвитку хмарних обчислень*. Львів: ЛНУ ім. Франка, 2021.
3. Коваленко, С. П. *Інфраструктурні моделі та безпека хмарних сервісів*. Харків: ХНУРЕ, 2019.
4. Тарасенко, Ю. О. *Хмарні технології та цифрова трансформація бізнесу*. Одеса: ОНУ ім. Мечникова, 2022.

Чорна В.В., асистент ПДАБА

Господицько М.А., студент ПДАБА

ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ ТА ШТУЧНИЙ ІНТЕЛЕКТ У СУЧАСНІЙ ПРОМИСЛОВІЙ ІНДУСТРІЇ

Постановка задачі. Промислова сфера сьогодні переживає масштабну цифрову модернізацію, у центрі якої — технології штучного інтелекту. Інтеграція ШІ в виробничі, логістичні, енергетичні та сервісні процеси відкриває можливості для підвищення продуктивності, оптимізації витрат та формування сучасних моделей управління підприємствами. Інтелектуальні алгоритми, включаючи нейромережеві системи, забезпечують автоматизацію рутинних операцій, підвищують точність контролю якості та дозволяють прогнозувати технічні збої обладнання. Завдяки цьому скорочується тривалість виробничих циклів, зменшується вплив людського чинника та підвищується загальна надійність промислових процесів.

Мета дослідження. Інформаційні технології стали базою для розвитку концепції «розумної фабрики», де ШІ інтегрується з цифровими двійниками

виробничих об'єктів, симуляційними платформами та системами моніторингу в реальному часі. Такий підхід дозволяє моделювати процеси, прогнозувати можливі несправності обладнання, оцінювати ефективність виробничих ланцюгів та запобігати ризикам ще до появи проблем у реальному виробництві. Використання цифрових двійників у поєднанні з аналітичними платформами забезпечує підприємствам більш точне планування виробництва, дозволяє скоротити час виходу нових продуктів на ринок та підвищує якість продукції.

Штучний інтелект у промисловій сфері також активно трансформує процеси управління. Інтелектуальні системи здатні обробляти величезні масиви даних і на основі цього формувати прогнози щодо потреб у матеріалах, оптимізувати логістику та розподіл ресурсів, а також виявляти вузькі місця виробничих ланцюгів. Такі платформи дозволяють керівникам підприємств приймати обґрунтовані стратегічні рішення, підвищувати ефективність праці та автоматизувати рутинні процеси управління. Це особливо важливо в умовах глобальної конкуренції та постійних змін ринкових умов, коли швидкість адаптації стає ключовим фактором успіху.

Результати дослідження. Особливу роль відіграє впровадження генеративних алгоритмів, які дозволяють створювати прототипи продукції, моделювати процеси та проводити цифрові випробування до початку реального виробництва. Такі технології забезпечують швидке тестування нових рішень, дозволяють передбачити поведінку матеріалів і компонентів, а також підвищують точність розрахунків. Впровадження інновацій на основі ІІІ сприяє підвищенню конкурентоспроможності підприємств та стимулює розвиток індустріальних екосистем.

Водночас, активне застосування ІІІ у промисловості пов'язане з певними викликами. Основними серед них є підготовка висококваліфікованих фахівців, здатних ефективно працювати з інтелектуальними системами та аналітичними платформами, забезпечення кібербезпеки та захисту промислових даних. Етичні питання залишаються не менш актуальними: потрібно визначати відповідальність за автоматизовані рішення, оцінювати вплив роботизації на зайнятість, а також створювати політики прозорості у використанні даних.

Перспективи розвитку промисловості на основі ІІІ та ІТ безпосередньо пов'язані з інтеграцією Інтернету речей (IoT), великих даних (Big Data), хмарних обчислень та роботизованих систем. Такий комплексний підхід дозволяє створювати інтелектуальні виробничі середовища, які здатні самостійно адаптуватися до внутрішніх та зовнішніх змін, оптимізувати ресурси, підвищувати продуктивність і забезпечувати високу якість продукції. Інтелектуальні системи стають не лише інструментами автоматизації, а й партнерами у прийнятті стратегічних рішень, управлінні ризиками та впровадженні інноваційних рішень.

У сучасних умовах цифрової трансформації підприємств особливу увагу слід приділяти взаємодії між людьми та машинами. Підприємства, що впроваджують ІІІ, повинні поєднувати розвиток компетенцій персоналу з автоматизацією виробництва, створюючи синергію між людським інтелектом і можливостями інтелектуальних систем. Крім того, важливим є розвиток

стандартів безпеки та протоколів для захисту даних, що гарантує стабільність і надійність промислових процесів.

Висновки. Таким чином, інтеграція інформаційних технологій та штучного інтелекту у промислову сферу є ключовим чинником підвищення ефективності, конкурентоспроможності та інноваційності підприємств. Вона забезпечує можливість прогнозування, оптимізації та управління складними виробничими процесами, сприяє створенню інтелектуальних виробничих середовищ та розвитку нових бізнес-моделей. Ефективне впровадження таких систем потребує підготовки висококваліфікованих фахівців, дотримання етичних норм та стандартів безпеки, що забезпечує стійкий розвиток промислових підприємств у цифрову епоху.

Список використаних джерел

1. Choudhury A., et al. Intelligent Manufacturing and AI: Emerging Trends and Case Studies. Springer, 2025.
2. Lee J., Bagheri B., Kao H. A. A Cyber-Physical Systems Architecture for Industry 4.0-based Manufacturing Systems. Manufacturing Letters, 2022.
3. Li B., Hou B., Yu W., Lu X., Yang C. Applications of Artificial Intelligence in Industry 4.0: A Review. Journal of Industrial Information Integration, 2023.
4. Pan S. J., Yang Q. A Survey on Transfer Learning. IEEE Transactions on Knowledge and Data Engineering, 2024.

Чорна В.В., асистент ПДАБА

Калашникова К.С., студентка ПДАБА

ІНФОРМАЦІЙНО-ПОШУКОВІ ТА ЕКСПЕРТНІ СИСТЕМИ: ІННОВАЦІЙНІ ПІДХОДИ В ОСВІТНІЙ СФЕРІ ВІД ІТ-КОМПАНІЙ

Постановка задачі. У тезах досліджується застосування інформаційно-пошукових та експертних систем у сучасній освіті завдяки впровадженню інноваційних рішень від ІТ-компаній. Розглядаються способи автоматизації пошуку навчальної інформації, аналізу знань студентів, персоналізованих рекомендацій та інтеграції експертних платформ у навчальні програми. Особлива увага приділяється ролі штучного інтелекту у створенні адаптивного навчання та оптимізації освітніх процесів. Окремо розглядаються переваги й виклики впровадження таких систем, включаючи питання підготовки педагогів, забезпечення конфіденційності даних та економічні аспекти. Тези орієнтовані на викладачів, освітніх менеджерів, розробників цифрових платформ та дослідників у сфері освітніх технологій.

Мета дослідження. Сучасна освіта переживає етап інтенсивної цифровізації, коли традиційні методи навчання стають недостатньо ефективними для студентів і викладачів. Потреба у швидкому доступі до актуальної інформації, аналізі навчальних результатів та персоналізації навчального процесу зумовлює розвиток

інформаційно-пошукових та експертних систем. ІТ-компанії активно пропонують комплексні рішення, що поєднують алгоритми пошуку, машинне навчання та експертну логіку, інтегруючи їх у навчальні платформи. Ці системи оптимізують навчальні траєкторії студентів, сприяють підвищенню ефективності освітніх процесів та забезпечують інноваційний підхід до підготовки фахівців.

Інформаційно-пошукові системи дозволяють швидко знаходити потрібні знання, ресурси та навчальні матеріали. Вони виконують індексацію даних, семантичний пошук, ранжування інформації та інтеграцію з навчальними платформами. Сучасні ППС від ІТ-компаній забезпечують автоматичне підключення до LMS (Learning Management Systems), що дозволяє студентам отримувати релевантні навчальні матеріали без додаткових зусиль. Рекомендаційні алгоритми пропонують курси та матеріали, виходячи з попередніх результатів студентів, а AI-асистенти надають швидкі відповіді на запитання та підтримку під час самостійного навчання. Семантичний аналіз текстів допомагає оцінювати якість матеріалів та відокремлювати актуальну інформацію від застарілої. Впровадження таких систем зменшує час на пошук знань, підвищує ефективність навчання і сприяє формуванню навичок самостійного дослідження.

Результати дослідження. Експертні системи моделюють процес прийняття рішень фахівця та застосовуються для автоматизованого оцінювання знань, рекомендацій щодо навчальних траєкторій і контролю успішності студентів. У навчальному процесі такі системи дозволяють аналізувати індивідуальні результати студентів і формувати персоналізовані завдання, підтримувати адаптивне навчання відповідно до рівня підготовки, створювати інтерактивні навчальні кейси та симуляції професійних ситуацій, а також допомагати викладачам у виборі оптимальних методик викладання та оцінювання складних дисциплін. Використання експертних систем стимулює розвиток критичного мислення, підвищує якість навчання та сприяє активному залученню студентів у процес навчання.

Сучасні ІТ-компанії пропонують освітнім закладам готові рішення, що поєднують функції інформаційно-пошукових та експертних систем. До них належать платформи Microsoft Education із AI-інструментами для аналізу навчальних досягнень, Google Workspace for Education з інтегрованими AI-асистентами та пошуковими алгоритмами, Coursera та edX із системами рекомендацій і адаптивними тестами, а також локальні стартапи, що розробляють експертні платформи з автоматизованим оцінюванням і персоналізацією навчання. Такі рішення дозволяють викладачам зосередитися на творчій роботі, а студентам отримувати персоналізовану підтримку та оптимізовані навчальні траєкторії.

Переваги застосування інформаційно-пошукових та експертних систем включають швидкий доступ до великого обсягу знань і матеріалів, індивідуалізацію навчання, підвищення мотивації студентів, покращене управління ресурсами та оцінювання результатів, а також інтеграцію з різними цифровими платформами для комплексного навчання. Водночас існують виклики, серед яких фінансові витрати на ліцензії та технічну підтримку, потреба у навчанні педагогів для ефективного використання систем, забезпечення

конфіденційності та захисту даних студентів, а також ризик надмірної залежності від автоматизованих рішень.

Висновки. Інформаційно-пошукові та експертні системи, впроваджені ІТ-компаніями, стають ключовим інструментом цифровізації освіти. Вони дозволяють підвищити ефективність навчального процесу, забезпечують персоналізовану підтримку студентів та інтегрують сучасні технології у щоденну освітню практику. Незважаючи на технічні та організаційні виклики, застосування таких систем є стратегічно важливим кроком для підготовки конкурентоспроможних фахівців XXI століття.

Список використаної літератури

1. Bostrom, N. *Superintelligence: Paths, Dangers, Strategies*. Oxford University Press, 2017.
2. Luckin, R., Holmes, W., Griffiths, M., Forcier, L. *Intelligence Unleashed: An Argument for AI in Education*. Pearson, 2016.
3. Siemens, G., Long, P. *Learning Analytics and Educational Data Mining*. Routledge, 2020.
4. UNESCO. *Artificial Intelligence in Education: Challenges and Opportunities*. 2023.

Чорна В.В., асистент ПДАБА
Нескреба М.С., студент ПДАБА

ІНТЕЛЕКТУАЛЬНІ АДАПТИВНІ СИСТЕМИ ПРОГНОЗУВАННЯ ПОВЕДІНКИ КОРИСТУВАЧІВ У ЦИФРОВИХ ПЛАТФОРМАХ НА ОСНОВІ МУЛЬТИМОДАЛЬНИХ AI-МОДЕЛЕЙ

Постановка задачі. Полягає у визначенні підходів до створення інтелектуальних адаптивних систем, здатних прогнозувати поведінку користувачів на цифрових платформах на основі мультимодальних моделей штучного інтелекту. Необхідно дослідити можливості комплексної обробки різнорідних даних — текстових, візуальних, поведінкових та біометричних — як ключового чинника підвищення точності прогнозування. Важливим є виявлення основних викликів, пов'язаних із впровадженням таких систем, зокрема етичних аспектів, питань приватності, ризику упередженості моделей та потреби у значних обчислювальних ресурсах. Також постає завдання обґрунтувати концепцію адаптивної архітектури, що здатна до самооптимізації та навчання в режимі реального часу. Сформульована проблематика має практичну цінність для розробників цифрових продуктів, систем рекомендацій, маркетингових сервісів, рішень у сфері кібербезпеки та інтелектуальної аналітики.

Мета дослідження. Сучасні цифрові сервіси активно переходять від статичних алгоритмів персоналізації до глибинних інтелектуальних систем, які здатні обробляти мільйони сигналів користувацької активності та будувати

прогнози з високою точністю. Поява потужних мультимодальних AI-моделей — таких як GPT-5, Claude, Gemini чи LLaMA-vision — відкрила можливість аналізувати контент у різних форматах, що значно підвищує точність прогнозування поведінкових патернів.

Результати дослідження. Актуальність теми зумовлена стрімким ростом цифрової взаємодії в е-комерції, освіті, соціальних мережах, мобільних застосунках, а також у системах безпеки. У той час як традиційні моделі враховують лише окремі типи даних (наприклад кліки або текстові відгуки), мультимодальні рішення об'єднують контекст, інтонацію голосу, рух очей, візуальні об'єкти, метрики активності, психологічні індикатори тощо. Це створює підґрунтя для появи нового покоління інформаційних технологій — інтелектуальних адаптивних систем прогнозування, які здатні динамічно змінювати свою логіку відповідно до користувацької поведінки.

Основний зміст

1. Мультимодальні AI-моделі як рушій прогнозних систем.

Мультимодальні моделі здатні одночасно обробляти різноманітні дані, створюючи глибокі узагальнення та прогнозні висновки. На відміну від класичних алгоритмів машинного навчання, вони інтегрують:

- 1) текстові дані (повідомлення, коментарі, пошукові запити);
- 2) візуальний контент (зображення, відео, скріншоти);
- 3) поведінкові патерни (тривалість сесій, маршрути навігації);
- 4) сенсорні дані (голос, тембр, мікроміміки);
- 5) контекстуальну інформацію (місцезнаходження, час доби).

Наприклад, платформа може передбачити, коли користувач шукає допомогу, коли він незадоволений, коли готовий до покупки, коли ризикує покинути сервіс або коли потребує додаткової підтримки. Це підвищує ефективність маркетингу, якість клієнтського досвіду та рівень безпеки цифрових продуктів.

Окремим напрямом розвитку є створення адаптивних інформаційних систем, здатних підлаштовуватися у реальному часі. Такі платформи характеризуються динамічним оновленням моделей без зупинки роботи, безперервним навчанням на основі нових користувацьких даних, наявністю механізмів самокорекції для зменшення помилок прогнозування, використанням нейронних ансамблів, які переключаються залежно від контексту користувача, а також застосуванням lightweight-моделей на пристроях, коли повні моделі недоступні.

Застосування таких систем спостерігається у банківській сфері (виявлення шахрайства), освітніх платформах (адаптивне навчання), медицині (аналіз стану користувача), кібербезпеці, ігровій індустрії та комерційних сервісах.

Разом із розвитком технологій зростає потреба в етичному регулюванні, оскільки прогнозні системи можуть ненавмисно створювати низку ризиків, зокрема втручання у приватність, алгоритмічні упередження щодо окремих груп користувачів, надмірну персоналізацію, яка формує «інформаційні бульбашки», маніпулятивні рекомендації та відсутність прозорості логіки прийняття рішень. Для мінімізації цих ризиків важливо впроваджувати підходи explainable AI, що забезпечують пояснюваність моделей, дотримуватися принципів етичного

дизайну, розробляти чіткі політики збору даних, проводити регулярний аудит моделей та встановлювати обмеження на використання чутливих параметрів.

У майбутньому очікується поява повністю автономних цифрових асистентів, які зможуть передбачати складні наміри користувача, проводити багатокрокові сценарні прогнози, адаптувати інтерфейс у реальному часі, підтримувати багатомодальну взаємодію та забезпечувати когнітивну підтримку в навчанні й професійній діяльності.

Висновки. Мультимодальні моделі та адаптивні інформаційні системи формують новий етап розвитку інтелектуальних технологій. Їх впровадження дозволяє цифровим платформам прогнозувати поведінку користувачів із безпрецедентною точністю та гнучкістю. Незважаючи на низку етичних та технічних викликів, потенціал таких систем є надзвичайно високим: від покращення персоналізації до зміцнення безпеки та підвищення якості взаємодії людини з технологіями.

Список використаної літератури

1. Goodfellow I., Bengio Y., Courville A. Deep Learning. MIT Press, 2016.
2. Google DeepMind. Multimodal Intelligence Reports, 2023–2024.
3. IBM AI Ethics Guidelines. IBM Research, 2022.
4. Müller V. Philosophy of Artificial Intelligence. Oxford University Press, 2020.

Чорна В.В., асистент ПДАБА

Смірнов Д.Є., студент ПДАБА

ШТУЧНИЙ ІНТЕЛЕКТ В ІНДУСТРІЇ: ВІД АВТОМАТИЗАЦІЇ ПРОЦЕСІВ ДО ЦИФРОВОЇ ТРАНСФОРМАЦІЇ ПІДПРИЄМСТВ

Постановка задачі. Полягає у всебічному дослідженні впливу цифрової трансформації на розвиток та модернізацію промислових систем, у межах яких ключова роль належить технологіям штучного інтелекту. Завданням є визначення механізмів, за допомогою яких інтеграція ШІ у виробничі, логістичні, енергетичні та сервісні процеси підвищує ефективність роботи підприємств, сприяє оптимізації витрат і формує нові бізнес-моделі. Окрему увагу слід приділити аналізу роботи інтелектуальних алгоритмів і нейромереж, що забезпечують автоматизацію рутинних операцій, вдосконалення контролю якості та прогнозування технічних несправностей обладнання. Усе це безпосередньо впливає на скорочення виробничих циклів, підвищення надійності систем та мінімізацію людського фактору.

Мета дослідження. У сфері логістики та управління ланцюгами постачання ШІ-аналітика допомагає прогнозувати попит, оптимізувати маршрути доставки та зменшувати запаси, що дозволяє підприємствам підвищувати швидкість обслуговування клієнтів та знижувати витрати. Інтелектуальні системи управління енергоспоживанням забезпечують ефективне використання ресурсів,

дозволяють передбачати пікові навантаження та зменшують витрати на енергію, сприяючи сталому розвитку підприємств. Крім того, генеративні моделі застосовуються для створення прототипів продукції, оптимізації конструкцій та пошуку нових рішень у дизайні, що скорочує час від ідеї до реалізації готового продукту та стимулює інновації.

Результати дослідження. ШІ також активно впроваджується у сфері обслуговування клієнтів. Використання інтелектуальних чат-ботів, голосових асистентів та систем прогнозування дозволяє покращити взаємодію з клієнтами, підвищити лояльність до бренду та швидко реагувати на запити ринку. Переваги інтеграції ШІ в промислові процеси очевидні: підвищується продуктивність за рахунок автоматизації рутинних завдань, зменшуються витрати на ресурси та енергію, покращується якість продукції завдяки постійному контролю та аналітиці, а також прискорюється процес впровадження інновацій і адаптації до змін ринкових умов.

Водночас впровадження ШІ супроводжується низкою викликів. Для ефективної роботи систем необхідно підготувати персонал, який буде володіти новими цифровими компетенціями та навичками роботи з інтелектуальними технологіями. Використання ШІ потребує захисту даних, оскільки алгоритми обробляють великі обсяги конфіденційної інформації. Вартість розробки та інтеграції інтелектуальних рішень може бути високою, що потребує значних фінансових вкладень. Також важливим аспектом є дотримання етичних норм, особливо при використанні ШІ у прийнятті управлінських рішень, автоматизації процесів і роботизації виробництва.

Перспективи розвитку ШІ у промисловості включають інтеграцію з Інтернетом речей (IoT), цифровими двійниками (Digital Twin) та аналітикою великих даних (Big Data), що дозволяє створювати автономні та високоефективні виробничі екосистеми. Таке поєднання технологій забезпечує підвищення адаптивності підприємств до швидких змін на ринку, сприяє розвитку інноваційних продуктів і формує умови для більш ефективного управління ресурсами. Штучний інтелект у промисловості стає не просто інструментом оптимізації, а повноцінним партнером у процесах виробництва, проектування та обслуговування, що відкриває нові можливості для підвищення конкурентоспроможності підприємств на національному та міжнародному рівнях.

Висновки. Отже, ШІ виступає ключовим фактором цифрової трансформації промисловості. Його застосування дозволяє підвищити ефективність виробничих процесів, стимулює інновації та створює умови для розвитку креативних підходів у роботі підприємств. Успішна інтеграція штучного інтелекту залежить від підготовки кадрів, забезпечення безпеки даних та дотримання етичних норм. Індустрія майбутнього - це середовище, де люди та інтелектуальні системи працюють у тандемі, поєднуючи аналітику, автоматизацію та творчість, що сприяє сталому розвитку та підвищенню конкурентоспроможності підприємств.

Список використаних джерел

1. Brynjolfsson E., McAfee A. The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies. Norton, 2014.

2. Gartner Research. AI in Industry 2023: Trends and Case Studies. 2023.
3. OpenAI. AI Applications in Manufacturing and Industry. Documentation, 2024.
4. Russell S., Norvig P. Artificial Intelligence: A Modern Approach. 4th Edition. Pearson, 2021.
5. Schwab K. The Fourth Industrial Revolution. Currency, 2017.

Чорна В.В., асистент ПДАБА

Сліпко Д.Р., студентка ПДАБА

ВИКОРИСТАННЯ ІНФРАСТРУКТУРНИХ ПЛАТФОРМ HEWLETT PACKARD ENTERPRISE ДЛЯ ПІДВИЩЕННЯ ЕФЕКТИВНОСТІ ІТ-СИСТЕМ

Постановка задачі. У сучасних умовах цифровізації та активного впровадження гібридних та хмарних інфраструктур особливої актуальності набуває забезпечення стабільної, масштабованої та безпечної роботи ІТ-систем. Традиційні операційні системи не завжди здатні ефективно управляти зростаючими обсягами даних, потребами автоматизації та швидкого розгортання сервісів, що зумовлює необхідність застосування спеціалізованих інфраструктурних платформ.

Мета дослідження. Метою роботи є аналіз функціональних можливостей інфраструктурних платформ HPE та визначення їх ролі у підвищенні ефективності сучасних ІТ-систем, а також оцінка доцільності використання цих технологій у навчальному процесі.

Результати дослідження. HPE iLO забезпечує віддалене адміністрування серверів на апаратному рівні, моніторинг стану обладнання та автоматизацію оновлень. HPE OneView реалізує концепцію «інфраструктура як код», забезпечуючи централізоване керування серверами і мережами. HPE Ezmeral орієнтована на контейнеризацію та Big Data, підтримує Kubernetes. ArubaOS забезпечує

Висновки. Інфраструктурні платформи Hewlett Packard Enterprise суттєво розширюють можливості традиційних операційних систем та забезпечують автоматизацію, аналітику, надійність і безпеку сучасних ІТ-систем.

Список використаних джерел

1. Aruba Networks. ArubaOS 10 – Data Sheet. 2024.
2. Hewlett Packard Enterprise. HPE iLO 5 User Guide. 2023.
3. Hewlett Packard Enterprise. HPE OneView – Technical Overview. 2024.
4. Hewlett Packard Enterprise. HPE Ezmeral Software Platform – Architecture and Capabilities. 2023.

Чорна В.В., асистент ПДАБА

Куліков Є.О., студент ПДАБА

РОЗРОБКА ТА ВПРОВАДЖЕННЯ WEB-СИСТЕМИ РЕКОМЕНДАЦІЙ ФІЛЬМІВ MOVIEVIBE

Постановка задачі. Сьогодні коли не завжди є можливість переглянути фільм у кінотеатрі, більшість обирає зручний спосіб перегляду вдома. Для когось це не просто спосіб провести вечір, а цілий світ емоцій, історій і натхнення. Фільми та серіали допомагають розслабитися, дізнатися щось нове чи навіть переосмислити якісь моменти в житті. Зараз серед нескінченної кількості фільмів, знайти щось дійсно цікаве серед величезної кількості контенту дуже складно. Існує багато платформ з фільмами, але не всі вони зручні та орієнтовані на українськомовних користувачів.

Мета дослідження. Виникла ідея розробки веб-додатку «MovieVibe» - системи рекомендацій фільмів, яка полягає в тому, щоб створити зручний інструмент, який буде аналізувати інтереси користувача і на основі цього пропонувати фільми для перегляду.

Результат дослідження. Для створення веб-додатку «MovieVibe» - системи рекомендацій фільмів використано три основні сучасні інформаційні технології. HTML відповідає за структуру сторінок, тобто створює каркас, на якому базується весь контент — списки фільмів, кнопки, форми тощо. CSS використовується для оформлення графічного дизайну, що дозволить обрати кольори і шрифти, а також розташування основних елементів на екрані. При розробці використана темна тема (вона стильна й не напружує очі) екрану та додані плавні анімації при переходах між сторінками. Розроблений додаток адаптивний для мобільної та десктопної версії. JavaScript, у свою чергу, забезпечує інтерактивність користувача з додатком, обробка натискань кнопок, фільтрацію списків у реальному часі та відправлення даних на сервер. Застосування цих технологій створюють повноцінний он-лайн додаток, який працює прямо в браузері без необхідності встановлення додаткового програмного забезпечення. Веб-додаток має основні важливі функції, які роблять його функціональним і зручним у використанні. Користувачі можуть зареєструватися та авторизуватися, щоб отримати доступ до персоналізованих можливостей. Після входу в систему відкривається доступ до перегляду списку фільмів, де можна застосовувати фільтри: обирати жанр (комедія, драма, трилер, тощо), рік виходу на екран, рейтингом, назви, популярність, категорії. Фільтр дозволяє швидко звузити вибір до тих фільмів, які найбільше відповідають настрою чи інтересам. До кожного фільму надається повна інформація — опис сюжету, імена акторів та режисера, тривалість та навіть постером чи тизером, що допомагає краще уявити, про що фільм. Кожен користувач має можливість оцінити переглянуті фільми від 1 до 5 зірок після перегляду, що в подальшому стане основою для вдосконалення рекомендацій. Такими чином, система запам'ятовує вподобання, а потім може запропонувати фільм певного жанру або з найулюбленішим актором в головній ролі. Наступна важлива та зручна функція — можливість додавати фільми до

закладок, щоб не забути про ті, які хочеться подивитися пізніше. Додаток підтримує дві мови українську та англійську, що робить його доступним для ширшої аудиторії. Пагінація, розподіл списку фільмів на сторінки, забезпечує комфортний перегляд навіть за великої кількості позицій, тобто якщо в каталозі десятки чи сотні фільмів, вони розміщені на сторінки, щоб не гортати нескінченний список. Наступним кроком в розробці веб-додатку є впровадження алгоритмів машинного навчання, які дозволять створювати персоналізовані рекомендації. Система зможе аналізувати, які фільми користувач оцінив високо, і пропонувати схожі за стилем чи тематикою. Корисною функцією стане, коментування та обговорення фільмів, щоб користувачі могли ділитися враженнями й обмінюватися думками про фільм. Можливість інтеграції з базами даних фільмів, такими як IMDb чи TMDb, допоможе автоматично оновлювати інформацію про нові релізи, акторів чи рейтинги, що значно спростить підтримку актуальності даних.

Висновки. Розроблений веб-додаток «MovieVibe» присвячений системі рекомендацій фільмів - це зручний інструмент, який аналізує інтереси користувача. Веб додаток дозволить переглядати інформацію про фільми, сортувати їх за різними параметрами та зберігати обране в особистому списку. Це робить його корисним як для кіноманів, так і для тих, хто просто хоче швидко знайти щось цікаве для вечірнього перегляду.

Список використаних джерел

1. Crockford, Douglas. JavaScript: The Good Parts: The Good Parts. "O'Reilly Media, Inc.", 2008. Вилучено з: <https://surl.lu/nqlyxs>
2. Гризун, Людмила, and Єгор Міщенко. "Огляд інструментів frontend розробки." Тези доповідей (2022): 28. Вилучено з: <https://surl.li/xofaiu>
3. Mitchell, Scott. Create your own website. Sams Publishing, 2006. Вилучено з: <https://surl.lu/notncw>

Стойко М.М., аспірант ДУІКТ

ІІІ-МЕТОДИ ІНТЕГРАЦІЇ ГЕТЕРОГЕННИХ МЕРЕЖ У СПЕЦІАЛЬНИХ АСУ СП

Постановка Завдання

Автоматизовані системи управління спеціального призначення (АСУ СП) критично залежать від гетерогенних мереж (ГМ) — комбінації різномірних технологій (радіо, супутниковий зв'язок, мережі БПЛА, оптичні лінії).

Проблема Гетерогенності: Складність полягає у забезпеченні безшовного зв'язку, надійної маршрутизації та кіберстійкості в умовах, коли різні мережеві домени мають різні протоколи, пропускну здатність, затримку та рівень завад.

Необхідність Автоматизації: Традиційні статичні методи управління не здатні вчасно реагувати на динамічні зміни в операційному середовищі. Це

вимагає впровадження алгоритмів автоматизації, побудованих на Штучному Інтелекті (ШІ) та Машинному Навчанні (МН).

Мета: Інтеграція має забезпечити єдиний інформаційний простір для реалізації складних спеціальних операцій.

Мета Дослідження

Дослідження спрямоване на розробку та оцінку ШІ-орієнтованих методів інтеграції, які дозволяють автономно управляти гетерогенними спеціальними мережами для підвищення їхньої адаптивності, надійності та якості обслуговування (QoS) в АСУ СП.

Результати Дослідження: Ключові ШІ-Методи Інтеграції

Ефективна інтеграція ГМ у спеціальних АСУ СП реалізується через централізовану архітектуру (наприклад, SDN), можливості якої автоматизуються ШІ:

А. ШІ-Адаптивна Маршрутизація та Управління Потоками

Навчання з Підкріпленням (RL) для Маршрутизації: Агенти RL навчаються вибирати оптимальні шляхи через гетерогенні мережі для максимізації доставки даних та мінімізації затримки і ризику перехоплення/глушіння.

Прогнозування Стану Каналу: Алгоритми МН прогнозують погіршення QoS (зростання затримки чи втрат) у конкретному каналі, дозволяючи АСУ СП проактивно перенаправляти критичний трафік.

Б. Когнітивне Управління Ресурсами

Когнітивне Радіо (CR): ШІ використовується для автоматичного сканування, ідентифікації завад та динамічного перемикавання між частотними діапазонами і радіотехнологіями. Це забезпечує виживання зв'язку в умовах РЕБ і дозволяє різним мережам адаптуватися до єдиного ворожого середовища.

Автоматизована Агрегація: ШІ балансує навантаження та агрегує пропускну здатність різних мережевих технологій (наприклад, оптимізація МРТСП) для максимізації загальної доступної швидкості.

В. Автономна Кібербезпека та Управління

Злиття Гетерогенних Даних (Data Fusion): Глибоке навчання застосовується для автоматичної обробки та злиття різнорідної інформації від інтегрованих мереж, створюючи єдину ситуаційну обізнаність.

Автономне Реагування: ШІ-системи, інтегровані з SDN-контролерами, можуть автоматично генерувати та виконувати правила (зміна маршрутів, активація резервів) у відповідь на виявлені загрози чи відмови, забезпечуючи кіберстійкість.

Висновки

Застосування ШІ-алгоритмів є життєздатним і необхідним рішенням для подолання складнощів інтеграції гетерогенних мереж у спеціальних АСУ СП.

ШІ забезпечує високий ступінь автоматизації управління мережею, дозволяючи системі самостійно адаптуватися до змінних умов, динамічно керувати ресурсами та забезпечувати пріоритетну доставку критичних даних.

SDN та NFV слугують архітектурною основою, що надає ШІ необхідні програмні інтерфейси для централізованого та гнучкого управління різнорідними елементами мережі.

Список Використаних Джерел

- Wang, F., et al. (2020). Deep Reinforcement Learning for Adaptive Routing in Dynamic Heterogeneous Networks. *IEEE Transactions on Network Science and Engineering*, 7(3), 1629-1643.
- Sezer, S., et al. (2018). Software Defined Networking for Military Communications. *IEEE Communications Magazine*, 56(9), 104-110.
- Bizanis, N., & Kuipers, F. (2016). SDN and virtualization solutions for the Internet of Things: A survey. *IEEE Access*, 4, 5591–5606.
- Li, Y., & Liu, Y. (2019). AI-Based Spectrum Sensing and Management in Cognitive Radio Networks. *Journal of Communications and Information Networks*, 4(2), 158-168.
- Chen, M., et al. (2017). Machine Learning for Data Fusion and Situational Awareness in Heterogeneous Sensor Networks. *IEEE Internet of Things Journal*, 4(6), 2322-2330.
- Al-Fuqaha, A., et al. (2019). Cyber-Resilience and Trust Management in IoT and Heterogeneous Networks Using Blockchain and Machine Learning. *Computer Networks*, 153, 1-14.

Гніденко М.М., аспірант ДУІКТ

ВПЛИВ РОЗМІЩЕННЯ КОНТРОЛЕРІВ SDN НА ЕФЕКТИВНІСТЬ МЕРЕЖІ

Постановка завдання.

Прихильники стверджують, що архітектури на основі контролерів SDN спрощують проектування площини керування, покращують конвергенцію та забезпечують гнучку, еволюціонуючу мережу [1,2,3] . Критики висловлюють занепокоєння щодо затримки прийняття рішень, масштабованості та доступності [4,5]. Щоб кількісно оцінити проблеми продуктивності, звузимо увагу щодо архітектури на основі контролерів SDN до двох основних питань: скільки контролерів SDN потрібно для забезпечення надійного функціонування мережі?; де в топології повинні бути розміщені контролери SDN?

Цей вибір дизайну, проблема розміщення контролера, впливає на кожен аспект розв'язаної площини керування, від варіантів розподілу станів до відмовостійкості та показників продуктивності. У глобальних мережах з довгою затримкою поширення це накладає фундаментальні обмеження на час доступності та конвергенції. Це має практичні наслідки для проектування програмного забезпечення, впливаючи на те, чи можуть контролери реагувати на події в режимі реального часу, чи вони повинні заздалегідь передавати дії пересилання елементам пересилання.

Мета дослідження.

Метою дослідження є зосередженість на двох конкретних питаннях: скільки контролерів потрібно, враховуючи топологію і куди їх слід розміщувати? Щоб відповісти на ці питання, необхідно розглянути фундаментальні обмеження затримки поширення площини керування в майбутньому розгортанні виробничої мережі Internet, а потім розширити сферу дослідження до понад 100 загальнодоступних топологій WAN.

Результати дослідження.

У роботі пропонується нове рішення для розміщення контролерів, яке враховує як затримку, так і навантаження комутаторів для SDN. Воно має такі три особливості:

1. Модуль контролера: Замість розподілу певної кількості контролерів в SDN, пропонується використовувати обмеженої кількості модулів контролера. Модуль контролера – це набір контролерів. Як модуль контролера можна використовувати багатоядерний процесор. На основі нової кількості потоків, що генеруються комутаторами, модулі контролера активують необхідну кількість контролерів.

2. Необмежений пошук місця розташування: Існуючі дослідницькі роботи, наприклад, вибирають місце розташування контролера з набору вибраних місць розташування (місця розташування комутаторів). Але запропонована система шукає по всій області, що містить місця розташування комутаторів, щоб отримати місце розташування контролера.

3. Динамічне навантаження трафіком: Більшість нещодавніх дослідницьких робіт, наприклад з проблеми розміщення контролерів, розглядають як затримку, так і навантаження трафіком комутаторів сукупно для розміщення контролерів у програмно-визначених мережах. Вони враховують навантаження комутаторів лише один раз у поєднанні із затримкою. Хоча затримка приблизно фіксована, навантаження комутаторів змінюється з часом. Оскільки контролер розміщується на основі як затримки, так і навантаження комутаторів. Таким чином, при цьому постійно змінному навантаженні контролери необхідно постійно замінювати. Але це практично нездійсненно. Враховуючи цю дилему, затримка та навантаження комутаторів розглядаються окремо.

Висновки.

Запропоноване рішення — це одне із перших рішень для розміщення контролерів, в якому враховується динамічне навантаження комутаторів трафіком. Хоча результати досліджень розміщення контролерів враховували як затримку, так і навантаження трафіком комутаторів сукупно для визначення розташування, у роботі представлено обґрунтований аргумент на користь окремого розгляду цих двох обмежень. Пропонується комбіноване рішення проблеми розміщення контролерів та проблеми планування контролерів SDN. Він приймає два вхідні дані: розташування комутаторів SDN та необхідне обмеження затримки. Рішення видає розташування модулів контролерів та кількість активних контролерів кожного модуля контролера. Хоча перший вихідний сигнал є статичним, другий — динамічним.

Список використаних джерел

1. Гніденко М.П., Вишнівський В.В., Ільїн О.О. Побудова SDN мереж. – Навчальний посібник. – Київ: ДУТ, 2019. – 190 с.
2. Гніденко М.М., Бай Я.В. Програмно-визначена мережа: сучасний стан і виклики. IV Всеукраїнська науково-практична конференція «Сучасні інтелектуальні інформаційні технології в науці та освіті» /травень/ Київ: ДУІКТ, - 2024 р. https://duikt.edu.ua/uploads/p_2661_45318838.pdf
3. Гніденко М.П., Гніденко М.М., Вишнівський О.В., Зінченко В.В. Забезпечення енергоефективності програмно визначених мереж (SDNs) при впровадженні різних схем безпеки. / Київ: ДУІКТ. Наукові записки ДУІКТ – 2024. – №2(6).
<https://journals.dut.edu.ua/index.php/sciencenotes/article/view/3092/2982>
4. Гніденко М.М., Шепетун О.О. Проблеми масштабованості програмно-визначених мереж / Київ: ДУІКТ. Наукові записки ДУІКТ – 2025. – №1(7), с 148-154 . <https://journals.duikt.edu.ua/index.php/sciencenotes/article/view/3269/3155>
5. Прокопов С.В., Гніденко М.М., Серих С.О. Вишнівський О.В., Проблеми, вирішення яких впливають на функціональну стійкість програмно-визначених мереж (SDNs). / Київ: ДУІКТ. Зв'язок – 2024. – №6 (172).
<https://con.dut.edu.ua/index.php/communication/article/view/2823/2713>

Худік Б.О., к.т.н., доцент
Гашійук А.О., студент ДУІКТ

РОЗРОБКА МЕТОДИКИ ВИЯВЛЕННЯ ФІШИНГОВИХ ЕЛЕКТРОННИХ ЛИСТІВ ЗА ДОПОМОГОЮ АЛГОРИТМІВ МАШИННОГО НАВЧАННЯ

Постановка задачі. Електронна пошта залишається одним із основних каналів комунікації в інформаційних системах, що робить її привабливою ціллю для фішингових атак. Фішингові електронні листи використовуються для викрадення конфіденційних даних, компрометації облікових записів та поширення шкідливого програмного забезпечення. Сучасні фішингові повідомлення часто маскуються під легітимні листи та використовують складні текстові конструкції, що значно ускладнює їх виявлення традиційними методами фільтрації. Більшість існуючих систем виявлення фішингу ґрунтуються на сигнатурних або rule-based підходах, які не забезпечують достатньої адаптивності до нових та модифікованих атак. Такі системи мають обмежену здатність до узагальнення та часто демонструють підвищену кількість хибнопозитивних і хибнонегативних спрацьовувань. Це негативно впливає як на рівень безпеки, так і на зручність користування електронною поштою. У зв'язку з цим актуальним є наукове завдання розробки ефективного методу виявлення фішингових електронних листів на основі алгоритмів машинного навчання, здатного адаптуватися до змінних шаблонів атак та забезпечувати стабільну якість

класифікації. Додатковою проблемою є інтеграція таких методів у реальні системи електронної пошти, що потребує стандартизованого механізму взаємодії між моделями аналізу та системними компонентами.

Мета дослідження. Підвищення ефективності виявлення фішингових електронних листів шляхом розробки та дослідження методу, що поєднує алгоритми машинного навчання з використанням уніфікованого інтерфейсу АМСІ. Досліджується вплив застосування різних моделей машинного навчання на якість класифікації електронних листів за стандартними метриками оцінювання, зокрема Accuracy, Precision, Recall та F1-score.

Основна ідея запропонованого підходу полягає у використанні алгоритмів машинного навчання для аналізу текстового вмісту електронних листів з подальшою інтеграцією результатів класифікації через інтерфейс АМСІ. Такий підхід дозволяє стандартизувати процес аналізу електронних повідомлень, зменшити залежність від статичних правил та забезпечити можливість масштабування і впровадження методу в реальні системи електронної пошти.

Результати дослідження. У ході дослідження було встановлено, що при використанні традиційних методів фільтрації електронної пошти ефективність виявлення фішингових листів знижується в умовах змінних шаблонів атак та складних текстових структур. Це зумовлено обмеженими можливостями таких методів щодо аналізу контексту та прихованих закономірностей у даних.

Для вирішення зазначеної проблеми було застосовано низку алгоритмів машинного навчання, зокрема Multinomial Naive Bayes, Logistic Regression, Random Forest, Decision Trees, k-nearest neighbors, Support Vector Classifier та Multilayer Perceptron. Оцінювання ефективності моделей проводилося з використанням стандартних метрик класифікації, що дозволило комплексно проаналізувати якість прогнозування та порівняти результати між різними підходами. Результати експериментального дослідження показали, що використання підходу ML у поєднанні з уніфікованим інтерфейсом АМСІ забезпечує більш збалансовані показники точності та повноти виявлення фішингових електронних листів. Зокрема, спостерігалось зменшення кількості хибнопозитивних і хибнонегативних спрацьовувань у порівнянні з окремими класичними моделями машинного навчання, що підтверджується значеннями метрик Precision, Recall та F1-score. Інтеграція моделей машинного навчання з АМСІ дозволила стандартизувати процес передачі результатів аналізу та спростити впровадження розробленого методу у системи електронної пошти без необхідності зміни їх базової архітектури. Крім того, результати дослідження підтверджують, що запропонований підхід є перспективним рішенням для підвищення рівня безпеки електронної пошти та може бути використаний при розробці сучасних систем захисту від фішингових атак.

Список використаних джерел

1. Alshamrani A., Myneni S., Chowdhary A., Huang D.
A survey on phishing attacks and countermeasures for email security // IEEE Access. – 2024. – Vol. 12. – P. 112345–112370.

2. Nguyen T., Kang J., Kim H.
Phishing email detection using transformer-based language models // Expert Systems with Applications. – 2024. – Vol. 236. – Article 121435.
3. Zhang Y., Liu Q., Wang S.
Context-aware phishing email detection based on BERT and attention mechanisms // Knowledge-Based Systems. – 2024. – Vol. 284. – Article 111320.
4. Kumar S., Jain A., Bansal A.
Machine learning approaches for phishing email detection: A comparative study // Applied Soft Computing. – 2024. – Vol. 151. – Article 111095.
5. Aljabri M., Alqahtani A.
Deep learning-based phishing email detection in enterprise environments // Computers & Security. – 2024. – Vol. 136. – Article 103438.
6. Chen L., Zhao Y., Sun J.
Hybrid phishing detection using NLP and ensemble learning techniques // Information Sciences. – 2024. – Vol. 646. – P. 119–134.
1. Microsoft Corporation.

Чичкарьов Є.А., професор ДУІКТ
Радковський Д.О., студент ДУІКТ

СИСТЕМА МОНІТОРИНГУ ПСИХОЕМОЦІЙНОГО СТАНУ ЛЮДИНИ З ВИКОРИСТАННЯМ ШТУЧНОГО ІНТЕЛЕКТУ

Анотація. У роботі розглядається концепція створення системи моніторингу психоемоційного стану людини з використанням технологій штучного інтелекту. Основна мета — забезпечити автоматизоване розпізнавання емоцій на основі фізіологічних, поведінкових та візуальних показників, що дозволить своєчасно виявляти ознаки стресу, емоційного вигорання або депресивних станів.

Вступ. Сучасні умови життя, високий рівень інформаційного навантаження та постійні стресові фактори суттєво впливають на психоемоційний стан людини. Тому розроблення систем, здатних аналізувати стан користувача в реальному часі, набуває особливого значення. Застосування методів штучного інтелекту відкриває можливості для більш точної оцінки емоцій на основі даних із різних джерел — відео, аудіо, біометричних сенсорів тощо.

Основна частина.

Запропонована система складається з кількох модулів:

1. *Збір даних.* Система отримує дані з камери (міміка, рухи очей), мікрофона (тон голосу) та сенсорів (пульс, рівень насичення крові киснем, варіабельність серцевого ритму).

2. *Попередня обробка.* Отримані сигнали проходять фільтрацію, нормалізацію та перетворення у вхідні параметри для моделей машинного навчання.

3. *Модуль аналізу.* На цьому етапі використовується нейронна мережа, навчена на відкритих наборах даних (наприклад, AffectNet або DEAP), для класифікації емоцій за категоріями — радість, смуток, злість, страх, нейтральний стан тощо.

4. *Візуалізація результатів.* Користувач отримує зручний інтерфейс, який відображає динаміку психоемоційного стану у вигляді графіків або індикаторів.

Система може бути реалізована як окремий застосунок або вбудований модуль у фітнес-трекери чи розумні гаджети. Особливу увагу приділено безпеці персональних даних: усі вимірювання можуть оброблятися локально без передачі інформації у хмару.

Результати та перспективи. Очікується, що впровадження такої системи сприятиме покращенню якості життя користувачів, своєчасному виявленню психоемоційних порушень і підвищенню ефективності персонального здоров'я-моніторингу. Надалі планується оптимізація моделей машинного навчання для роботи в режимі реального часу та підвищення точності розпізнавання емоцій за індивідуальними особливостями користувача.

Висновки. Розробка системи моніторингу психоемоційного стану людини з використанням штучного інтелекту є актуальним напрямом у галузі інформаційних технологій і біоінженерії. Поєднання алгоритмів комп'ютерного зору, аналізу звуку та біосигналів дозволяє створити ефективний інструмент для контролю психологічного стану та профілактики емоційних розладів.

Список використаних джерел:

1. Picard R. W. *Affective Computing*. MIT Press, 1997.
2. Koelstra S. et al. *DEAP: A database for emotion analysis using physiological signals*. IEEE Trans. Affective Computing, 2012.
3. Poria S. et al. *Multimodal Sentiment Analysis: A Review*. Information Fusion, 2017.
4. Li X., Zhao G. *Automatic facial expression recognition with deep learning: A survey*. Pattern Recognition, 2020.
5. Huang Y., Yang H. *Emotion recognition using physiological signals and deep learning methods*. Sensors, 2021.

Вишнівський В.В., д.т.н., професор ДУІКТ
Мороз М.В., аспірант ДУІКТ

ПОРІВНЯЛЬНИЙ АНАЛІЗ МОДЕЛЕЙ U-NET, DPT ТА MIDAS ДЛЯ ОЦІНКИ ГЛИБИНИ У SLAM-СИСТЕМАХ

Анотація. Проведено порівняльний аналіз точності та швидкодії моделей U-Net, DPT та MiDaS для задачі оцінки глибини сегментованого зображення у системах SLAM за допомогою модульного поєднання з існуючим алгоритмом ORB-SLAM3. Результати показали, що DPT забезпечує найвищу точність оцінювання, тоді як MiDaS є більш придатною для задач реального часу з обмеженими ресурсами. Було запропоновано критерії вибору моделі на основі вимог до точності обчислення та ресурсних можливостей системи.

Ключові слова: SLAM; U-Net; DPT; MiDaS; оцінка глибини; машинне навчання.

Вступ

Точна оцінка глибини отриманого зображення є важливою підзадачею у системах, що реалізують одночасну оцінку положення та побудову мапи середовища (SLAM). Конвенційні методи використовують стереозображення, або апаратне забезпечення як LiDAR, однак їх ефективність суттєво знижується у складних динамічних умовах або за наявності лише монокулярної камери.

Нейромережеві моделі, такі як U-Net, DPT та MiDaS, показують здатність до відновлення геометрії та глибини з монокулярного RGB-зображення [1–3]. Проте задача інтеграція таких моделей у SLAM потребує компромісу між точністю, швидкістю та стабільністю роботи.

Мета роботи – провести порівняння трьох вибраних моделей оцінювання глибини для визначення їх придатності до використання у візуальних монокулярних системах SLAM.

Основна частина

Для дослідження було використано існуючу модульну систему ORB-SLAM3, у якій було інтегровано модуль оцінки глибини. Кожна модель отримувала RGB зображення та генерувала оцінку глибини для кожного конкретного зображення. Експерименти проводилися на наборі даних Low Viewpoint Forest Depth [4]. Метрики оцінювання включали середню абсолютну похибку (MAE), а також швидкість обробки кадрових зображень (FPS).

Таблиця 1.

Модель	MAE	FPS
DPT	0.093	12
U-Net	0.128	30
MiDaS	0.106	22

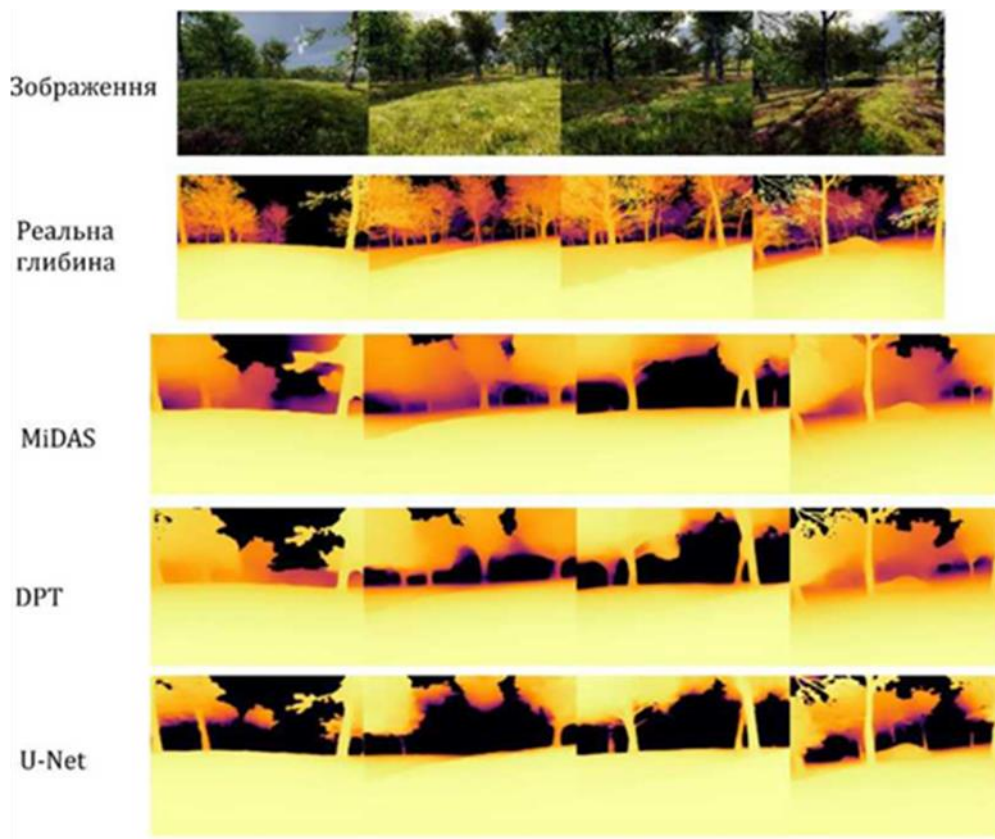


Рисунок 1. Візуальне порівняння згенерованих мап глибини та реального зображення.

Еспериментальні результати показали (таблиця 1) (рис. 1), що модель DPT досягла найменшої середньої абсолютної похибки, але при цьому використовує значний обчислювальний ресурс. Модель U-Net забезпечує задовільну точність із високим ступенем швидкодії. Модель MiDaS продемонструвала компроміс між точністю і швидкістю.

Висновки

Було проведено порівняльний аналіз моделей U-Net, DPT та MiDaS, для задачі оцінки глибини зображення у системах SLAM. Вибір конкретної моделі повинен залежати від цільового застосування, як от модель DPT для задач високоточної оцінки глибини зображень із значними обчислювальними ресурсами, MiDaS - для динамічних середовищ та систем реального часу, U-Net – для автономних роботизованих платформ із обмеженими ресурсами. Подальші дослідження можуть включати об'єднання переваг конкретних моделей у гібридному рішенні та інтеграцію з різними типами систем SLAM.

Список використаних джерел

- [1] Ronneberger, O., Fischer, P., & Brox, T. (2015). U-Net: Convolutional Networks for Biomedical Image Segmentation. MICCAI, 234–241.
- [2] Ranftl, R., Bochkovskiy, A., & Koltun, V. (2021). Vision Transformers for Dense Prediction. ICCV, 12159–12168.

[3] Ranftl, R., Lasinger, K., Hafner, D., Schindler, K., & Koltun, V. (2020). Towards Robust Monocular Depth Estimation: Mixing Datasets for Zero-shot Cross-dataset Transfer. arXiv:1907.01341.

[4] Niu, C., Tarapore, D., & Zauner, K.-P. (2020). Low-viewpoint forest depth dataset for sparse rover swarms. arXiv:2003.04359.

Moroz M.V., Vyshnivskyi V.V.

Comparative analysis of U-Net, DPT and MiDaS models for depth estimation in SLAM systems Abstract.

A comparative analysis of the accuracy and performance of the U-Net, DPT, and MiDaS models for the task of segmented image depth estimation in SLAM systems using a modular combination with the existing ORB-SLAM3 algorithm was conducted. The results showed that DPT provides the highest estimation accuracy, while MiDaS is more suitable for real-time tasks with limited resources. The model selection criteria were proposed based on the requirements for computational accuracy and system resource capabilities.

Keywords: SLAM; U-Net; DPT; MiDaS; depth estimation; machine learning.

Вишнівський В.В., д.т.н., професор ДУІКТ
Мороз М.В., аспірант ДУІКТ

ПОЄДНАННЯ ГЛИБИННИХ МАП І СЕМАНТИЧНОГО РОЗПІЗНАВАННЯ ДЛЯ ПІДВИЩЕННЯ ТОЧНОСТІ ЛОКАЛІЗАЦІЇ У ЗАДАЧАХ SLAM

Постановка задачі. Задача SLAM (Simultaneous Localization and Mapping) полягає в одночасній локалізації безпілотної платформи та створенні деталізованої мапи середовища, що є критичною функцією для багатьох безпілотних платформ, які оперують у невідомих умовах. У динамічних середовищах без доступу до GPS, традиційні методи SLAM швидко накопичують похибку через обмежену кількість надійних орієнтирів. Хоча отримані результати можна обмежено використовувати для завдань автономної навігації, системі не вистачає семантичної інформації про середовище. Також, наявність рухомих об'єктів, людей, техніки та іншого вносить додаткову невизначеність, оскільки без фільтрації таких динамічних елементів, системи SLAM часто дають нестабільні результати [1]. Водночас збільшення обсягу відстежуваної інформації підвищує навантаження на обчислювальні ресурси, що потребує особливих рішень щодо оптимізації. Тому постає наукова задача впровадження семантичної та глибинної оцінки спостережуваних об'єктів у системах, що використовують SLAM.

Мета дослідження. Метою роботи є дослідження методів, які поєднують інформацію з глибинних мап зображення і семантичної інформації отриманої від згорткових нейронних мереж (CNN) для покращення точності локалізації та мапи SLAM. Передбачається застосування існуючих нейромережових алгоритмів для

отримання семантичних масок розпізнаваних об'єктів, а також проєкція отриманих масок у простір за допомогою глибинної карти зображення.

Результати дослідження. Було використано існуючу архітектуру ORB-SLAM3 в комбінації з семантичною сегментацією зображення, інтегрованою у єдине рішення. Семантичні маски, отримані нейромережею Mask R-CNN та Vision Transformer, проєктувалися у вихідну мапу SLAM за допомогою глибинних мап зображення [2]. Це дало змогу ідентифікувати простір, зайнятий тими чи іншим об'єктом, і враховувати цю інформацію при оптимізації мапи середовища. Зокрема, виявлено, що при фільтрації рухомих об'єктів суттєво зменшується накопичена помилка локалізації [3]. Система порівнює отримані ознаки між кадрами, щоб забезпечити стабільність розпізнавання в умовах руху камери даних. Експерименти на відомих датасетах (TUM RGB-D) показали зменшення похибок позиціонування. Зокрема, підхід із поліпшеною сегментацією забезпечив зниження середньої помилки (RMSE) на 17.1% порівняно з класичним методом DynaSLAM [4]. Отримані результати узгоджуються з висновками інших досліджень: семантична сегментація зображень в поєднанні з глибинними мапами зображення показує істотне покращує стабільності SLAM-систем у динамічних та складних умовах [2][3][4].

Висновки. Комбінація інформації з глибинних мап зображень і семантичної інформації підтверджує свою ефективність для покращення точності локалізації та мапи в задачах SLAM. Аналіз результатів показав, що гібридні системи, які об'єднують декілька типів інформаційних джерел і згорткових нейронних мереж (CNN), забезпечують найкращі результати та стабільну поведінку в динамічних середовищах. Семантична інформація дозволяє фільтрувати вплив рухомих об'єктів і виділяти повторювані орієнтири, що сприяє зниженню покращенню точності і швидкодії системи. Подальші дослідження можуть бути спрямовані на оптимізацію нейронних моделей і розробку адаптивних архітектур для роботи в реальному часі.

Список використаних джерел

1. Усов О.С. (2025). Методи візуальної SLAM для навігації у внутрішньому середовищі без доступу до глобального позиціонування, Вчені записки ТНУ імені В.І. Вернадського. Серія: Технічні науки, т.36 (75), №3, 411–418. DOI: 10.32782/2663-5941/2025.3.2/54
2. Martín F., González F., Guerrero J.M., Fernández M., Ginés J., 2021, Semantic 3D Mapping from Deep Image Segmentation, Applied Sciences, 11(4):1953. DOI: 10.3390/app11041953
3. Beghdadi A., Mallem M., Beji L., 2023, D2SLAM: Semantic Visual SLAM based on the Depth-Related Influence on Object Interactions for Dynamic Environments, arXiv preprint arXiv:2210.08647
4. Chen M., Guo H., Qian R., Gong G., Cheng H., 2024, Visual simultaneous localization and mapping (vSLAM) algorithm based on improved Vision Transformer semantic segmentation in dynamic scenes, Mechanical Sciences, 15, 1–16. DOI: 10.5194/ms-15-1-2024.

Вишнівський В.В., д.т.н., професор ДУІКТ
Мороз М.В., аспірант ДУІКТ

ВИКОРИСТАННЯ НЕЙРОМЕРЕЖЕВИХ ОЗНАКОВИХ ДЕСКРИПТОРІВ У ВІЗУАЛЬНИХ ЗАДАЧАХ SLAM

Постановка задачі. У конвенційних системах, що вирішують задачу SLAM, а саме одночасну побудову мапи та локалізацію роботизованої платформи у середовищі, ефективність не в останню чергу залежить від точності сегментації ключових точок і ознакових дескрипторів. Класичні алгоритми SIFT, SURF, ORB тощо забезпечують високу швидкість та інваріантність до повороту, проте їх точність значно падає за умов обмеженої текстуризації, низької освітленості чи наявності динамічних об'єктів у середовищі [1] [2]. Такі методи використовують інтенсивність зображення для виділення ключових точок, що при певних умовах, дає мало повторюваних точок у порожніх областях і часто генерують помилкові збіги при різкій зміні світла, розмитті чи слабких контурах на зображенні. Така поведінка часто призводить до втрати трекінгу об'єкту і накопичення похибки локалізації в динамічному середовищі.

Мета дослідження. Метою роботи є оцінити можливості сучасних нейромережових дескрипторів ознак, таких як SuperPoint, R2D2, LF-Net, D2-Net для поліпшення відповідності ознак у візуальному SLAM та загальному підвищенні точності локалізації.

Результати дослідження. У ході дослідження було з'ясовано, що нейромережові алгоритми для сегментації ознакових дескрипторів здатні навчатися на великих множинах даних і генерувати інваріантні до геометричних та фотометричних перетворень ознаки [4] [5]. Наприклад, мережа SuperPoint одночасно виділяє ключові точки і 256-вимірні дескриптори, показуючи високу повторюваність та точність співставлення при зміні середовища [3]. У свою чергу, R2D2 також показує високий ступінь надійності та повторюваності точок, а LF-Net і D2-Net націлені на відновлення згенерованих раніше стабільних дескрипторів точок [4]. Результати інтеграції у існуючу систему на базі ORB-SLAM3, свідчать про значні переваги нейромережових дескрипторів перед класичними методами візуального SLAM. Так, у SuperPoint на датасеті KITTI середня трансляційна помилка знизилась зі 4.15 до 1.45%, порівняно з базовим варіантом. Це забезпечує рівень точності, недосяжний для конвенційного методу ORB [3]. У динамічних RGB-D сценах (датасет TUM) система з R2D2-дескрипторами досягла 70% повторюваності ознак проти 40% у класичного ORB-SLAM3 [4]. При цьому середня квадратична похибка знизилась приблизно на 90% порівняно з базовою реалізацією. Загалом впровадження таких дескрипторів суттєво покращує трекінг об'єктів, система витримує зміни динамічного середовища, у той час як конвенційні методи генерують значну похибку з часом.

Висновки. Було підтверджено, що нейромережові алгоритми для сегментації ключових точок і ознакових дескрипторів ефективніше сегментують інформацію в складних динамічних середовищах, що підвищує точність локалізації порівняно

з класичними алгоритмами. Проте вони мають деякі обмеження, як от навантаження на апаратні ресурси та потребу у великій вибірці даних для навчання. Інтеграція таких алгоритмів з іншими компонентами SLAM, як от глибинними мапами зображення чи семантичною інформацією, може також значно підвищити точність. Наявність семантичної інформації дозволить фільтрувати нерелевантні об'єкти.

Список використаних джерел

1. Peng Q., Xiang Z., Fan Y., Zhao T., Zhao X. (2022). RWT-SLAM: Robust Visual SLAM for Highly Weak-textured Environments. ArXiv preprint arXiv:2207.03539.
2. Xie C., Liu Q., Chen B., Hao Z. (2024). Evaluation and Analysis of Feature Point Detection Methods based on vSLAM Systems. Image and Vision Computing, 146:105015. DOI:10.1016/j.imavis.2024.105015.
3. DeTone D., Malisiewicz T., Rabinovich A. (2018). SuperPoint: Self-Supervised Interest Point Detection and Description. Proc. IEEE CVPR Workshops, CVPRW 2018.
4. Yang W., Chao H., Zhang Y., Tan S. (2025). 2HR-Net VSLAM: Robust visual SLAM based on dual high-reliability feature matching in dynamic environments. PLoS ONE 20(7):e0328052. DOI:10.1371/journal.pone.0328052.
5. Tareen S.A.K., Tareen F.K. (2025). DeepDetect: Learning All-in-One Dense Keypoints. ArXiv preprint arXiv:2510.17422.

Чичур А.І., к.т.н., доцент ДУІКТ
Подгорбунський В.В., студент ДУІКТ

ОЦІНКА ЛІНГВІСТИЧНОЇ НОВИЗНИ ТЕКСТІВ, ЗГЕНЕРОВАНИХ ВЕЛИКИМИ МОВНИМИ МОДЕЛЯМИ

Сучасні великі мовні моделі демонструють вражаючі результати в генерації текстів, однак залишається відкритим питання: чи дійсно ці моделі розуміють мовні закономірності, чи просто відтворюють фрагменти з навчальних даних? Традиційні метрики оцінки якості тексту (зв'язність, граматика, стилістична коректність) можуть вводити в оману, якщо модель переважно копіює матеріал із тренувальної вибірки замість справжнього узагальнення лінгвістичних правил.

Постановка задачі

Розробити методологію систематичної оцінки новизни текстів, створених нейронними мовними моделями, та з'ясувати, якою мірою моделі спираються на копіювання навчальних даних порівняно з продуктивною генерацією нових структур.

Мета дослідження

Проаналізувати механізми генерації тексту в сучасних мовних моделях через призму лінгвістичної новизни на різних рівнях: лексичному, синтаксичному

та морфологічному. Визначити сильні та слабкі сторони моделей у створенні справді нового змісту.

Результати дослідження

Дослідження проводилось на основі інструментарію RAVEN (RAting VErbal Novelty), який дозволяє багаторівнево оцінити новизну згенерованих текстів. Експерименти охопили чотири архітектури: LSTM, Transformer, Transformer-XL та GPT-2, із використанням корпусів Wikitext-103 і WebText.

Виявлено цікаву закономірність: для коротких фраз (до 6 слів) моделі виявляються досить консервативними. Наприклад, базова вибірка містить приблизно 6% нових пар слів, тоді як моделі генерують лише 2-3% таких пар. Це говорить про те, що моделі «перестраховуються» при створенні локальних структур.

Натомість для довших послідовностей (понад 6 слів) спостерігається протилежна картина – моделі часто перевищують базові показники новизни. Однак тут виникає несподіваний феномен «надкопіювання»: іноді моделі дослівно відтворюють фрагменти з навчальних даних обсягом понад 100, а в окремих випадках – навіть понад 1000 слів. Особливо це характерно для GPT-2 при роботі з текстами, що повторюються в корпусі.

Що стосується загальної структури речень, результати виглядають оптимістично: близько 83% речень, згенерованих моделлю Transformer-XL, мають синтаксичну будову, яка не зустрічається в навчальних даних. Це свідчить про те, що моделі справді здатні комбінувати синтаксичні елементи по-новому.

Проте на рівні окремих синтаксичних зв'язків між словами картина інша. GPT-2 створює лише 3% нових типів залежностей проти 5% у природних текстах. Знову ж таки спостерігається консерватизм на локальному рівні при новаторстві на глобальному.

GPT-2 продемонструвала непогані результати в морфології: 96% нових слів є морфологічно правильними (для порівняння, у природних текстах цей показник становить 99%). Модель добре справляється з граматичними правилами – наприклад, правильно узгоджує дієслова з новими іменниками у 25 випадках із 26, навіть коли в реченні є відволікаючі елементи.

Точність вибору між закінченнями -s/-es для множини складає 97%, а для присвійних форм -'s/-' – навіть 99%. Однак виявилась систематична проблема з акронімами: лише 28% згенерованих скорочень адекватно відповідають повним формам (проти 81% у природних текстах).

Найслабшим місцем виявилась семантика. Тільки 21% нових слів є семантично чіткими у своєму контексті (порівняно з 33% у базовій вибірці). Ще 59% можна вважати потенційно прийнятними, але не беззаперечними. Особливо помітні проблеми виникають при роботі з числами: неправильні конверсії одиниць виміру, фізично неможливі величини, непослідовні обмінні курси. Трапляються також самосуперечливі твердження та порушення базової логіки.

Вдалось ідентифікувати два основних механізми, які GPT-2 використовує для створення нових форм:

Композиційний підхід – модель додає префікси та суфікси до основ слів, які раніше не зустрічалися в навчальних даних, але присутні в поточному контексті. Це нагадує те, як люди інтуїтивно створюють нові слова за аналогією з відомими правилами.

Аналогічне узагальнення – заміна схожих компонентів слів. Характерний приклад: модель бачить слово «torero» і систематично замінює частину «tore» на різні форми дієслова «tear» («tear», «torn», «tearing», «tears»), хоча такі варіанти семантично не мають сенсу.

Експерименти показали чіткий компроміс між новизною та якістю: методи декодування, що підвищують новизну (вищі значення параметрів k , p , temperature), водночас погіршують загальну якість тексту. Метод top-40 sampling забезпечує високу якість, але знижує новизну для малих структур.

Цікаво, що розмір моделі не має послідовного впливу на новизну. GPT-2 XL виявилась найбільш новаторською серед версій GPT-2, але GPT-2 Medium перевершила GPT-2 Large за цим показником.

Моделі з механізмом рекурентності (LSTM, Transformer-XL) менш схильні до новизни на рівні коротких фраз через ефект недавності (recency bias). GPT-2 демонструє певне «розуміння» концепції нового слова – схильна брати такі слова в лапки (1,9% нових слів у лапках проти 0,16% для всіх слів загалом). Також помітна здатність моделі до впорядкування: вона може продовжувати числові послідовності та алфавітні списки.

Висновки та перспективи

Дослідження спростовує твердження про те, що мовні моделі є просто «стохастичними папугами». Моделі дійсно здатні до продуктивної генерації нових структур, особливо на рівні загальної організації речень та морфологічних комбінацій. Водночас вони демонструють надмірний консерватизм на локальному рівні – створюють значно менше нових пар і трійок слів порівняно з природними текстами.

Результати підтверджують, що моделі засвоїли багато лінгвістичних абстракцій: синтаксичні структури, правила утворення слівформ, узгодження граматичних категорій. Проте виявлено і систематичні обмеження – труднощі з акронімами, числовими значеннями та підтриманням семантичної послідовності.

Практичне значення роботи полягає у розробці методології перевірки новизни як обов'язкового кроку при оцінці можливостей мовних моделей. Якщо згенерований текст копіюється з навчальних даних, він не може слугувати доказом узагальнювальних здібностей моделі. Отримані результати можуть бути корисними для вдосконалення методів дедуплікації даних, оптимізації стратегій декодування та розробки архітектур із вбудованими композиційними механізмами.

Перспективним напрямом подальших досліджень є поглиблений аналіз семантичної новизни та вивчення зв'язку між явним плануванням змісту і здатністю моделей генерувати справді оригінальний текст.

Список використаних джерел

1. McCoy R. T., Smolensky P., Linzen T., Gao J., Celikyilmaz A. How Much Do Language Models Copy from Their Training Data? Evaluating Linguistic Novelty in Text Generation Using RAVEN // Transactions of the Association for Computational Linguistics. 2023. Vol. 11. P. 652–670. URL: https://direct.mit.edu/tacl/article-pdf/doi/10.1162/tacl_a_00567/2141008/tacl_a_00567.pdf
2. Bender E. M., Koller A. Climbing towards NLU: On meaning, form, and understanding in the age of data // Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics. 2020. P. 5185–5198.
3. Carlini N., Tramer F., Wallace E. et al. Extracting training data from large language models // 30th USENIX Security Symposium. 2021. P. 2633–2650.
4. Lee K., Ippolito D., Nystrom A. et al. Deduplicating training data makes language models better // Proceedings of the 60th Annual Meeting of the ACL. 2022. P. 8424–8445.

Чичур А.І., к.т.н., доцент ДУІКТ
 Подгорбунський В.В., студент ДУІКТ

АНАЛІЗ СИСТЕМ КОНТРОЛЮ ТА МОНІТОРИНГУ LLM

Сучасні великі мовні моделі (англ. Large Language Model, LLM) потребують контролю та моніторингу якості їх роботи, оскільки звичайні інструменти моніторингу не спроможні в повній мірі задовольнити ці потреби.

Необхідно оперативно моніторити як саме модель поводить себе в різних ситуаціях, із складними ланцюжками промтів (англ. prompts), проміжними викликами та контекстами, відстежувати що її «спонукає» помилятися і як зміни між різними варіантами налаштувань впливають на ефективність її роботи.

Постановка задачі

Проаналізувати ринок сучасних ІТ систем і визначити найпопулярніші системи моніторингу роботи LLM.

Мета дослідження

Оцінити сучасний стан ринку ІТ систем в галузі систем моніторингу роботи LLM. Оцінити перспективи розвитку, ефективність застосування, а також визначити їх переваги та недоліки.

Результати дослідження

Визначено сильні та слабкі сторони застосування систем моніторингу роботи LLM на прикладі найбільш популярних сервісів: Langfuse та Comet. Проведено їх оцінку їх роботи для упорядковування та покращення роботи застосунків із LLM.

Langfuse – це платформа з відкритим кодом (англ. open source), що важливо для проєктів із підвищеними вимогами до конфіденційності або самостійного

розгортання і підтримки (у власній хмарі), а також із різними ліцензійними планами для хмарного розміщення.

Вона легко (з коробки) інтегрується із найпопулярнішими комерційними і некомерційними LLM застосунками та моделями: OpenAI, DeepSeek, Mistral, Ollama, OpenWebUI, тощо. Підтримує інтеграцію з no-code builders (переклад з англ. конструктори без коду), такі як: n8n, LobeChat, Flowise, тощо. Langfuse має близько 19 тисяч «зірок» на GitHub.

Вирізняється тим, що досить детально розглядає кожен запит до моделі: що було передано, як швидко опрацьовано, які проміжні кроки виконувалися. Моніторинг продуктивності LLM-додатку здійснюється за допомогою комплексних панелей інструментів метрик та API. Відстежує витрати (вартість запитів), затримки, використання токенів та показники якості для різних моделей, користувачів та періодів часу.

В Langfuse на високому рівні реалізовано моніторинг роботи з промтами: спільне керування їх версіями, миттєве розгортання/відкат у різних середовищах, підтримка шаблонів, змінних та A/B-тестування. Кешування на стороні клієнта для нульовий вплив на затримку/доступність. Langfuse дозволяє з легкістю ставити і проводити експерименти та поступово вдосконалювати налаштування, навіть у великих командах.

Сервіс Langfuse дозволяє виконувати онлайн/офлайн-оцінювання через інтерфейс користувача (експерименти із запитамі/моделями) та через SDK (Python, JS/TS), створювати набори даних (англ. datasets) для постійного покращення оцінок.

Також є можливість створювати анотації вручну, щоб надавати відгуки, виправлення та покращення до результатів LLM. Це допомагає в створенні високоякісних наборів даних та встановлення базової лінії для автоматизованих оцінок.

Comet – це також платформа з відкритим кодом, проте історично створювалась як платформа для MLOps (англ. Machine Learning & DevOps), тому її можливості виходять за межі лише LLM. Comet дозволяє зберігати різні моделі, експерименти, метрики й артефакти, а також відстежувати їх поведінку під час реального використання. Comet має близько трохи більше ніж 15 тисяч «зірок» на GitHub.

З появою окремого модуля для роботи з LLM Comet став зручнішим для проєктів, де важливо контролювати саме продуктивність у продуктовому середовищі: наприклад, відловлювати ознаки деградації моделі або проблему з дрейфом даних. Через більш «корпоративну» спрямованість Comet краще підходить для великих команд або довготривалих проєктів, де потрібен порядок у великій кількості експериментів.

Загалом Langfuse дає гнучкіший контроль над конкретними LLM й підходить там, де потрібно багато експериментувати з промтами й архітектурою агентів. Comet же доцільніший, коли мова йде про комплексні ML-процеси або проєкти, що вже працюють у промисловому середовищі.

Висновки та перспективи

Отже, сучасний ринок ІТ систем для моніторингу роботи LLM та інших моделей штучного інтелекту невпинно розвивається. Найпопулярніші сервіси Langfuse та Comet є лідерами цього сегменту ринку і поставляють свої продукти як платформа з відкритим кодом, завдяки чому мають швидко ростучу спільку інтузіастів.

Ці платформи не є прямими конкурентами, а скоріше за все, доповнюють одна одну, орієнтуючись на різні потреби та вирішення проблем ринку.

Langfuse більше підходить для розробників, які хочуть детально бачити відстежувати всі процеси роботи LLM-застосунків і мати можливість швидко змінювати окремі елементи без надлишкових обмежень. Проводити експерименти та накопичувати інформацію.

Comet орієнтований на команди, яким важливо вести облік багаточисленних експериментів, контролювати моделі в реальному середовищі та підтримувати цілісну ML-інфраструктуру.

Перспективним напрямом може бути й комбіноване їх використання, коли трасування та робота з промтами здійснюється через Langfuse, а загальний життєвий цикл моделей контролює Comet.

Список використаних джерел

1. Langfuse. *Langfuse*. URL: <https://langfuse.com/> (дата звернення: 21.11.2025).
2. Comet - monitor LLM applications. *Comet*. URL: <https://www.comet.com/site/> (дата звернення: 21.11.2025).
3. Reliable Weak-to-Strong Monitoring of LLM Agents : Paperprint / N. Kale та ін. Massachusetts : Carnegie Mellon University, Massachusetts Institute of Technology, 2025. 48 с.
URL: https://www.researchgate.net/publication/395034628_Reliable_Weak-to-Strong_Monitoring_of_LLM_Agents , <https://arxiv.org/abs/2508.19461> (дата звернення: 21.11.2025).

Чичур А.І., к.т.н., доцент ДУІКТ

Метелюк А.В., студент ДУІКТ

ПОРІВНЯННЯ МЕТОДІВ ПЕРЕХОДУ ВІД МОНОЛІТНОЇ ДО МІКРОСЕРВІСНОЇ АРХІТЕКТУРИ ДЛЯ СИСТЕМИ МОНІТОРИНГУ ТА АЛЕРТИНГУ РЕКЛАМНИХ КАМПАНІЙ

Монолітні архітектури часто обмежують можливості масштабування, доцільно переходити до мікросервісів. Це дозволяє розділити складну систему на самостійні компоненти, що дозволяє локалізувати помилки та значно підвищує загальну відмовостійкість системи.

Постановка задачі

Проаналізувати методи міграції від монолітної до мікросервісної архітектури та оцінити їх ефективність застосування для систем моніторингу. Виокремити оптимальні архітектурні патерни для модуля алертингу (англ. alerting) та збору даних.

Мета дослідження

Визначити найбільш ефективні стратегії та патерни проектування мікросервісної архітектури для системи моніторингу рекламних кампаній.

Результати дослідження

Визначено відмінності монолітної та мікросервісної архітектури для систем моніторингу, визначено їх переваги та недоліки. Основними перевагами переходу до мікросервісної архітектури є покращення масштабованості та спрощення технічного обслуговування. Для систем моніторингу рекламних розміщень, де критичним є відмовостійкість, важливим фактором є можливість оновлювати окремі модулі без повної зупинки системи.

Найбільш поширеною стратегією декомпозиції системи є поділ на основі бізнес-можливостей. Такий підхід є переважаючим (майже 50,9% практик). Для системи моніторингу це означає виділення окремих сервісів для збору статистики, аналізу ефективності та відправки сповіщень (англ. alert). При цьому більшість систем моніторингу (майже 58%) використовують архітектурний підхід «одна база даних на сервіс» для забезпечення автономності.

Проведено аналіз методів міграції з монолітної на мікросервісну архітектуру для систем моніторингу. Порівняльний аналіз патернів (англ. patterns), що лежать в основі методів міграції, розглядав патерни: Strangler Fig, Parallel Run та Decorating Collaborator.

Strangler Fig Pattern або Strangler Pattern (переклад з англ. «Фікус-душитель») демонструє низький ризик помилок (8%) та незначну деградацію продуктивності (4%), що робить його оптимальним для поступової заміни функціоналу моніторингу без зупинки системи. Ефективність цього методу оцінюється коефіцієнтом 0,6 і вказує на баланс між надійністю та часом реалізації. Цей термін сформульований Мартіном Фаулером, побачивши особливий вид фікусів, що оселяється на верхніх гілках дерев, обволікає його та використовує як тимчасову опору й джерело поживних речовин. Згодом, фікус спускається до землі і пускає коріння, а дерево гине. На той момент фікус стає самостійним.

Parallel Run (переклад з англ. «Паралельний запуск») характеризується найвищою надійністю для критичних модулів, таких як фінансовий моніторинг рекламних бюджетів, оскільки дозволяє порівнювати результати роботи старої та нової систем. Однак, цей метод, через високу складність тестування та подвійне навантаження на інфраструктуру, має найнижчий показник коефіцієнту ефективності – 0,27.

Decorating Collaborator (переклад з англ. «Співробітник з декорування») показав найвищий коефіцієнт ефективності – 11,62. Все завдяки швидкій реалізації та мінімальному втручанню в код моноліту, що може бути використано для додавання нових функцій алертингу без зміни основної логіки збору даних.

Для оцінки доцільності вибору методу міграції варто скористатись оцінкою методу її ефективності:

$$E = \frac{1}{T \times S \times R \times \left(1 + \frac{P}{100}\right)},$$

де, E – оцінка ефективності методу міграції; T – час, що необхідно витратити на міграцію, год; S – складність тестування (кількість записів або файлів); R – ризик помилок; та P – частка помилок, %.

Для успішного функціонування мікросервісної архітектури для системи моніторингу важливо також впровадити автоматизовані механізми спостереження. Їх використовують 83% компаній, що дозволяє гарантувати своєчасну відправку сповіщень про критичні зміни в рекламних кампаніях.

Висновки та перспективи

Отже, перехід від монолітної до мікросервісної архітектури для системи моніторингу рекламних кампаній є дієвим способом підвищення їх ефективності. Це розширює можливості масштабування такої системи і підвищує її надійність.

Найефективнішим методом є Strangler Fig. Він найбільш збалансований, оскільки передбачає поступове заміщення функціоналу, тоді як для критично важливих вузлів рекомендується Parallel Run.

Перспективним напрямком є розробка механізмів підвищення відмовостійкості, що гарантуватимуть стабільну роботу при зростанні кількості користувачів та активних рекламних кампаній.

Список використаних джерел

1. Insights on Microservice Architecture Through the Eyes of Industry Practitioners / V. L. Nogueira et al. 2024 IEEE International Conference on Software Maintenance and Evolution (ICSME), Flagstaff, AZ, USA, 6–11 October 2024. 2024. P. 765–777. URL: <https://doi.org/10.1109/icsme58944.2024.00080> (date of access: 19.11.2025).
2. Kornaha Y., Hubariev O. Improving the effectiveness of monolith architecture to microservices migration using existing migration methods. Information, Computing and Intelligent systems. 2024. No. 5. P. 159–171. URL: <https://doi.org/10.20535/2786-8729.5.2024.316288> (date of access: 19.11.2025).
3. Tuusjärvi K., Kasurinen J., Hyrynsalmi S. Migrating a Legacy System to a Microservice Architecture. e-Informatica Software Engineering Journal. 2024. Vol. 18, no. 1. P. 240104. URL: <https://doi.org/10.37190/e-inf240104> (date of access: 19.11.2025).

Чичур А.І., к.т.н., доцент ДУІКТ
Брагінець Т.В., студент ДУІКТ

ПЕРЕВАГИ ТА НЕДОЛІКИ ВПРОВАДЖЕННЯ ТЕХНОЛОГІЙ ДОПОВНЕНОЇ РЕАЛЬНОСТІ В ТУРИСТИЧНІ ГАЛУЗІ

Сучасний ринок туристичних послуг активно розвивається та інтегрує передові цифрові технології, серед яких перспективним напрямком є технології доповненої реальності (англ. Augmented Reality, AR). Вони покликані збагатити туристичний досвід, покращуючи комунікацію та привабливість туристичних послуг.

Постановка задачі

Проаналізувати переваги та недоліки застосування AR технологій для туристичних послуг і визначити перспективи розвитку.

Мета дослідження

Виявити параметри для оцінки якості використання AR технологій в туристичній галузі, та фактори, що обмежують їх впровадження.

Результати дослідження

AR технології суттєво трансформують сучасну туристичну індустрію, впливаючи як на якість туристичного досвіду, так і на ефективність роботи туристичних підприємств. Дослідження показало, що AR системи здатні значно підвищити інформативність туристичних послуг, шляхом створення інтерактивів – накладання цифрового контенту на реальний простір. Такі інструменти спрощують навігацію користувачів, надають доступ до контексту про туристичні об'єкти та створюють нові форми залучення аудиторії.

Виявлено, що інтеграція AR сприяє підвищенню привабливості туристичних локацій, зокрема за рахунок створення віртуальних екскурсій, реконструкцій історичних подій та оновлених форматів взаємодії з культурною спадщиною. Крім того, AR технології довели свою ефективність у сфері маркетингу, оскільки інтерактивний контент стимулює інтерес потенційних відвідувачів і формує позитивний імідж туристичних операторів.

Разом з тим, виявлено низку обмежень, які стримують широкомасштабне впровадження доповненої реальності. Основними бар'єрами є високі технічні та фінансові вимоги до розробки AR-додатків, залежність їхньої роботи від продуктивності мобільних пристроїв, а також ризики, пов'язані з конфіденційністю та безпекою користувачів. Значна частина потенційних обмежень пов'язана із збиранням персональних даних, геолокації та авторським правом контенту, що вимагає суворого контролю над дотриманням міжнародних стандартів у сфері захисту інформації.

Крім цього, інформаційне перевантаження користувача, яке виникає внаслідок надмірної кількості цифрових об'єктів на екрані мобільного пристрою, є неабиякою перепорою.

Дослідження підкреслює необхідність ретельного проектування інтерфейсів AR систем, для забезпечення достатнього рівня зручності, інтуїтивності та безпеки під час їхнього використання. Також, мобільний пристрій є визначальним

у створенні інтерфейсу таких систем, і як результат накладає обмеження на нього. В перспективі, розвиток технологій сприятиме розробці більш доступних інтерфейсів для взаємодії AR технологій із користувачем, не тільки в туристичній галузі.

Загалом, результати дозволяють комплексно оцінити потенціал та обмеження AR у туристичній галузі, а також визначають напрямки для подальшого їх вдосконалення.

Висновки та перспективи

Технології доповненої реальності є одним з перспективних інструментів цифрової трансформації туристичної галузі, спрямований суттєво підвищити якість послуг і створювати унікальний користувацький досвід.

Реалізація повного потенціалу AR вимагає подолання технологічних бар'єрів, удосконалення законодавчих механізмів захисту даних та адаптації інфраструктури туристичних об'єктів.

Перспективами подальших досліджень є розробка методики впровадження AR, оптимізація інформаційних потоків для уникнення перевантаження користувача, а також аналіз економічної ефективності використання AR у різних сегментах туристичної індустрії.

Список використаних джерел

1. Ivanov, S., & Webster, C. Augmented Reality in Tourism: Concepts, Applications and Challenges. *Journal of Tourism Technologies*, 2019, №4(2), с. 45–58.
2. Rauschnabel, P. A., Felix, R., & Hinsch, C. Augmented Reality Marketing: What We Know, What We Don't Know, and Where to Go. *International Journal of Technology in Tourism*, 2020, №7(1), с. 11–29.
3. Tussyadiah, I., Jung, T. H., & tom Dieck, M. Embodied Experiences in Augmented Reality: Implications for Tourism Destination Engagement. *Tourism Management Perspectives*, 2021, №37, с. 100–115.

Товсточуб І. С. доктор філософії, доцент ДУІКТ
Ярохно Д.В., студент ДУІКТ

ЕВОЛЮЦІЯ ПЛАСКОГО ТА СКЕВОМОРФНОГО ДИЗАЙНІВ В ІГРОВИХ UI

Задача дослідження полягає в аналізі еволюції скевоморфного та плоского дизайну в ігрових користувацьких інтерфейсах, починаючи від домінування скевоморфізму у ранніх мобільних ОС і до поширення плоского підходу, а також у виявленні того, як ці зміни відобразилися на візуальній мові HUD та меню в іграх різних жанрів.

Мета дослідження полягає у виявленні ролі скевоморфного та плоского дизайну в формуванні атмосферності, впізнаваності та функціональності ігрових

UI, а також у визначенні умов, за яких доцільно використовувати плоский, скевоморфний або гібридний підхід для підтримки жанрової естетики (наприклад, sci-fi проти середньовіччя) та очікувань гравців.

Еволюція дизайну користувацьких інтерфейсів (UI – User Interface), зокрема в ігровій індустрії, є віддзеркаленням ширших тенденцій у сфері UI-дизайну. Приблизно до 2012 року, UI намагалися імітувати реальні об'єкти - кнопки виглядали об'ємними, іконки мали тіні та градієнти, а меню нагадували фізичні панелі управління. Цей підхід мав на меті зробити цифрові інтерфейси більш знайомими для користувачів, але часто призводив до візуального безладу та зайвої складності, як зазначають деякі критики [1]. Трансформація від скевоморфізму, до плоского дизайну істотно вплинула на сприйняття та функціональність UI. Треба зазначити те, що дизайн UI змінювався в іграх згідно тенденцій дизайну в інших сферах, і в першу чергу завжди цю тенденцію надавали операційні системи. Основні кроки еволюції UI визначилися ключовими подіями: впровадження скевоморфізму в iOS у 2007 році та переходом до плоского дизайну з виходом Windows 8 у 2012 році, що задав нову парадигму адаптивного дизайну, зручного для сенсорних екранів і різних розмірів пристроїв[2]. Ці зміни відбивалися і в ігрових UI, де операційні системи традиційно задають тон і задають напрямок розвитку UI-дизайну загалом. Хоча спочатку плоский дизайн Windows 8 і не був прийнятий спільнотою, згодом Microsoft виправили помилки занадто грубого дизайну у Windows 10, та переваги адаптивності на різні пристрої та різні розміри екранів плоского дизайну все ж стали домінуючими.

Розглянемо як це відобразилося в ігровій індустрії. Плоский дизайн приніс із собою філософію, де функціональність має пріоритет над візуальними ефектами. У ігрових UI це означає чіткі іконки без зайвих текстур та об'ємних тіней, прості геометричні форми та обмежену кольорову палітру. Така простота дозволяє гравцям швидше обробляти інформацію та зосереджуватися на ігровому процесі замість декодування UI. Хорошим прикладом є UI Destiny та Destiny 2, де ігровий Heads Up Display (HUD) і меню побудовані на м'яких плоских кольорах, простих блоках без скевоморфних ефектів і мінімальній кількості декоративних деталей, що підкреслюють читабельність і футуристичну естетику[3].

Плоский дизайн в UI забезпечує кілька критичних переваг для користувацького досвіду:

- Зменшує когнітивне навантаження, допомагаючи гравцям швидко знаходити потрібні елементи та дані;
- Відсутність надлишкових деталей сприяє фокусуванню уваги на основних показниках, таких як здоров'я або ресурси;
- Спрощує масштабування UI під різні розміри екрану і пристрої, що важливо для кросплатформених ігор;
- Полегшує локалізацію завдяки структурній простоті;
- Підвищує доступність для людей з порушеннями зору через високий контраст та прості іконки;

- Легко інтегрує функції доступності, як режими для дальтоніків або змінні елементи управління.

Проте, скевоморфізм переживає несподіване відродження у 2025 році, набуваючи нових форм завдяки технологічним досягненням та відчуттю ностальгії серед користувачів. Цей стиль адаптується до сучасних очікувань, поєднуючи скевоморфічні деталі з мінімалізмом плоского дизайну, створюючи комфортні та знайомі користувацькі досвіди. Як приклад: Tesla app з елементами приладової панелі, що імітують реальні матеріали, або інтерфейси VST плагінів з візуальним відтворенням аналогових пристроїв, а також iOS 23, що черпає натхнення зі стилю Windows Vista[4]. Дизайнери тепер використовують скевоморфізм як спосіб створення більш комфортних, знайомих та захоплюючих цифрових користувацьких досвідів.

Kingdom Come: Deliverance 2 демонструє майстерний приклад такого відродження, поєднуючи різні дизайнерські підходи у межах однієї гри. Ігровий HUD використовує спрощений скевоморфічний підхід, де елементи зберігають середньовічну естетику через тонкі рамки та текстури, але при цьому залишаються функціональними та читабельними поєднуючи плоскі стилізовані елементи статусів, іконок на компасі та текстом(рис.1). Особливо помітні динамічні плиткі біля смуги здоров'я, які розширюються як шухляди при зміні зброї, що створює естетичний і функціональний ефект, роблячи цей UI одним із найкращих у RPG жанрі. [5].



Рис.1 – UI у відкритому світі в грі Kingdom Come: Deliverance 2. Одиницями позначені скевоморфні елементи інтерфесу, двійками – плоскі

Підбиваючи підсумки про дизайн інтерфейсу, можна сказати що скевоморфічний дизайн ідеально пасує середньовічним та фентезі іграм, створюючи автентичність через дерев'яні панелі, металеві клепки та пергаментні фони, що занурюють гравця в епоху. Навпаки, sci-fi проекти обирають плоский дизайн з геометричними формами, неоновими кольорами, голограмами та прозорими панелями для відчуття високих технологій. Такий контраст підсилює тематику та полегшує адаптацію гравців.

Список використаних джерел

1. Elements of Game Design. ilogos. URL: <https://ilogos.biz/elements-of-game-design/>.
2. Exposito I. Що таке скевоморфізм і чому він має тенденцію зникати?. creativosOnline. URL: <https://salolli/f441bD7>.
3. Dawson D. Destiny: Futura and Helvetica. dwsn3ee3. URL: <https://dwsn3ee3.wordpress.com/2014/09/16/destiny-futura-and-helvetica/>.
4. Skeuomorphism: an unexpected comeback in 2025 | Kryzalid. Kryzalid. URL: <https://kryzalid.net/en/web-marketing-blog/skeuomorphism-an-unexpected-comeback-in-2025/>.
5. Loreworx. I Redesigned the Kingdom Come Deliverance 2 UI (I'm in love), 2024. YouTube. URL: <https://www.youtube.com/watch?v=CaqBWjXyVQ4>.

Товсточуб І. С. доктор філософії, доцент ДУІКТ
Ярохно Д.В., студент ДУІКТ

МЕТРИКИ ЕФЕКТИВНОСТІ ВІЗУАЛЬНИХ МАРКЕРІВ НА ПРИКЛАДІ ЖОВТОЇ ПІДКРАСКИ

Задача полягає в розробленні формалізованої метрики для комплексної оцінки ефективності візуальних маркерів у геймдизайні, зокрема на прикладі жовтої підкраски, з урахуванням критеріїв успішності навігації, видимості, швидкості виявлення та збереження занурення гравця, з метою знаходження оптимального балансу між функціональністю та атмосферою ігрового досвіду.

Мета дослідження полягає в кількісній оцінці впливу різних типів візуальних маркерів на навігацію та занурення гравця у відеоіграх і виявленні умов, за яких жовта підкраска є доцільною або, навпаки, шкідливою для користувацького досвіду (UX – User Experience).

З розвитком графіки у відеоіграх зросла проблема навігаційної ясності у тривимірних просторах. Ігрові розробники почали використовувати жовту підкраску інтерактивних об'єктів та шляхів, щоб полегшити орієнтацію гравців — яскравий жовтий колір привертає увагу завдяки високій видимості й контрасту з природними відтінками. Ця практика стала популярною, зокрема після ремастерів Resident Evil 4 та Final Fantasy 7 Rebirth[1]. Вона вирішує проблему загублення у складних світах, але викликає критику за порушення ігрового занурення, руйнуючи атмосферу реалістичності (рис. 1). Деякі вважають, що жовта підкраска знижує відчуття дослідження, створюючи враження, що розробники сумніваються у здібностях гравця.



Рис. 1 – Приклад рішень для відображення природних інтерактивних об'єктів ландшафту підкраскою (1) та натуралістичний (2) в сучасних відеоіграх

Жовта підкраска являє собою візуальний орієнтир, або як їх ще називають – візуальний маркер, це елементи дизайну, які спрямовують увагу гравця на важливі об'єкти або напрямки[2].

У цьому дослідженні запропоновано власну формалізовану метрику ефективності візуальних маркерів (Visual Marker Effectiveness) для кількісної оцінки підказок у відеоігровому дизайні. Формула VME об'єднує критичні аспекти UX[3], а саме: успішність навігації, видимість, швидкість виявлення та збереження занурення у єдиний критерій, що дозволяє порівнювати різні типи маркерів між собою:

$$VME = \frac{Success_{rate} \times Visibility}{Discovery_{time} \times Immersion_{break}} \quad 1$$

Розглянемо параметри даної формули:

- $Success_{rate} \in [0,1]$ – відсоток гравців, які правильно інтерпретували маркер і досягли мети (розраховується як відношення успішних спроб до всіх спроб);
- $Visibility \in [0,1]$ – видимість маркера з урахуванням контрасту та розміру;
- $Discovery_{time}$ – час у секундах від появи маркера до його помітки гравцем;
- $Immersion_{break} \in [0,10]$ – суб'єктивна оцінка порушення занурення за 10-бальною шкалою (1 = жодного порушення, 10 = критичне порушення).

Для розрахунку компонента $Visibility$ було розроблено формулу, яка враховує контрастність, різницю яскравості та відносний розмір маркера $Size_{norm}$, що вперше комплексно узагальнює вплив цих факторів для задач геймдизайну й UX-аналізу ігрових інтерфейсів:

$$Visibility = CR \times \left(1 - \left[\frac{L_{marker} - L_{background}}{L_{max}} \right] \right) \times Size_{norm},$$

2

де:

- CR – контрастне співвідношення за WCAG[4] (нормалізоване до діапазону [0, 1];
- $L_{marker}, L_{background}$ – відносна яскравість маркера та фону[5];
- $Size_{norm}$ – нормалізований розмір маркера відносно поля зору.

Результати дослідження:

Було виявлено, що маркери (світлові акценти, кольорові підсвітки, геометричні орієнтири) ефективно скорочують час навігації, але при надмірному використанні (як "жовта підкреска") руйнують занурення. Левел-дизайн потребує балансу між підказками та свободою гравця — адаптивні системи чи стилістично інтегровані маркери можуть стати рішенням. Розроблено метрику Visual Marker Effectiveness (VME) для оцінки навігаційної успішності, видимості, часу виявлення та збереження занурення.

Список використаних джерел

1. IGN. Yellow Paint in Video Games Needs to Stop, 2024. YouTube. URL: <https://www.youtube.com/watch?v=K2wAYQsrnn8>.
2. Henry A. Visual Cues In Level Design. Game Developer. URL: <https://saloli/943F78c>.
3. Winning Game UX Formula: Unite Usability and Engageability. The Acagamic. URL: <https://saloli/47Bd843>.
4. What are The Web Content Accessibility Guidelines (WCAG)? – updated 2025. The Interaction Design Foundation. URL: <https://saloli/De327D5>.
5. Relative luminance - WCAG WG. W3C. URL: <https://saloli/c9A24B2>.

Товсточуб І. С. доктор філософії, доцент ДУІКТ
Колесник В.Я., студент ДУІКТ

ІНТЕЛЕКТУАЛЬНА СИСТЕМА РАНЬОГО ВИЯВЛЕННЯ АНОМАЛІЙ ДОСТУПУ ДО КОРПОРАТИВНОГО VPN

У сучасних умовах цифрової трансформації корпоративних інфраструктур зростає значення систем безпечного віддаленого доступу, серед яких VPN (Virtual Private Network) є одним з ключових компонентів. VPN-сервіси забезпечують захищену комунікацію між співробітниками та внутрішніми ресурсами, однак водночас стають цілью для широкого спектра кібератак, зокрема підбору облікових даних, викрадення сесій, експлуатації ненадійних клієнтів і спроб несанкціонованого проникнення. Актуальні події у сфері кібербезпеки демонструють, що компрометація VPN-аккаунтів часто є першим етапом складних багатоступеневих атак на організації [1][2].

Проблема полягає в тому, що традиційні засоби моніторингу, зокрема класичні SIEM-системи, побудовані на фіксованих сигнатурах і простих правилах

кореляції, не здатні ефективно виявляти багатфакторні, нетипові або повільні за розвитком аномалії [3]. Часто зловмисники використовують тактику імітації поведінки легітимних користувачів, поєднання сесій з різних IP-адрес, появу у незвичних часових інтервалах та поступову зміну активності, що не потрапляє під статичні правила виявлення. Це потребує переходу від реактивних підходів до моделей, здатних прогнозувати відхилення на основі поведінкового профілю користувача.

Метою дослідження є розроблення інтелектуальної системи раннього виявлення аномалій у доступі до корпоративного VPN шляхом застосування комплексного підходу, що поєднує автоматизований лог-аналіз, поведінкову аналітику та методи машинного навчання. Запропонована система орієнтована на виявлення нетипових, потенційно шкідливих патернів доступу до ресурсу на ранніх етапах, що дає змогу запобігти ескалації атаки.

У межах роботи проведено аналіз структури логів популярних VPN-рішень, зокрема OpenVPN, WireGuard та IPSec [4][5]. Окрему увагу приділено відмінностям у форматах журналів подій, алгоритмах встановлення та завершення сесій, способах фіксації помилкових входів, видачі токенів автентифікації та параметрах транспортних з'єднань. На основі цього описано процедуру ETL-обробки логів: попереднє очищення, нормалізацію часових міток, уніфікацію полів, виділення сесій користувача, фільтрацію службових записів та агрегування у єдину модель подій.

Для побудови поведінкових векторів користувача виділено низку ключових ознак, що дозволяють моделювати нормальні патерни роботи: географічне походження трафіку, час доби та день тижня, швидкість переходу між різними IP-адресами (velocity anomaly), стабільність параметрів клієнта (User-Agent, версія протоколу), кількість невдалих та успішних спроб автентифікації, тривалість сесій, використання нестандартних портів або шифрувальних наборів. Також запропоновано низку кореляційних ознак, побудованих на груповій поведінці користувачів, що дозволяє виявляти синхронізовані аномальні дії.

У дослідженні аналізуються методи виявлення аномалій, до яких належать як класичні підходи (Isolation Forest [1], Local Outlier Factor [2]), так і сучасні глибинні моделі (автокодери, LSTM та 1D-CNN [3]). Оцінено можливість застосування алгоритмів для роботи з часовими рядами, зокрема послідовні моделі LSTM для виявлення нетипової зміни поведінки у часовому вимірі, а також згорткові нейронні мережі для виокремлення повторюваних структур у VPN-трафіку. Обґрунтовано актуальність використання гібридної системи, яка поєднує ML-методи з традиційними SIEM-правилами, що дозволяє зменшити кількість хибнопозитивних результатів та підвищити точність виявлення складних атак.

Серед типових сценаріїв, що можуть бути ефективно виявлені інтелектуальною системою, виділено такі: спроби підбору пароля з використанням географічно розподілених IP-адрес; нетипові нічні або святкові підключення користувачів; різка зміна регіону доступу («подорож користувача» між країнами за короткий час); зловживання сесіями тривалої дії; активність з аномальними параметрами клієнта; появи у логах незвичайних шаблонів помилкових входів; ознаки lateral movement у корпоративній мережі. Окремо розглянуто можливість виявлення доступів із заражених пристроїв за зафіксованими аномальними характеристиками трафіку.

Попередні результати свідчать, що застосування поведінкової аналітики та машинного навчання дозволяє підвищити ймовірність раннього виявлення аномалій, забезпечити більш повне охоплення складних сценаріїв та зменшити навантаження на аналітиків SOC. Очікується, що інтеграція запропонованої інтелектуальної системи до корпоративного середовища сприятиме реалізації принципів Zero Trust та забезпечить вищий рівень адаптивного захисту у порівнянні з фіксованими правилами.

Список використаних джерел

1. Chandola V., Banerjee A., Kumar V. Anomaly Detection: A Survey. ACM Computing Surveys. 2009.
2. Ahmed M., Mahmood A., Hu J. A survey of network anomaly detection techniques. Journal of Network and Computer Applications. 2016.
3. Sommer R., Paxson V. Outside the Closed World: On Using Machine Learning for Network Intrusion Detection. IEEE Symposium on Security and Privacy.
4. Documentation: OpenVPN. <https://openvpn.net>.
5. WireGuard Technical Whitepaper. <https://www.wireguard.com>

Товсточуб І. С. доктор філософії, доцент ДУІКТ
Колесник В.Я., студент ДУІКТ

ІНТЕЛЕКТУАЛЬНА СИСТЕМА РАНЬОГО ВИЯВЛЕННЯ АНОМАЛІЙ ДОСТУПУ ДО КОРПОРАТИВНОГО VPN

Зростання кількості кіберінцидентів, пов'язаних із несанкціонованими підключеннями до корпоративних ресурсів, зумовлює необхідність запровадження адаптивних систем контролю доступу. VPN-технології, що забезпечують захищений канал передачі даних, водночас є точкою підвищеного ризику, адже компрометація облікових даних користувача відкриває зловмиснику шлях до внутрішньої інфраструктури компанії. Поширення атак типу password spraying, використання анонімізуючих сервісів, автоматизованих бот-мереж і методів прихованої ескалації робить традиційні системи моніторингу малоефективними [1][2].

Основною проблемою є те, що класичні SIEM-рішення аналізують окремі події, а не комплексну поведінку користувача, і опираються переважно на статичні правила. Такі підходи добре працюють для простих сценаріїв порушень, але не здатні виявляти багатоступеневі атаки, під час яких зловмисник поступово змінює власні патерни активності, адаптується під звичний профіль користувача або діє з використанням скомпрометованих пристроїв [3]. Понад те, швидка зміна географічного розташування, використання нетипових версій клієнтів або багаторазові помилки автентифікації можуть залишатися непоміченими, якщо розглядати їх у відриві від інших характеристик доступу.

Метою дослідження є розроблення інтелектуального підходу до раннього виявлення аномалій у VPN-доступі на основі аналізу послідовностей подій та моделювання поведінкових профілів користувачів. На відміну від сигнатурних систем, запропонований підхід орієнтується на комплексне оцінювання моделей активності, включно з часовими, географічними та технічними характеристиками, що дозволяє виявляти як очевидні, так і приховані аномалії.

Дослідження включає аналіз логів VPN-протоколів OpenVPN [4] і WireGuard [5], порівняння їх структури, способу формування сесій та специфіки реєстрації помилок автентифікації. На основі отриманих даних сформовано ETL-процес, що дозволяє перетворити неструктуровані журнали в уніфіковані моделі подій. Особливу увагу приділено створенню поведінкового вектора, який містить багатовимірні ознаки: стабільність IP-адрес, інтенсивність входів, часові шаблони, швидкість переміщення між географічними точками, сталість параметрів клієнтського ПЗ, кількість невдалих та успішних спроб підключення.

Для виявлення аномальної активності розглянуто класичні (Isolation Forest [1], LOF [2]) та сучасні нейронні підходи (автокодери, LSTM, 1D-CNN [3]). Окремо оцінено застосовність рекурентних моделей до послідовностей подій для виокремлення нетипових динамічних змін, а також можливість використання згорткових мереж для аналізу структурних послідовностей запитів. Запропоновано гібридну схему поєднання машинного навчання з правилами кореляції SIEM, що дозволяє зменшити кількість хибнопозитивних результатів.

Отримані результати засвідчують, що комплексний поведінковий аналіз у поєднанні з ML-підходами забезпечує суттєве підвищення ефективності виявлення аномалій. Такий підхід дає змогу виявляти критичні сценарії, включно зі швидкою зміною географії, множинними помилками автентифікації, підозрілими шаблонами входів і потенційними ознаками lateral movement. Інтеграція запропонованої системи до корпоративної інфраструктури сприятиме реалізації стратегій Zero Trust, підвищить рівень автоматизації SOC-процесів і знизить ризики несанкціонованого доступу.

Список використаних джерел

6. Chandola V., Banerjee A., Kumar V. Anomaly Detection: A Survey. ACM Computing Surveys. 2009.
7. Ahmed M., Mahmood A., Hu J. A survey of network anomaly detection techniques. Journal of Network and Computer Applications. 2016.
8. Sommer R., Paxson V. Outside the Closed World: On Using Machine Learning for Network Intrusion Detection. IEEE Symposium on Security and Privacy.
9. Documentation: OpenVPN. <https://openvpn.net>.
10. WireGuard Technical Whitepaper. <https://www.wireguard.com>

Тестов Т., Чичкарьов Є., Вишневський В., Зінченко О., ДУІКТ

ВИДІЛЕННЯ КЛЮЧОВИХ СЛІВ З БАГАТОМОВНОГО ТЕКСТУ ПРИ ОРГАНІЗАЦІЇ KEYWORD DRIVEN TESTING

Тестування програмного забезпечення – це останній етап, який необхідно пройти в розробці програмного забезпечення. Якість програмного забезпечення можна визначити лише на етапі тестування. Завдяки технологічному прогресу в усьому світі відбулося багато вдосконалень у методах і методах, що використовуються для тестування програмного забезпечення [1-3].

Одним з підходів, який можна використовувати, є тестування на основі ключових слів у процесі тестування, що є концепцією з ISO/IEC/IEEE 29119 [4-5]. Тестування на основі ключових слів – це метод, за якого тестові випадки створюються за допомогою ключових слів високого рівня, що представляють дії користувача. Кожне ключове слово відповідає певній функції, наприклад, натисканню кнопки, введенню тексту або перевірці результату. Ці дії зазвичай перераховані в табличному форматі, такому як електронна таблиця, де кожен рядок визначає крок тестування.

Ця робота має на меті проаналізувати можливість виділення ключових слів в описі вимог до програмного забезпечення для створення тестових випадків.

Потужністю цього підходу є його абстракція. Замість того, щоб писати код для кожного тесту, тестувальники просто впорядковують низку попередньо визначених ключових слів. Базовий код, що стоїть за кожним ключовим словом, виконує фактичні операції, дозволяючи зосередитися на тому, «що» робить тест, а не на тому, «як» він це робить.

У тестуванні на основі ключових слів ключові слова складаються з 2 рівнів, описаних наступним чином [4-5]:

1. Низький рівень, на цьому рівні ключі пов'язані з набором однієї або кількох дій, які пояснюють, які кроки необхідно виконати в процесі тестування. Приклад: клік, встановлення тексту, вибір та інші.

2. Високий рівень, на цьому рівні ключові слова вимагають набору вхідних параметрів, які також включені до структури. Ключові слова та параметри

утворюють високорівневий опис дій, пов'язаних з тестовим випадком. Приклад: відкриття браузера, вхід та інші.

Можна створювати ключові слова, що представляють дії на різних рівнях абстракції. Зазвичай розрізняють рівень домену та рівень графічного інтерфейсу користувача [4-5].

Ключові слова в шарах домену відповідають бізнесу або діяльності, пов'язаній з доменом, і відображають термінологію, яку використовують експерти в предметній області. Ключові слова, розроблені на шарах домену, зазвичай не залежать від реалізації.

Ключові слова в шарах інтерфейсу тестування відносяться до типу інтерфейсу, що використовується в певному тесті. Дії, необхідні для вирішення тестових завдань, зазвичай легко визначити.

Тестування на основі ключових слів вимагає ідентифікації та визначення ключових слів. Існує кілька джерел, які можна використовувати для ідентифікації та визначення ключових слів, включаючи наступні [4-6]:

1. Дослідницьке тестування

Під час дослідницького тестування екзаменатор спостерігає, які кроки виконуються. Нове ключове слово визначається шляхом надання змістовної назви. Якщо послідовність кроків можна використовувати з різними даними, ключове слово матиме відповідні параметри для цих даних.

2. Бізнес-експертиза

Ключові слова можна визначити шляхом проведення інтерв'ю з бізнес-експертами. Це питання може бути «Що слід зробити для перевірки програми?» або «що потрібно протестувати?». Відповіді, дані експертами, будуть використані для ідентифікації ключових слів шляхом пошуку термінів, які можуть часто зустрічатися.

3. Тестування інтерфейсу

Ключові слова можна вказати через тестовий інтерфейс, але оскільки елементи інтерфейсу обмежені та зазвичай невеликі, для цих елементів інтерфейсу можна визначити певну кількість ключових слів. Цей підхід визначить ключові слова низького рівня в шарах тестового інтерфейсу.

4. Документовані процедури тестування та тестові випадки

Наявні процедури тестування та тестові випадки також можуть бути використані як матеріал для визначення ключових слів. Якщо два або більше знайдених ключових слів стосуються однієї й тієї ж діяльності, ці ключові слова будуть замінені лише тим ключовим словом, яке найкраще описує діяльність.

Щоб вирішити проблему вилучення ключових слів з тексту опису програмного забезпечення або проблемно-орієнтовного тексту, у цій роботі запропоновано метод вилучення ключових слів або n-грам за допомогою моделей Bert з вибором мовно-орієнтованої моделі.

За думкою [7], ключові слова вибираються зі статті та використовуються для вираження теми статті. За ключовими словами ми можемо швидко та точно зрозуміти центральний зміст довгого тексту, що забезпечує велику зручність пошуку. Тому вилучення ключових слів з тексту привертає все більше уваги в останні роки та широко використовується в багатьох галузях, пов'язаних з

обробкою природної мови, таких як обчислення подібності тексту, генерація діалогових систем, аналіз тем тексту, класифікація тексту.

Дослідження процесу виділення ключових слів було виконано в середовищі Google Colaboratory з використанням мови програмування python і відповідних пакетів і бібліотек.

Встановлено, що для вилучення ключових слів з україномовного тексту найкращі результати забезпечує використання пакету keybert з завантаженням моделі «lang-uk/ukr-paraphrase-multilingual-mpnet-base» [8-9]. Це модель трансформації речень, точно налаштована для української мови: вона відображає речення та абзаци у 768-вимірний щільний векторний простір і може бути використана для таких завдань, як кластеризація або семантичний пошук. За даними [8], підхід до задачі розв'язання сенсу слів українською мовою, заснований на контрольованому точному налаштуванні попередньо навченої моделі великої мови (LLM) на наборі даних, згенерованому без нагляду, для отримання кращих контекстних вбудовувань для слів з кількома значеннями.

Для багатомовних текстів було використано розбиття тексту на речення з встановлення мови речення (було використано пакет langdetect). Для розбиття тексту на речення було використано пакет spacy. Для виділення ключових слів російською або англійською мовами достатню надійність і точність забезпечили моделі «all-mpnet-base-v2» та «DeepPavlov/rubert-base-cased-sentence».

Перевірка виділення ключових слів та пошуку основного тексту за рефератом показало, що точність пошуку становила 85-90%.

Перелік посилань

1. Arumugam, Arun. (2020). Software Testing Techniques New Trends. International Journal of Engineering Research and. V8. 10.17577/IJERTV8IS120318.
2. Tetteh, Samuel Gbli. (2024). Software Testing Techniques and Levels in Software Development. International Journal of Advancements in Computing Technology. 2. 10-19. 10.56472/25838628/IJACT-V2I1P102.
3. Jia, Xiaoqing. (2025). Research on Software Security Testing Mode Based on Open Network Environment. Journal of Electronic Research and Application. 9. 290-296. 10.26689/jera.v9i4.11472.
4. ДСТУ ISO/IEC/IEEE 29119-1:2017 Інженерія систем і програмних засобів. Тестування програмних засобів. Частина 1. Поняття та визначення (ISO/IEC/IEEE 29119-1:2013, IDT) URL: https://online.budstandart.com/ua/catalog/doc-page?id_doc=75488
5. International Electrotechnical Commission, Institute of Electrical and Electronics Engineers, and International Electrotechnical Commission. Technical Committee 91, Standard for automatic test markup language (ATML) test adapter description. URL: <https://www.iso.org/obp/ui/en/#iso:std:iso-iec-ieee:29119:-5:ed-2:v1:en>
6. Octavially, Reynaldi & Riskiana, Rosa & Laksitowening, Kusuma & Kusumo, Dana & Adrian, Monterico & Selviandro, Nungki. (2022). Test Case Analysis with Keyword-Driven Testing Approach on Angkasa Website Using Katalon Studio Tools. Ultimatics : Jurnal Teknik Informatika. 13. 134-141. 10.31937/ti.v13i2.2391.

7. Yili Qian et al (2021) Bert-Based Text Keyword Extraction. J. Phys.: Conf. Ser. 1992 042077 URL: <https://dx.doi.org/10.1088/1742-6596/1992/4/042077>
8. Yurii Laba, Volodymyr Mudryi, Dmytro Chaplynskyi, Mariana Romanyshyn, and Oles Dobosevych. 2023. [Contextual Embeddings for Ukrainian: A Large Language Model Approach to Word Sense Disambiguation](#). In *Proceedings of the Second Ukrainian Natural Language Processing Workshop (UNLP)*, pages 11–19, Dubrovnik, Croatia. Association for Computational Linguistics.
9. Model lang-uk/ukr-paraphrase-multilingual-mpnet-base URL: <https://huggingface.co/lang-uk/ukr-paraphrase-multilingual-mpnet-base>

Гніденко М.П., к.т.н., доцент ДУІКТ
Шепетун О.О., аспірант ДУІКТ

ВИКОРИСТАННЯ МОДУЛЬНИХ НЕЙРОМЕРЕЖЕВИХ МОДЕЛЕЙ ДЛЯ ПОКРАЩЕННЯ ЯКОСТІ ІДЕНТИФІКАЦІЇ ЛЮДИНИ НА ФОТОЗОБРАЖЕННЯХ

Вступ. Значного розвитку досягли методи покращення якості ідентифікації та реєстрації особи за її фотозображенням, а також алгоритми візуальної ідентифікації особи за її зображеннями у відеопотоці на базі алгоритмів Віоли-Джонса та FisherFace, продовженням яких на новому рівні стало використання модульних нейромережових моделей.

Постановка завдання. Дослідити використання нейронних мереж для підвищення точності ідентифікації особи за фотозображенням у сфері дистанційної роботи (контролю присутності особи в певному місці (дистанційне робоче місце, дистанційне навчання)).

Мета дослідження. Описати сутність візуальної ідентифікації особи з допомогою та обґрунтувати доцільність її застосування з допомогою модульних нейромережових моделей для покращення якості, зменшити кількість похибок в процесі візуальної ідентифікації із застосуванням візуальних біометричних ідентифікаторів, що зберігаються в фотозображеннях.

Результати дослідження. Розглянуто та визначено доцільною ідентифікацію особи за її візуальними біометричними показниками із застосуванням нейронних мереж як допоміжних інструментів. На стадії обробки кадрів з відеоряду здійснюється покращення якості зображення для подальшого більш точного виявлення та розпізнавання обличчя в кадрі, порівнюючи його із набором зображень із бази даних. Дана система може мати своє застосування в сфері контролю персоналу за віддаленими місцями роботи.

Висновки. Дослідження має продовжуватись для напрацювання найбільш ефективних комбінацій модульних нейромережових моделей для отримання максимального ефекту для свого часу. Інновації за даним напрямом мають потенціал для практичного застосування та, як наслідок, забезпечення вивільнення програмних та фінансових ресурсів з процесів ідентифікації особи (за

рахунок їх оптимізації) та спрямування їх на інші процеси вдосконалення робочого процесу.

Список використаних джерел

1. Джавед, А., 2013 Розпізнавання обличчя на основі аналізу головних компонентів.. International Journal of Image, Graphics and Signal Processing, 2, pp.38-44.
2. Кобець, Н. та Ковалюк, Т., 2020. Метод розпізнавання та індексування обличч людей у відео за допомогою моделі машинного навчання. Досягнення інтелектуальних систем і обчислень, 1247, pp.534-544.
3. Тін, Х. та Хтаке, Х., 2012. Сприйнята гендерна класифікація за зображеннями обличчя. International Journal of Modern Education and Computer Science, 1, pp.12-18.
4. Віола, П. та Джонс, М.І., 2001. Швидке виявлення об'єктів за допомогою розширеного каскаду простих функцій. Proceedings IEEE Conference on Computer Vision and Pattern Recognition, pp.1-9.
5. Wang, Y-Q., 2014. Аналіз алгоритму виявлення обличчя Віоли-Джонса. Обробка зображень в режимі онлайн, 4, pp.128-148.

Іщеряков С.М., доцент ДУІКТ
Дубовицький Д.С., аспірант ДУІКТ

ВПЛИВ СЕРВЕРНОГО РЕНДЕРУ НА ПРОДУКТИВНІСТЬ ВЕБ-СТОРІНОК

Постановка задачі

Більшість веб-сайтів на сьогоднішній день будуються за допомогою виконання скриптів JavaScript. Виконання JavaScript коду при старті сторінки негативно впливає на продуктивність сторінки, таким чином що після завантаження сторінки, основний потік JavaScript є заблокованим під час виконання вкладених скриптів. Заблокування основного потоку JavaScript призводить до того, що сторінка не відповідає на дії користувача. Параметр який вираховує час до того, коли сторінка стає інтерактивною після першого свого відображення називається Total Blocking Time, та відповідає за 30% загальної оцінки веб-сайту в категорії продуктивність. Тому необхідно вирішити наукове завдання з пошуку способу зниження часу блокування та поліпшення загальної продуктивності сторінки.

Мета дослідження

Підвищити ефективність відображення веб-сайту за допомогою рендеру веб-сайту на боці серверу. Дослідити вплив використання технології на загальну продуктивність відображення веб-сайту, зокрема параметр Total Blocking Time. Основна ідея даного підходу полягає у виконанні JavaScript на боці серверу, позбавляючи браузер клієнта від необхідності його виконання під час

відображення сторінки. Крім того, результатом рендеру на боці серверу є статична HTML сторінка, що забезпечує можливість індексації пошуковою машиною для веб-сайтів які раніше не мали цієї можливості.

Результати дослідження

Під час завантаження веб-сторінки, яка наповнюється динамічно за рахунок виконання скриптів JavaScript, сторінка є заблокованою для взаємодії до тих пір, поки основний потік виконання JavaScript не закінчить свою роботу. Дана умова викликана тим, що JavaScript в браузері має один основний потік який відповідає за виконання скриптів, малювання інтерфейсу, та обробку подій. Таким чином виконання JavaScript при старті сторінки для наповнення веб-сайту вмістом блокує взаємодію користувача з сайтом до звільнення основного потоку. Метрика, яка заміряє час від того як веб-сайт завантажений до того, як основний потік стане вільним і сайт доступним для взаємодії називається Total Blocking Time (TBT). Метрика TBT впливає на загальну оцінку продуктивності веб-сайту, таким чином що оцінка TBT впливає на 30% загальної оцінки продуктивності веб-сайту при вимірах за допомогою системи Google Lighthouse. При виконанні JavaScript веб-сайту завчасно на боці серверу, досягається результат, при якому після завантаження веб-сайт вже є наповнений вмістом, таким чином що браузеру не потрібно виконувати більшість скриптів. За рахунок звільнення від виконання скрипту для наповнення веб-сайту, досягається зниження Total Blocking Time, що має позитивний ефект на оцінку продуктивності веб-сайту пошуковими машинами та сприйманні веб-сайту користувачем. Крім того, оскільки сайт тепер є сгенерованим завчасно, він стає доступним для індексації пошуковою машиною, що також є значним покращенням SEO. Дослідження показало, що для веб-сайту з наповненням в 60 елементів, за допомогою рендеру на боці серверу параметр TBT було знижено з 200 мілісекунд до 20. Таким чином, технологія рендеру на боці серверу є перспективним рішенням для покращення продуктивності відображення веб-сайту.

Висновки та перспективи

Використання рендеру на боці серверу дає можливість значно підвищити продуктивність відображення сторінки. Також, дана технологія дає можливість індексації веб-сайту, вміст якого створюється динамічно за допомогою JavaScript. Перспективами розвитку є кешування завчасно оброблених сторінок та покращення швидкості обробки.

Список використаних джерел

1. First Contentful Paint | Lighthouse | Chrome for Developers URL: <https://developer.chrome.com/docs/lighthouse/performance/first-contentful-paint>
2. Total Blocking Time (TBT) URL: <https://web.dev/articles/tbt>
3. Strazzullo, F. (2023). Web Components. In: Frameworkless Front-End Development. Apress, Berkeley, CA. https://doi.org/10.1007/978-1-4842-9351-5_5

Дубовицький Д. С., аспірант ДУІКТ

Емуляція браузерного середовища для виконання скриптів JavaScript на сервері: методи та виклики

Державний університет телекомунікацій

Постановка задачі. На сьогоднішній день веб-сайти є одним з основних методів створення графічного інтерфейсу для більшості систем. Сучасний веб-сайт складається з сотень елементів, кожен з яких динамічно додається або видаляється зі сторінки відповідно до умов, і так само динамічно наповнюється змістом будь то через запити до бази даних, до інших веб-сайтів, або генерація за допомогою власної логіки. Подібні задачі з динамічного відображення вмісту сторінки виконуються за допомогою JavaScript. Подібне динамічне створення вмісту сторінок є відносно новим феноменом, а сама технологія веб-сайтів збудована навколо статичної розмітки HTML. Таким чином, динамічно створений веб-сайт неможливо проіндексувати пошуковою машиною та він буде пустим після завантаження. Тому необхідно вирішити наукове завдання з розробки методу, який би міг зробити динамічно генеруємий JavaScript сайт доступним як статична сторінка для клієнту.

Мета дослідження.

Підвищити ефективність відображення веб-сайту за допомогою методу емуляції браузерного середовища на боці серверу, таким чином щоб JavaScript код сторінки, написаний для браузерів, міг коректно відпрацювати на боці серверу. Дослідити проблеми сумісності, що виникають при виконанні клієнтського скрипта в серверному середовищі, та запропонувати оптимальні рішення.

Результати дослідження.

Основними проблемами при виконанні скриптів призначених для браузерного середовища на стороні серверу є несумісність або відсутність API браузера. Оскільки сервер не працює в середовищі, де не існує графічного інтерфейсу, в ньому відсутні API які відповідають за парсинг HTML, створення та маніпуляцію document object model. Також, якщо веб-сайт використовує технологію веб-компонентів, необхідно додати емуляцію реєстру власних компонентів, та підтримку заміни тегу власного компоненту його вміст, який зазначається під час реєстрації власного компоненту. Для того щоб зробити емуляцію document object model або DOM, необхідно визначити її структуру. DOM це граф типу дерево, він складається вузлів елементів або текстових вузлів. Кожен текстовий вузол є кінцевим вузлом, а вузол елемент може бути або кінцевим, або вузлом розгалуження. Вузли елементів мають підтримувати такі методи як getByld, для пошуку вузлів з відповідним атрибутом id, які є нижчими за рівнем відповідно до поточного. Кожен вузол має підтримувати властивість innerHTML, яка задає можливість змінювати та отримувати вміст HTML конкретного вузла. Також необхідно реалізувати емуляцію технології веб-компонентів. Вона складається з реєстру власних компонентів customRegistry, яка представляє собою структуру даних Map, де ключем є назва тегу, а значенням

клас власного компоненту. Під час парсингу та рендеру HTML розмітки, теги які було додано до реєстру власних тегів необхідно замінити на вміст власного компоненту. Елементи які використовують технологію Shadow Dom мають бути розміщені в тег `template` з атрибутом `shadowrootmode` зі значенням `open`. Емуляція браузерного середовища на сервері дозволяє виконати JavaScript завчасно, таким чином що результатом роботи буде статична HTML сторінка. За рахунок завчасного виконання JavaScript на боці серверу, цей крок не потрібно виконувати на клієнті під час відображення веб-сайту, що знижує загальний час відображення та знижує час блокування інтерфейсу. Також статична сторінка додає можливість індексації сторінки.

Висновки та перспективи

Емуляція браузерного середовища на боці серверу дає можливість виконувати скрипти створені для виконання в браузері на сервері. Таким чином, ця технологія дозволяє змістити навантаження з браузера клієнта до серверу, зменшуючи загальний час та кількість JavaScript які мають бути виконані для відображення сторінки. Перспективами розвитку є емуляція додаткових АПІ браузера таких як XMLHttpRequest для підтримки генерації додатків які для свого відображення використовують інформацію завантажену з сторонніх ресурсів.

Шепетун О.О., аспірант ДУІКТ

Гніденко М.П., к.т.н., доцент ДУІКТ

ВИКОРИСТАННЯ МОДУЛЬНИХ НЕЙРОМЕРЕЖЕВИХ МОДЕЛЕЙ ДЛЯ ПОКРАЩЕННЯ ЯКОСТІ ІДЕНТИФІКАЦІЇ ЛЮДИНИ НА ФОТОЗОБРАЖЕННЯХ

Вступ. Значного розвитку досягли методи покращення якості ідентифікації та реєстрації особи за її фотозображенням, а також алгоритми візуальної ідентифікації особи за її зображеннями у відеопотоці на базі алгоритмів Віюлі-Джонса та FisherFace, продовженням яких на новому рівні стало використання модульних нейромережових моделей.

Постановка завдання. Дослідити використання нейронних мереж для підвищення точності ідентифікації особи за фотозображенням у сфері дистанційної роботи (контролю присутності особи в певному місці (дистанційне робоче місце, дистанційне навчання)).

Мета дослідження. Описати сутність візуальної ідентифікації особи з допомогою та обґрунтувати доцільність її застосування з допомогою модульних нейромережових моделей для покращення якості, зменшити кількість похибок в процесі візуальної ідентифікації із застосуванням візуальних біометричних ідентифікаторів, що зберігаються в фотозображеннях.

Результати дослідження. Розглянуто та визначено доцільною ідентифікацію особи за її візуальними біометричними показниками із застосуванням нейронних мереж як допоміжних інструментів. На стадії обробки

кадрів з відеоряду здійснюється покращення якості зображення для подальшого більш точного виявлення та розпізнавання обличчя в кадрі, порівнюючи його із набором зображень із бази даних. Дана система може мати своє застосування в сфері контролю персоналу за віддаленими мсцями роботи.

Висновки. Дослідження має продовжуватись для напрацювання найбільш ефективних комбінацій модульних нейромережевих моделей для отримання максимального ефекту для свого часу. Інновації за даним напрямом мають потенціал для практичного застосування та, як наслідок, забезпечення вивільнення програмних та фінансових ресурсів з процесів ідентифікації особи (за рахунок їх оптимізації) та спрямування їх на інші процеси вдосконалення робочого процесу.

Список використаних джерел

1. Джавед, А., 2013 Розпізнавання обличчя на основі аналізу головних компонентів.. International Journal of Image, Graphics and Signal Processing, 2, pp.38-44.
2. Кобець, Н. та Ковалюк, Т., 2020. Метод розпізнавання та індексування обличч людей у відео за допомогою моделі машинного навчання. Досягнення інтелектуальних систем і обчислень, 1247, pp.534-544.
3. Тін, Х. та Хтаке, Х., 2012. Сприйнята гендерна класифікація за зображеннями обличчя. International Journal of Modern Education and Computer Science, 1, pp.12-18.

Шикула О.М. , професор ДУІКТ
Целованський Т.Р., аспірант ДУІКТ

КОМПЛЕКСНІ МЕТОДИ ОПТИМІЗАЦІЇ ПІДСИСТЕМИ ПАМ'ЯТІ НА СЕРВЕРНИХ БАГАТОСОКЕТНИХ СИСТЕМАХ

Постановка задачі.

Сучасні центри обробки даних використовують багатосокетні серверні системи архітектури NUMA (Non-Uniform Memory Access), де кожен сокет має власний контролер пам'яті. Така архітектура створює критичну диспропорцію затримок доступу: до локальної пам'яті 60–100 наносекунд, до віддаленої — 200–400 наносекунд (різниця в 3–4 рази). Це призводить до значної деградації продуктивності додатків.

Розвиток штучного інтелекту та аналізу великих даних посилює проблему непередбачуваними шаблонами доступу до пам'яті. Попередні дослідження зосередились на окремих методах оптимізації, ігноруючи синергетичні ефекти при одночасному застосуванні кількох методик (NUMA оптимізація, величезні сторінки пам'яті, чередування каналів, попереджувальне завантаження) та їхній вплив на енергоефективність.

Мета дослідження.

Встановити синергетичні ефекти комбіноване впровадження сучасних технологій оптимізації взаємодії з ОЗП на багатосокетних серверних системах та розробити практичні рекомендації для впровадження в промислових центрах обробки даних.

Результати дослідження.

Експериментальне дослідження проведено на двох типах сучасних серверних конфігурацій: Intel Xeon та AMD EPYC під управлінням ОС Linux kernel 6.5+ з повною підтримкою NUMA.

Дослідження реалізовано у вигляді шести послідовних експериментальних фаз: (1) базовий аналіз без оптимізацій як еталон; (2) NUMA оптимізація з прив'язанням потоків та розміщенням даних; (3) включення величезних сторінок пам'яті (Transparent Huge Pages); (4) оптимізація чередування каналів пам'яті; (5) активація попереджувального завантаження даних (Prefetching); (6) комбінована оптимізація зі всіма методиками одночасно.

Вимірювання здійснювались за допомогою: Intel VTune Profiler для детального аналізу архітектурних подій, Linux perf для контролю апаратних лічильників продуктивності, LIKWID для комплексного аналізу продуктивності та енергоспоживання.

Комбінована оптимізація забезпечила значні поліпшення всіх ключових метрик: затримка доступу до пам'яті скоротилась на 38% (320→195 нс), пропускна здатність поліпшена на 62% (48→78 GB/s), частота промахів кешу зменшена на 67% (18,7%→6,2%), енергоефективність повисилась на 72% (1,41→2,43 GFLOPS/W).

Диференційований аналіз показав найбільший ефект оптимізації на пам'ять-інтенсивних додатках: обробка графів +67% продуктивності, машинне навчання TensorFlow +55%, бази даних PostgreSQL +42%, Key-Value сховища Redis +45%, HPC LINPACK +28%, веб-сервери Nginx +18%.

Встановлена синергія між методами оптимізації: комбінований ефект 67% перевищує математичну суму індивідуальних вкладів 61%, демонструючи важливість комплексного підходу. Коефіцієнти взаємодії: NUMA+Huge Pages=0,04; NUMA+Prefetch=0,06; Huge Pages+Prefetch=0,03.

Детальний моніторинг системи показав: частота попадань в L1 кеш збільшилась з 94,2% до 96,8%, L2 кеш — з 78,5% до 89,2%, L3 кеш — з 45,3% до 68,7%. Енергоспоживання системи зменшилось на 20% (850→680 Вт), контекстні перемикання скоротились на 72,5% (1245→342 на секунду).

Висновки.

Синергія виникає через три механізми: NUMA оптимізація зменшує промахи L3 кеша, роблячи Prefetching більш ефективним; Huge Pages розвантажують TLB, вивільняючи місце в кеші L1/L2 для корисних даних; чередування каналів пам'яті забезпечує рівномірний розподіл навантаження по каналам. Комбінований підхід до оптимізації ОЗП є критичним для досягнення максимальної продуктивності (67% > 61%). Оптимізація вимагає координованих дій на рівні BIOS, ОС, компілятора та прикладного коду. Ефект залежить від типу додатка: максимум для пам'ять-інтенсивних операцій (графи +67%, ML +55%), мінімум для мережевих сервісів (+18%).

Список використаних джерел

1. OpenGauss Documentation. Kunpeng Numa Architecture Optimization [Електронний ресурс]. — 2025. — Режим доступу: <https://docs.opengauss.org/en/docs/3.1.0/docs/CharacteristicDescription/kunpeng-numa-architecture-optimization.html>
2. Simcentric. How to Optimization Memory Latency for Server Performance [Електронний ресурс]. — 2024. — Режим доступу: <https://www.simcentric.com/hong-kong-dedicated-server/how-to-optimization-memory-latency-for-server-performance/>

Миронов Д.В.

Верещака А.О., студент ДУІКТ

ВПРОВАДЖЕННЯ АРХІТЕКТУРИ НУЛЬОВОЇ ДОВІРИ (ZERO TRUST) ДЛЯ ПІДВИЩЕННЯ БЕЗПЕКИ СУЧАСНИХ КОМП'ЮТЕРНИХ МЕРЕЖ

Постановка задачі. Дослідити недоліки традиційних моделей мережевої безпеки, заснованих на захисті периметра, та проаналізувати архітектурні принципи побудови мереж за методологією Zero Trust для мінімізації ризиків несанкціонованого доступу та витоку даних.

Мета дослідження: Визначити ключові компоненти архітектури Zero Trust, їх вплив на проектування корпоративних мереж та розробити алгоритм переходу від класичної топології до мережі з «нульовою довірою».

Актуальність проблеми. Сучасні комп'ютерні системи стикаються з розмиванням мережевого периметра через масовий перехід на віддалену роботу, використання хмарних сервісів (SaaS, IaaS) та концепції BYOD (Bring Your Own Device). Традиційні методи захисту, що довіряють усьому трафіку всередині локальної мережі, стають неефективними проти інсайдерських загроз та складних цільових атак (APT).

Основна мета дослідження. Аналіз архітектурних змін, необхідних для реалізації моделі Zero Trust, та оцінка їх ефективності у протидії латеральному (горизонтальному) переміщенню зловмисників у мережі.

Результати аналізу. Впровадження архітектури Zero Trust (ZTA) змінює фундаментальний підхід до проектування комп'ютерних мереж. Якщо класична модель керується принципом «довіряй, але перевіряй», то ZTA базується на постулаті «ніколи не довіряй, завжди перевіряй», незалежно від того, де знаходиться користувач або пристрій - всередині чи зовні корпоративної мережі. Суть трансформації полягає у відмові від поняття «довіреної зони». У мережі Zero Trust кожен запит на доступ до ресурсу повинен проходити автентифікацію, авторизацію та шифрування в реальному часі. Це вимагає перебудови архітектури з використанням мікросегментації. Мікросегментація дозволяє розділити мережу на дрібні ізольовані зони аж до рівня окремих робочих

навантажень. Це критично важливо для обмеження радіусу атаки: якщо зломисник компрометує один сегмент, він не отримує автоматичного доступу до всієї інфраструктури. Ключовим елементом нової архітектури стає ідентифікація (Identity). У системах це означає інтеграцію потужних систем IAM (Identity and Access Management) з мережевим обладнанням. Рішення про надання доступу приймаються динамічно, базуючись не лише на паролі, а й на контексті: стан безпеки пристрою, геолокація, час доби та поведінкові аномалії.

Важливу роль відіграє технологія програмно-визначеного периметра (SDP — Software Defined Perimeter). Вона дозволяє приховувати мережеві ресурси від неавторизованих користувачів, роблячи інфраструктуру «невидимою» для сканування портів. З'єднання встановлюється лише після підтвердження довіри брокером доступу.

Однак перехід на Zero Trust супроводжується технічними викликами. Це потребує повної інвентаризації всіх активів мережі, що часто є складною задачею для великих підприємств із застарілим (legacy) обладнанням. Крім того, впровадження наскрізного шифрування та постійної перевірки політик створює додаткове навантаження на пропускну здатність мережі та вимагає оптимізації мережевих контролерів. Варто також зазначити, що архітектура Zero Trust тісно пов'язана з автоматизацією та оркестрацією безпеки (SOAR). Ручне керування політиками доступу в динамічному середовищі є неможливим, тому системи повинні автоматично реагувати на інциденти, ізолюючи підозрілі вузли без участі адміністратора.

Висновок. Таким чином, перехід до архітектури Zero Trust є еволюційною необхідністю для сучасних комп'ютерних систем та мереж. Ця модель дозволяє створити адаптивну систему захисту, яка відповідає вимогам децентралізованого ІТ-середовища. Попри складність реалізації, ZTA забезпечує значно вищий рівень стійкості мережі до зламів порівняно з периметральним захистом.

Перелік посилань.

1. NIST SP 800-207. Zero Trust Architecture. National Institute of Standards and Technology, 2020. - Офіційний стандарт, що визначає ключові принципи, логічні компоненти та моделі розгортання архітектури нульової довіри.
2. Kindervag, J. (2010). Build Security Into Your Network's DNA: The Zero Trust Network Architecture. Forrester Research. – Основоположна робота, в якій вперше було сформульовано концепцію та необхідність переходу до моделі Zero Trust.
3. Rose, S., et al. (2020). Zero Trust Architecture (ZTA) Guide. CISA. - Практичний посібник від Агентства кібербезпеки США щодо стратегій міграції від традиційних мереж до архітектури ZTA.
4. Buck, C., Olenberger, C., et al. (2021). Never Trust, Always Verify: A Survey of Zero Trust Networking Components and Solutions. IEEE Access. - Комплексний огляд технічних компонентів, протоколів та існуючих рішень на ринку для реалізації принципів нульової довіри.

5. Gilman, E., & Barth, D. (2017). Zero Trust Networks: Building Secure Systems in Untrusted Networks. O'Reilly Media. - Детальний технічний аналіз побудови захищених систем без використання традиційного периметра безпеки.

Звенигородський О.С., доцент ДУІКТ
Шаш М.С., аспірант ДУІКТ

МОДЕЛЬ АНАЛІЗУ ЮНІТ-ЕКОНОМІКИ SaaS ЗА ДОПОМОГОЮ ШТУЧНОГО ІНТЕЛЕКТУ

Вступ. Сучасні SaaS-компанії працюють у середовищі високої конкуренції, мінливої поведінки користувачів та тиску на ефективність кожного вкладеного долара. Традиційні підходи до аналітики часто дають лише ретроспективний погляд на бізнес, не дозволяючи передбачити майбутню динаміку зростання та ризику. У цьому контексті штучний інтелект стає ключовим інструментом для глибшого розуміння користувачів, виявлення закономірностей і побудови прогнозів, які безпосередньо впливають на юніт-економіку. Використання моделей ШІ у прогнозуванні LTV (Lifetime Value), CAC (Customer Acquisition Cost) та поведінкових сценаріїв дозволяє SaaS-продуктам ухвалювати більш точні рішення, оптимізувати залучення, покращувати утримання користувачів і досягати стійкого зростання.

Мета роботи. Мета роботи — розробити модель, що дозволяє за допомогою методів машинного навчання та ШІ прогнозувати ключові метрики юніт-економіки SaaS-продуктів, зокрема LTV, CAC та поведінкові патерни користувачів. Робота також спрямована на визначення груп користувачів із підвищеним ризиком відтоку, виявлення потенціалу для зростання та формування рекомендацій щодо оптимізації маркетингових, продуктових та операційних рішень на основі прогнозних даних.

Огляд літератури.

У сфері SaaS і підпискових бізнес-моделей зростає інтерес до застосування машинного навчання й ШІ для прогнозування показників юніт-економіки — таких як LTV, CAC, відтік (churn) і поведінкові патерни клієнтів. Дослідження A novel approach to predicting customer lifetime value in B2B SaaS companies (2023) пропонує гнучкий ML-фреймворк для прогнозування CLV (LTV) саме в B2B SaaS-середовищі. Воно виділяє такі проблеми, як гетерогенність клієнтів, мультипродуктовість, обмежені часові дані. Автори формують задачу як «lump-sum» прогноз майбутнього доходу і розробляють ієрархічну ансамблеву модель, яка показує кращу точність, ніж традиційні підходи та прості евристики [1]. [OBJ]

Зокрема, їх підхід дозволяє використовувати широкий набір фіч — від фінансових метрик (MRR, license data) до даних про використання продуктів та поведінку користувачів. Це відкриває можливості для гнучкої сегментації клієнтів,

обрання оптимального маркетингового бюджетування та планування ретеншн-стратегій на основі прогнозованого LTV. [ОБ]

Інше дослідження — Predicting Customer Lifetime Value Using Recurrent Neural Net (2024) — показує, що рекурентні нейронні мережі можуть давати кращу точність у прогнозуванні LTV для SaaS-додатків у порівнянні з класичними моделями (наприклад, LightGBM або простими CLV-моделями). Модель враховує такі часові виміри, як дата приєднання (kohort), «вік» користувача в системі та календарну дату, що дозволяє захоплювати динаміку поведінки користувачів з плином часу [2]. [ОБ]

Ще одне важливе дослідження — Using Machine Learning Algorithms to Analyze Customer Churn in the Software as a Service (SaaS) Industry (2022) — застосовує різні ML-алгоритми для аналізу відтоку клієнтів у SaaS. Вони виділили набір фіч, важливих для прогнозування churn, і проаналізували, як відтік впливає на життєву цінність клієнта. Ця робота демонструє, що аналіз churn — критичний компонент для оцінки реального LTV та для побудови стійкої юніт-економіки [3]. [ОБ]

Крім того, дослідження A data-driven approach to customer lifetime value prediction using probability and machine learning models (2025) працює над питанням, чи перевершують ML-підходи класичні ймовірнісні моделі CLV. Автори порівнюють обидва підходи й приходять до висновку, що в деяких випадках класичні моделі дають кращу точність за умови недостатньої історії чи нестабільної поведінки клієнтів, що показує межі та ризики застосування ML для прогнозування LTV[4].

Постановка задачі.

У межах дослідження формулюється така задача. Дано масив операційних даних SaaS-продукту з метрики залучення користувачів (активація, онбординг, частота використання, активні події), показники транзакцій та білінгу, дані про канали залучення й маркетингові витрати, історія відтоку й оновлень, а також дані про поведінкові патерни користувачів у продукті.

Необхідно побудувати модель, яка:

- прогнозує LTV користувачів з різних сегментів, враховуючи зміну поведінки, цикли використання та вплив маркетингових активностей;
- моделює та прогнозує SAC, у тому числі з урахуванням динаміки каналів, сезонності та ефективності кампаній;
- ідентифікує користувачів із підвищеним ризиком відтоку або низької монетизації, щоб SaaS-компанія могла запускати таргетовані втручання;
- визначає ключові фактори, що формують юніт-економіку, і дозволяє продукту й маркетингу приймати обґрунтовані рішення щодо оптимізації ціни, сегментації, каналів і пріоритетів розвитку.

Методи та модель.

Джерела даних у SaaS-компанії:

- Аналітика продукту / використання сервісу — активність користувачів, час сесій, кількість запусків функцій, взаємодії з контентом, частота логінів.
- Фінансові дані — доходи від користувача (MRR, ARR), історія платежів, підписки на різні пакети, знижки, акційні пропозиції.
- Маркетингові дані — канали залучення, вартість залучення клієнта (CAC), рекламні кампанії, кліки, конверсії.
- Демографічні / фонові дані — країна, сегмент клієнта, тип організації (B2B/B2C), розмір компанії, галузь.
- Поведінкові дані — реакції на оновлення продукту, відгуки, підтримка/тикети, churn.

Етапи очищення та підготовки даних

1. Очищення даних: усунення дублікатів, виправлення помилок у записах, нормалізація форматів (дата, валюта, категорії).
2. Обробка пропусків: аналіз причин відсутніх даних, інтерполяція, категорія “невідомо” або виключення рядків за потреби.
3. Стандартизація та нормалізація: фінансові показники, метрики активності та поведінки користувачів приводяться до єдиної шкали для порівнянності.
4. Вирівнювання даних у часі: узгодження часових точок для прогнозування (наприклад, щомісячні, квартальні або річні періоди).

Вибір ознак (features)

- Фінансові: середній дохід на користувача (ARPU), LTV, CAC, частота платежів.
- Поведінкові: час активності, кількість використаних функцій, повторні сесії, engagement score.
- Маркетингові / джерела трафіку: канали залучення, активація користувача, конверсія з пробної версії в платну.
- Демографічні / сегментаційні: тип підписки, розмір компанії, ринок, географія.
- Історичні та часові ряди: показники по місяцях або кварталах, поведінка перед відтоком, патерни апгрейдів/понижень тарифів.

Моделі прогнозування

- Регресія: передбачення LTV, CAC, доходу на користувача; оцінка майбутніх фінансових показників.
- Класифікація: передбачення відтоку (churn) — поділ користувачів на “високий ризик”, “середній ризик”, “низький ризик”; сегментація клієнтів за поведінкою.
- Методи машинного навчання: Random Forest, Gradient Boosting (XGBoost, LightGBM), нейронні мережі (особливо для великих наборів поведінкових даних), SVM.
- Моделі часових рядів / послідовностей: рекурентні нейронні мережі (RNN, LSTM) для прогнозування LTV та CAC на основі поведінки користувачів за часом.

- Інтерпретованість: Explainable AI (SHAP, LIME) для визначення найважливіших ознак, що впливають на LTV, CAC і відтік, що критично для прийняття бізнес-рішень.

Метрики оцінки моделей

- Регресія: RMSE, MAE, R^2 .
- Класифікація: Accuracy, Precision, Recall, F1-score, ROC-AUC.
- Додатково: тестування на “реальному часі” — прогнозування LTV і CAC на основі актуальних даних користувачів для оцінки моделі у продуктивному середовищі.

Застосування результатів. Результати моделей ШІ у прогнозуванні LTV, CAC і поведінки користувачів у SaaS застосовуються для вирішення кількох ключових задач. По-перше, вони дозволяють прогнозувати життєву цінність клієнта та вартість його залучення на ранніх етапах взаємодії, що допомагає оптимізувати маркетингові витрати та стратегії залучення. По-друге, моделі виявляють користувачів із високим ризиком відтоку, аналізують ключові фактори, які впливають на прибутковість та лояльність клієнтів — наприклад, низьку активність у продукті або недостатню взаємодію з функціями. На основі цих прогнозів можна давати рекомендації бізнесу та продуктовому менеджменту щодо коригування продукту, тарифних планів, маркетингових кампаній або персоналізованих пропозицій, а також підтримувати стратегічне фінансове планування і оцінку рентабельності різних сегментів клієнтів.

Висновки. Застосування ШІ у прогнозуванні юніт-економіки SaaS має великий потенціал для підвищення ефективності бізнесу. Моделі дозволяють не лише вимірювати показники, а й передбачати розвиток користувачів та ефективність маркетингових заходів, що сприяє оптимізації LTV і зниженню CAC. Ключовими факторами успіху є добір якісних та актуальних даних, регулярне оновлення моделей та інтерпретованість прогнозів, щоб менеджмент розумів причини певних рішень моделі. Важливими є також етичні та практичні аспекти, такі як захист конфіденційності даних клієнтів і прозорість методів прогнозування. Завдяки цьому SaaS-компанії можуть підвищувати стійкість бізнесу, зменшувати відтік клієнтів та забезпечувати рентабельне зростання.

Список використаних джерел

1. Curiskis, S., Dong, X., Jiang, F. *et al.* A novel approach to predicting customer lifetime value in B2B SaaS companies. *J Market Anal* 11, 587–601 (2023). <https://doi.org/10.1057/s41270-023-00234-6>
2. Chen, H., Ng, E., Smyl, S., & Steininger, G. Predicting Customer Lifetime Value Using Recurrent Neural Networks. *arXiv* (2024). <https://doi.org/10.48550/arXiv.2412.20295>
3. L. Çallı and S. Kasım, “Using Machine Learning Algorithms to Analyze Customer Churn in the Software as a Service (SaaS) Industry”, *APJESS*, vol. 10, no. 3, pp. 115–123, 2022, doi: 10.21541/apjess.1139862.
4. Wong, A., Vilorio Garcia, A., & Lim, Y.-W. (2025). A data-driven approach to customer lifetime value prediction using probability and machine learning

Серих С.О., доцент ДУІКТ
Попов А.О., аспірант ДУІКТ

ВИКОРИСТАННЯ CROSS-LINGUAL MULTI-SPEAKER TTS ТА VOICE CONVERSION ДЛЯ АУГМЕНТАЦІЇ ДАНИХ У ЗАДАЧАХ РОЗПІЗНАВАННЯ МОВИ

У статті розглянуто застосування багатомовного багатоголосового синтезу мовлення (TTS) та голосової конверсії (VC) для створення великих та різноманітних наборів аудіоданих у низькоресурсних мовах з метою підвищення точності систем автоматичного розпізнавання мовлення (ASR).

Постановка задачі

Необхідно дослідити можливості використання cross-lingual TTS та voice conversion для генерування синтетичних мовних даних у ситуаціях, коли кількість реальних аудіозаписів є обмеженою. Також важливо визначити ефективність такої аугментації для покращення роботи систем ASR у low-resource умовах.

Мета дослідження

Метою дослідження є проаналізувати потенціал синтетичної генерації мовлення й голосової конверсії для збільшення обсягу та різноманітності тренувальних даних, а також оцінити, наскільки такі методи можуть зменшити залежність ASR-систем від великих людських наборів даних.

Результати дослідження

Сучасні системи розпізнавання мовлення потребують значних обсягів аудіо даних, однак багато мов залишаються низько ресурсними, що ускладнює розробку якісних моделей. Стаття демонструє, що використання cross-lingual multi-speaker TTS у поєднанні з voice conversion дає можливість ефективно розширювати тренувальні корпуси навіть за наявності лише одного реального голосового прикладу цільової мови або спікера. Синтетичні дані забезпечують різноманітність голосів, інтонацій та умов запису, що покращує узагальнювальні властивості моделей ASR. Застосування голосової конверсії дозволяє адаптувати вже наявні аудіозаписи до потрібного голосу або мови, не вимагаючи масштабного збирання нових даних. Дослідження показує, що комбінування реальних аудіозаписів навіть у невеликій кількості з великою кількістю синтетичних зразків суттєво підвищує точність ASR-моделей. Водночас якість синтезованого мовлення має ключове значення: наявність артефактів або низька природність може зменшувати користь аугментації. Системи такого типу потребують адаптації до конкретної мови та варіантів вимови, а також здатності швидко оновлюватися відповідно до зміни мовних або акустичних умов. Однією з перспективних галузей дослідження є створення моделей, що здатні до ефективного до навчання з мінімальною кількістю даних.

Висновки та перспективи

Використання cross-lingual multi-speaker TTS і voice conversion є перспективним напрямом для розв'язання проблеми нестачі мовних даних у задачах автоматичного розпізнавання мовлення. Подальші дослідження мають бути спрямовані на підвищення природності синтезованого мовлення, зменшення артефактів, а також на створення систем, що можуть гнучко адаптуватися до конкретних мов та умов без потреби у масштабних повторних тренуваннях. Такі методи не зможуть повністю замінити реальні дані, проте суттєво зменшують їхню необхідну кількість та підвищують ефективність розробки ASR-моделей у low-resource сценаріях, що відкриває нові можливості для впровадження технологій розпізнавання мовлення у широкому спектрі мов світу.

1. Cross-speaker style transfer for text-to-speech using data augmentation
https://arxiv.org/abs/2202.05083?utm_source=chatgpt.com

2. Cross-lingual multi-speaker speech synthesis with limited bilingual training data
https://www.sciencedirect.com/science/article/abs/pii/S0885230822000584?utm_source=chatgpt.com

3. An Exhaustive Evaluation of TTS- and VC-based Data Augmentation for ASR
https://arxiv.org/abs/2503.08954?utm_source=chatgpt.com

Серих С.О., доцент ДУІКТ

Попов А.О., аспірант ДУІКТ

ПОПЕРЕДНЄ НАВЧАННЯ БЕЗ СПОСТЕРЕЖЕННЯ ДЛЯ ЕФЕКТИВНОГО СИНТЕЗУ МОВЛЕННЯ В УМОВАХ НИЗЬКОРЕСУРСНИХ МОВ

У статті розглядається методика **unsupervised pre-training** для моделей синтезу мовлення (TTS), що дозволяє значно зменшити обсяг необхідних розмічених даних для тренування моделей у низькоресурсних мовах.

Постановка задачі

Необхідно розробити підхід, який зменшує потребу у великих аудіо-текстових датасетах для навчання моделей TTS, оскільки більшість мов світу є низькоресурсними. Також важливо дослідити, чи може модель, тренувана на нерозмічених аудіо інших мов, покращити якість синтезу цільової мови.

Мета дослідження

Метою дослідження є оцінити ефективність несупервізованого попереднього навчання на великих масивах “дикого” аудіо без транскрипцій, а також проаналізувати, наскільки цей підхід скорочує необхідність у великих обсягах пар “аудіо-текст” для подальшого донавчання TTS-моделі.

Результати дослідження

Сучасні TTS-системи вимагають значних обсягів розмічених даних для досягнення природного та стабільного синтезу мовлення, однак для більшості мов світу такі ресурси є обмеженими. У своїй роботі Park та співавтори запропонували двоетапний підхід, що дозволяє суттєво зменшити потребу в аудіо-текстових парах. На етапі несупервізованого попереднього навчання модель тренується на великих корпусах нерозміченого мовлення, навчаючись реконструювати та «розтягувати» (speech de-warping) сигнал таким чином, щоб виокремлювати універсальні мовні ознаки, спільні для різних мов. На етапі супервізованого донавчання система адаптується до конкретної low-resource мови, використовуючи лише невеликий розмічений датасет, який інколи обмежується кількома годинами аудіо. Такий підхід демонструє високу ефективність, оскільки попередньо навчена модель значно краще відділяє фонетичні характеристики від шуму, тембральних особливостей та просодичних варіацій, що у підсумку підвищує здатність до засвоєння нової мови навіть за мінімальної кількості даних. Експериментальні результати показують суттєве покращення якості синтезу після fine-tuning навіть на 1–3 годинах даних, стабільнішу вимову, природнішу інтонацію та зменшення характерних для low-resource систем артефактів — таких як ритмічні порушення, нечіткі фонемі чи «пливучі» форманти. Додатково відзначається покращення загальної мовленнєвої узгодженості, що пов'язано з формуванням універсальних фонетичних репрезентацій під час попереднього навчання. Основна ідея роботи полягає в тому, що мовні системи мають низку універсальних структурних властивостей, які модель здатна засвоїти заздалегідь, ще до знайомства з конкретною цільовою мовою. Це дозволяє знизити потребу в ресурсах, зробивши синтез мовлення доступним навіть для мов із дуже малими корпусами. Несупервізоване попереднє навчання вже сьогодні демонструє значний потенціал у скороченні залежності TTS-моделей від дорогих і складних у збиранні аудіо-текстових даних. Перспективи подальших досліджень полягають у вдосконаленні архітектур для більш глибокого вилучення фонетичних ознак, адаптації під тональні чи фонетично відмінні мови, розробці більш ефективних self-supervised технологій, що зменшуватимуть потребу у текстовій розмітці, а також у комбінуванні попереднього несупервізованого навчання з методами аугментації, голосової конверсії та міжмовного переносу. Сукупність цих напрямів відкриває шлях до створення високоякісних систем синтезу мовлення для практично будь-якої мови, незалежно від рівня її ресурсної забезпеченості.

Висновки та перспективи

Несупервізоване попереднє навчання є перспективним напрямом для зменшення залежності TTS-моделей від великих парних аудіо-текстових корпусів, а отримані результати демонструють, що такі моделі можуть швидко й ефективно адаптуватися до нових мов навіть за мінімальної кількості даних. Подальші дослідження в цій галузі можуть фокусуватися на покращенні архітектур, які забезпечуватимуть глибше вилучення фонетичних ознак, а також на розширенні підходу для тональних мов або мов із суттєво відмінною фонетичною структурою. Важливим напрямом є розробка більш ефективних self-supervised методів, що дозволятимуть зменшити потребу у великих текстових наборах для етапу fine-tuning. Додатковий потенціал відкриває поєднання unsupervised pre-training із

техніками аугментації даних, голосової конверсії та міжмовного переносу. У сукупності це створює можливість розробляти високоякісні системи синтезу мовлення практично для будь-якої мови, навіть якщо вона має обмежений мовний корпус.

1. Low-Resource Expressive Text-To-Speech Using Data Augmentation https://openreview.net/forum?id=FxCealv5sa&utm_source=chatgpt.com
2. Text-to-speech system for low-resource language using cross-lingual transfer learning and data augmentation https://link.springer.com/article/10.1186/s13636-021-00225-4?utm_source=chatgpt.com
3. Advancing Text-to-Speech Systems for Low-Resource Languages: Challenges, Innovations and Future Directions https://www.researchgate.net/publication/395230321_Advancing_Text-to-Speech_Systems_for_Low-Resource_Languages_Challenges_Innovations_and_Future_Directions?utm_source=chatgpt.com

Шикула О.М., професор ДУІКТ
Коник С. П., студент ДУІКТ

ТЕЛЕГРАМ ЧАТ-БОТ НА МОВІ PYTHON З ІНТЕГРАЦІЄЮ В SRM СИСТЕМУ

Постановка задачі

У сучасних умовах автоматизації бізнес-процесів телеграм чат-боти стають одним із найефективніших інструментів комунікації між компанією та клієнтами. Вони дозволяють забезпечити швидкий і зручний зв'язок без участі людини, що значно скорочує час обробки запитів і підвищує рівень задоволеності користувачів. Завдяки використанню штучного інтелекту та сучасних алгоритмів обробки повідомлень, чат-боти можуть не лише відповідати на типові запитання, а й виконувати складні завдання, наприклад, приймати замовлення, перевіряти статуси, надсилати нагадування або пропонувати індивідуальні рішення.

Інтеграція чат-бота з CRM системою відкриває нові можливості для ефективного управління клієнтськими даними. Така взаємодія дозволяє не лише автоматизувати процес обробки звернень користувачів, а й зберігати всю історію комунікацій у єдиній централізованій базі даних. Це забезпечує повний контроль над взаємодією з клієнтами, підвищує точність аналітики та сприяє формуванню персоналізованих підходів до обслуговування.

Таким чином, поєднання функціоналу Telegram чат-бота та CRM системи створює єдину цифрову платформу, здатну ефективно вирішувати завдання бізнес-автоматизації. Такий підхід підвищує рівень обслуговування клієнтів, зменшує навантаження на менеджерів, оптимізує процеси збору та аналізу інформації, а також дозволяє компаніям швидше реагувати на потреби ринку.

Мета дослідження

Метою роботи є розробка телеграм чат-бота, який автоматично взаємодіє з CRM системою для збору, обробки та збереження даних клієнтів. Бот дозволить оптимізувати бізнес-процеси, покращити комунікацію з клієнтами та підвищити ефективність управління.

Розробка цього проекту спрямована на демонстрацію можливостей інтеграції месенджер-платформ із корпоративними системами управління. Це дозволить компаніям спростити процес роботи з клієнтами, зменшити навантаження на працівників підтримки та забезпечити швидкий обмін інформацією між користувачами й CRM. Таким чином, чат-бот стане не лише інструментом спілкування, а й важливим елементом автоматизації бізнес-процесів.

Результати дослідження

Для розробки чат-бота використано мову програмування Python, оскільки вона має велику кількість бібліотек для роботи з API та зручну структуру для побудови серверної логіки. Для взаємодії з Telegram застосовується офіційний Telegram Bot API у поєднанні з фреймворком Aiogram, який підтримує асинхронну обробку повідомлень і забезпечує високу швидкість роботи. Інтеграція з CRM системою реалізована через REST API, що дозволяє виконувати операції зі збереження, оновлення та отримання даних у режимі реального часу. Для зберігання інформації використовується реляційна база даних, взаємодія з якою здійснюється через ORM, що спрощує роботу з моделями та зменшує кількість помилок. Такий технологічний стек робить систему масштабованою, надійною та придатною до подальшого розвитку.

У результаті дослідження створено повнофункціональний Telegram чат-бот, інтегрований із CRM системою, який забезпечує автоматизовану взаємодію між компанією та клієнтами. Розроблений бот дозволяє збирати та обробляти дані клієнтів, автоматично зберігаючи їх у CRM системі. Він забезпечує можливість вести облік користувачів, формувати їхні профілі та відстежувати історію звернень.

Бот автоматично інформує клієнтів про оновлення, акції, нагадування або спеціальні пропозиції. Завдяки інтерактивній взаємодії чат-бот виконує запити користувачів, надає консультації та приймає замовлення, роблячи процес комунікації швидким і зрозумілим.

Система має гнучку структуру, дозволяє адаптувати її під потреби різних бізнесів і легко розширювати функціональні можливості. У CRM автоматично формуються звіти про активність користувачів, система управління проводить аналітику й приймає обґрунтовані рішення.

Висновки та перспективи

Результатом виконаної роботи є розробка телеграм чат-бота з інтеграцією в CRM систему, який прагне підвищити ефективність роботи з клієнтами. Такий підхід поєднує автоматизацію, аналітику та зручність комунікації, що робить його незамінним інструментом у сучасному цифровому середовищі.

Список використаних джерел

1. Telegram API. – [Електронний ресурс]. – URL: <https://telegra.ph/api>
2. Aiogram Documentation. – [Електронний ресурс]. – URL: <https://docs.aiogram.dev/>
3. CRM Integration Guide. – [Електронний ресурс]. – URL: <https://zapier.com/apps/telegram>
4. Приклади чат-ботів на Python. – [Електронний ресурс]. – URL: <https://github.com/topics/pythonchatbot>

Шикула О.М., професор ДУІКТ

Бугайченко А.Д., студент ДУІКТ

ІНФОРМАЦІЙНА СИСТЕМА ПРИЙОМУ ТА АНАЛІЗУ ЗАЯВОК ТЕХНІЧНОЇ ПІДТРИМКИ (Python/Django + Bootstrap + PostgreSQL)

Постановка задачі

У сучасних умовах ефективна робота служби технічної підтримки є критично важливою для забезпечення безперебійного функціонування будь-якої організації. Відсутність централізованої системи призводить до втрати заявок, ускладнення аналізу роботи фахівців та зниження загальної якості сервісу. Створення спеціалізованого програмного рішення, яке поєднує оперативну обробку запитів з глибоким аналітичним функціоналом, дозволяє автоматизувати рутинні операції, забезпечити прозорість усіх етапів роботи та накопичувати структуровані дані для прийняття управлінських рішень. Таким чином, розробка інформаційної системи, що інтегрує функції обліку заявок та аналітики, є актуальним завданням для підвищення ефективності ІТ-сервісів.

Мета дослідження

Метою роботи була розробка повнофункціональної інформаційної системи для автоматизації процесів прийому, обробки та аналізу заявок технічної підтримки. Система мала забезпечити зручний інтерфейс для користувачів та фахівців, автоматизувати рутинні операції з розподілу та супроводу заявок, а також накопичувати дані для аналітики. Крім того, метою була демонстрація ефективності використання сучасного веб-стеку (Python/Django, Bootstrap, PostgreSQL) для створення практичних, масштабованих рішень у сфері управління бізнес-процесами.

Результати дослідження

У результаті дослідження було створено повнофункціональну веб-систему. Система реалізує зручну форму для подачі заявок користувачами з автоматичним присвоєнням унікального номера. Для фахівців техпідтримки розроблено панель управління з переліком заявок, можливістю зміни статусів, пріоритетів та призначення відповідальних осіб. Реалізовано систему облікових записів із різними рівнями доступу (користувач, фахівець, адміністратор). Веб-інтерфейс

створено з використанням фреймворку Bootstrap, що забезпечує його адаптивність та зручність використання. Для зберігання даних розгорнуто базу даних PostgreSQL. Ключовою особливістю системи є вбудований аналітичний модуль, який дозволяє будувати звіти щодо навантаження на фахівців, аналізувати статистику вирішених заявок, час їх вирішення та виявляти типові проблеми. Наукова новизна роботи полягає в глибокій інтеграції механізмів аналітики в ядро системи обліку заявок.

Висновки та перспективи

Розроблена інформаційна система на основі Python/Django, Bootstrap та PostgreSQL є ефективним інструментом для організації роботи служби технічної підтримки. Вона не лише успішно автоматизує ключові процеси прийому та обробки заявок, але й завдяки вбудованим засобам аналітики надає цінну інформацію для покращення якості сервісу та оптимізації ресурсів. Отриманий результат є готовим до впровадження рішенням для середніх та великих організацій. Перспективою подальшого розвитку системи є впровадження додаткових функцій, таких як система сповіщень по електронній пошті, інтеграція з системами моніторингу та впровадження механізмів машинного навчання для передбачення типових проблем.

Список використаних джерел

1. Офіційна документація Django. – [Електронний ресурс]. – URL: <https://docs.djangoproject.com/>
2. Офіційна документація Bootstrap. – [Електронний ресурс]. – URL: <https://getbootstrap.com/docs/>
3. Офіційна документація PostgreSQL. – [Електронний ресурс]. – URL: <https://www.postgresql.org/docs/>
4. ITIL Foundation: ITIL 4 Edition. – AXELOS, 2019. – 214 p.

Шикула О.М., професор ДУІКТ
Івасюк А.І., студент ДУІКТ

ВЕБ-ЗАСТОСУНОК ДЛЯ МОНІТОРИНГУ ПОГОДНИХ ТА ЕКОЛОГІЧНИХ ПОКАЗНИКІВ НА MOBI JAVASCRIPT З ВИКОРИСТАННЯМ ФРЕЙМВОРКУ VUE.JS

Постановка задачі

На сучасному етапі глобальних кліматичних змін та зростаючої уваги до екологічної безпеки проблема оперативного та зручного моніторингу погодних і екологічних показників є вкрай актуальною. Існуючі системи часто є громіздкими, мають застарілий інтерфейс або обмежені можливості візуалізації даних. Розроблення сучасного, інтерактивного веб-застосунку на базі фреймворку Vue.js дозволить створити ефективний інструмент для швидкого доступу, аналізу та візуалізації цих критично важливих даних, сприяючи прийняттю обґрунтованих

рішень у різних сферах (сільське господарство, містобудування, екологічний контроль).

Мета дослідження

Метою даного дослідження є розробка функціонального та адаптивного веб-застосунку для моніторингу погодних та екологічних показників з використанням сучасного JavaScript-фреймворку Vue.js для клієнтської частини. Застосунок повинен забезпечувати ефективне збирання, обробку, наочну візуалізацію даних та створювати зручний канал комунікації з користувачами-аналітиками.

Результати дослідження

Проведено аналіз предметної області, визначено ключові погодні та екологічні показники, необхідні для моніторингу, та сформовано функціональні вимоги до веб-застосунку.

Обґрунтовано вибір технологій: Для розроблення клієнтської частини (Frontend) обрано фреймворк Vue.js (Vue 3) завдяки його реактивності, компонентній архітектурі та високій продуктивності. Обґрунтовано відмову від розробки власного Backend-сервісу на користь прямого використання публічних API для отримання даних.

Реалізовано інтеграцію з зовнішніми сервісами. Здійснено налаштування асинхронних запитів з клієнтської частини Vue.js до відкритих погодних та екологічних API для отримання актуальних даних у форматі JSON.

Висновки та перспективи

Створено інтерактивний інтерфейс за допомогою Vue.js. Розроблено багаторазові компоненти Vue для відображення ключових показників (температура, якість повітря) у вигляді карт, індикаторів та графіків.

Впроваджено функціонал розширеного аналізу: Додано можливість аналізу історичних даних, отриманих через API, з метою визначення:

- найхолоднішого та найтеплішого дня у вибраний період;
- середньої температури за конкретну дату чи період.

Впроваджено функціонал інтелектуального аналізу за допомогою Gemini API:

- Реалізовано модуль для надсилення даних про поточні та прогнозовані температури, а також рівень забрудненості повітря.
- Забезпечено отримання та відображення персоналізованих, згенерованих ШІ-рекомендацій для користувача, включаючи поради "як одягатись" залежно від погодних умов та пояснення і рекомендації щодо поведінки при поточному рівні забрудненості повітря (наприклад, "обмежити фізичну активність на вулиці").

Проведене тестування підтвердило стабільність, адаптивність застосунку та коректність відображення даних, отриманих через зовнішні API.

Список використаних джерел

1. Vue.js 3 Documentation. – [Електронний ресурс]. – URL: <https://vuejs.org/>
2. The Modern JavaScript Tutorial. – [Електронний ресурс]. – URL: <https://javascript.info/>

3. Chart.js Documentation. – [Електронний ресурс]. – URL: <https://www.chartjs.org/docs/latest/>

4. OpenWeatherMap API Documentation. – [Електронний ресурс]. – URL: <https://openweathermap.org/api>

Шикула О.М., професор ДУІКТ
Борода К.О., студент ДУІКТ

ЧАТ-БОТ-МАГАЗИНУ В TELEGRAM НА МОВІ PYTHON ІЗ ВИКОРИСТАННЯМ СКБД MYSQL

Постановка задачі

У сучасних умовах цифрової трансформації бізнесу все більшої популярності набувають автоматизовані системи взаємодії з клієнтами. Особливе місце серед них посідають чат-боти, що інтегруються у популярні месенджери. Telegram завдяки великій аудиторії, простоті використання та розвиненим можливостям API є ефективним інструментом для реалізації комерційних сервісів, зокрема інтернет-магазинів.

Попит на інтелектуальні системи онлайн-продажу зумовлює потребу у створенні рішень, які дозволяють бізнесу автоматизувати комунікації з покупцями, оптимізувати процеси формування замовлень, контролювати товарні залишки та вести базу клієнтів. Проте, незважаючи на значну кількість існуючих бот-платформ, багато з них є обмеженими або не дозволяють масштабувати систему відповідно до потреб бізнесу.

Отже, виникає необхідність у створенні Telegram чат-бота з повноцінним функціоналом магазину, інтегрованого з реляційною базою даних MySQL, що забезпечуватиме надійне зберігання інформації, обробку замовлень та можливість подальшого розвитку системи. Актуальність теми визначається потребою у сучасних інструментах електронної комерції та автоматизації сервісів клієнтської підтримки.

Мета дослідження

Метою дослідження є розробка телеграм чат-бота-магазину, здатного забезпечувати повний цикл взаємодії з користувачем: від реєстрації до оформлення та обробки замовлення. Основні завдання, що визначають досягнення поставленої мети, включають:

- розроблення архітектури системи та визначення функціональних компонентів чат-бота;
- створення системи збереження даних на основі реляційної бази MySQL;
- реалізацію можливості перегляду каталогу товарів, сортування та пошуку;
- забезпечення функціоналу кошика та оформлення замовлення;
- розробку панелі адміністратора для роботи з товарами, категоріями та замовленнями;
- тестування роботи системи та аналіз її ефективності.

Таким чином, мета роботи полягає у створенні інструменту, який забезпечує автоматизацію торговельних процесів та підвищує якість клієнтського обслуговування.

Результати дослідження

У процесі виконання роботи проведено аналіз існуючих методів реалізації чат-ботів, особливостей Telegram API та бібліотеки Aiogram. На основі отриманих даних було спроектовано структуру бази даних MySQL, яка включає таблиці користувачів, товарів, категорій, кошика, замовлень та статусів.

Розроблено архітектуру Telegram-бота, що передбачає розподіл функціональних модулів: модуль обробки команд користувачів, модуль адміністрування, система авторизації та логування, модуль взаємодії з базою даних. Реалізовано функціонал, що дозволяє користувачеві переглядати товари, додавати їх до кошика, редагувати вміст кошика, оформлювати замовлення та переглядати інформацію про покупки.

Для адміністратора передбачено можливість роботи з каталогом товарів, перегляду замовлень, керування статусами та перегляду бази користувачів. Проведене тестування підтвердило працездатність системи, стабільність її роботи та коректність функціонування інтерактивних елементів.

Отриманий програмний продукт може використовуватися малим та середнім бізнесом для автоматизації продажів та підвищення якості обслуговування клієнтів.

Висновки та перспективи

Розроблений Telegram чат-бот-магазин демонструє можливість ефективного застосування інструментів Python, Aiogram та MySQL для створення сучасної системи електронної комерції. Реалізоване рішення забезпечує автоматизацію процесу прийому й оброблення замовлень, надає інструменти адміністрування та дозволяє зберігати дані у структурованому вигляді.

Практичне значення роботи полягає в тому, що розроблений сервіс може бути використаний для покращення взаємодії з клієнтами, оптимізації бізнес-процесів та розширення каналів збуту.

Подальші перспективи розвитку проекту передбачають:

- створення веб-інтерфейсу адміністратора;
- розширення можливостей аналітики та статистики;
- впровадження механізмів кешування та оптимізації продуктивності.

Таким чином, реалізована система є фундаментом для розбудови повноцінної платформи електронної комерції на базі Telegram-бота.

Список використаних джерел

1. Офіційна документація Telegram Bot API. – URL: <https://core.telegram.org/bots>
2. Aiogram (офіційна документація). – URL: <https://docs.aiogram.dev/>
3. MySQL (офіційна документація). – URL: <https://dev.mysql.com/doc/>
4. mysql-connector-python. – URL: <https://dev.mysql.com/doc/connector-python/en/>
5. SQLAlchemy (ORM для Python). – URL: <https://www.sqlalchemy.org/>

Шикула О.М., професор ДУІКТ
Кирилов Д.В., студент ДУІКТ

РОЗРОБКА САЙТУ КІННОГО КЛУБУ НА МОВАХ ПРОГРАМУВАННЯ JAVASCRIPT ТА PHP З ВИКОРИСТАННЯМ СКБД MYSQL

Постановка задачі

На сучасному етапі розвитку ринкових відносин проблема організації та просування послуг кінних клубів залишається недостатньо дослідженою. У науковій літературі відсутні комплексні дослідження, присвячені вивченню маркетингових стратегій, інструментів просування та організаційних особливостей функціонування підприємств кінної галузі. Наявні роботи здебільшого розглядають окремі аспекти спортивного маркетингу або загальні питання управління послугами, не враховуючи специфіку кінного спорту як особливої сфери дозвілля, спорту та туризму. Саме тому обрана тема дослідження є *актуальною*, адже її розроблення сприятиме удосконаленню теоретичних підходів і практичних механізмів управління маркетинговою діяльністю кінних клубів.

Мета дослідження

Метою даного дослідження є розробка сайту кінного клубу, який виконуватиме функцію візитної картки організації, забезпечуватиме ефективне представлення її діяльності у мережі Інтернет, створюватиме зручний канал комунікації з користувачами та сприятиме залученню нових членів клубу.

Результати дослідження

У процесі дослідження проведено ретельний огляд та аналіз тематичної області дослідження: проаналізовано діяльність кінних клубів, їх сучасний стан і перспективи розвитку, визначено функціональні вимоги до сайту кінного клубу.

Зроблено детальний аналіз популярних технологій та інструментів для створення сайтів. Було детально обґрунтовано вибір технологій HTML, CSS і JavaScript для розроблення клієнтської частини веб-додатку, а також PHP і MySQL – для створення серверної частини. Таке поєднання інструментів забезпечує високу продуктивність роботи програми, стабільність її функціонування та зручність у використанні як для розробників, так і для кінцевих користувачів. HTML5 і CSS3 виступають основою сучасної веб-розробки, адже саме вони відповідають за структуру, логіку розташування елементів і візуальне оформлення веб-сторінок. JavaScript – це потужна мова програмування, яка забезпечує динамічність і інтерактивність веб-сторінок. Вона дає змогу змінювати вміст сторінки без її перезавантаження, реагувати на дії користувача, виконувати асинхронні запити до сервера та створювати більш живий і зручний інтерфейс. Для реалізації серверної логіки використовується PHP – мова програмування, призначена для створення динамічних веб-сайтів і автоматизації оброблення даних. MySQL – це система управління базами даних, призначена для створення, впорядкування та ефективного керування інформацією. Вона надає розширений набір засобів, які дозволяють зберігати, редагувати,

аналізувати та швидко знаходити необхідні дані, що особливо важливо під час розробки та підтримки веб-проектів.

За допомогою цих технологій було створено сайт кінного клубу. Під час розробки було виконано аналіз потреб користувачів і визначено основні завдання сайту – представлення діяльності клубу, інформування клієнтів про послуги, новини та акції, а також забезпечення швидкого зворотного зв'язку. На основі цього було створено структуру сайту, яка складається з головної сторінки, розділів «Про нас», «Ціни», «Новини» та «Контакти».

На етапі проектування розроблено макет дизайну сайту з урахуванням принципів юзабіліті та адаптивності. У процесі реалізації створено шаблонну сторінку з чітким розподілом структурних елементів – заголовка, основного контентного блоку, нижньої панелі та нижнього колонтитулу. Окремо розроблено навігаційну систему, яка забезпечує простоту переходів і логічну взаємодію між сторінками.

Особливу увагу приділено візуальній складовій сайту – використано контрастну, проте гармонійну колірну гаму, зображення в тематичному стилі та логотип клубу. Таке оформлення створює цілісний візуальний образ і підсилює впізнаваність бренду.

Проведене тестування підтвердило, що сайт функціонує стабільно, коректно відображається у різних браузерах і на пристроях із різною роздільною здатністю екрана. Усі інтерактивні елементи працюють без збоїв, а швидкість завантаження сторінок відповідає сучасним стандартам.

Висновки та перспективи

Результатом виконаної роботи є функціональний, інформативний та естетично привабливий сайт кінного клубу, що може використовуватись як ефективний інструмент для популяризації клубу, комунікації з клієнтами та розширення його присутності в інтернет-просторі.

Подальше вдосконалення ресурсу може передбачати інтеграцію системи онлайн-запису на заняття, можливість бронювання послуг через сайт, а також розширення мультимедійних можливостей для залучення нових користувачів.

Список використаних джерел

1. Мальчук Е.В. HTML и CSS. Самоучитель. М.: Издательский дом Вильямс”. – 2018. – 416 с.
2. Флэнаган Д. JavaScript: Посібник. Санкт-Петербург: Символ Плюс, 2008. 992 с.
3. Локхарт Дж. Сучасний PHP: нові можливості та хороші практики: Пер. с англ. – СПб.: Символ-Плюс, 2016. – 270с.
4. Welling L., Thomson L. PHP and MySQL Web Development. – Boston: Addison-Wesley Professional, 2016. – 688с.

Шикула О.М., професор ДУІКТ
Лоскучерява С.С., студент ДУІКТ

РОЗРОБКА САЙТУ ФІРМИ З ОРГАНІЗАЦІЇ СВЯТКОВИХ ЗАХОДІВ ДЛЯ ДІТЕЙ НА МОВАХ ПРОГРАМУВАННЯ JAVASCRIPT ТА PHP З ВИКОРИСТАННЯМ СКБД MYSQL

Постановка задачі

Особливого значення набуває створення веб-сайтів для підприємств сфери послуг, адже саме вони потребують постійного контакту з клієнтами. Зокрема, фірми, що займаються організацією святкових заходів, повинні швидко, яскраво й ефективно презентувати свої послуги, демонструючи професійність, креативність і відкритість до нових ідей. У цьому контексті сайт-візитка виступає не лише інструментом реклами, а й ключовим елементом корпоративного іміджу, який формує довіру потенційних клієнтів. Уміння створювати ефективні, естетично приємні та технічно грамотні веб-ресурси сьогодні є важливим показником конкурентоспроможності будь-якої компанії.

Саме тому обрана тема дослідження є актуальною, адже розробка сучасного, зручного та функціонального веб-ресурсу забезпечить фірмі ефективну присутність в інформаційному просторі, дозволить оптимізувати процеси взаємодії з клієнтами та підвищити конкурентоспроможність на ринку послуг.

Мета дослідження

Метою даного дослідження є розробка та створення сайту-візитки фірми, що спеціалізується на організації святкових заходів для дітей.

Результати дослідження

У процесі дослідження проведено ретельний огляд та аналіз тематичної області дослідження: оглянуто сайти як візитівки підприємства: основні способи, цілі та критерії створення сайтів, архітектура сайту та основні складові процесу розробки; оглянуто сайти компаній по організації свят, виявлено їх переваги та недоліки, визначено вимоги до власного сайту.

За допомогою мов програмування JavaScript, PHP з використанням СКБД MySQL було розроблено функціональний, структурований і візуально привабливий ресурс, який забезпечує користувачеві зручну взаємодію з інтерфейсом та доступ до основної інформації про діяльність компанії. Сайт має адаптивний дизайн, що дозволяє коректно відображати контент на пристроях із різними розмірами екранів – від настільних комп'ютерів до мобільних телефонів.

Завдяки використанню сучасних вебтехнологій реалізовано низку інтерактивних елементів, серед яких:

- анімація числових показників, що підвищує динамічність сторінки;
- ефекти підсвічування при наведенні курсору, які покращують зручність навігації;
- плавні переходи між розділами за допомогою фіксованої панелі навігації;
- застосування ефекту паралаксу, що надає сторінкам глибини та привабливого візуального сприйняття.

Користувач може легко знайти необхідну інформацію завдяки логічно побудованій структурі сайту: від ознайомлення з головною сторінкою та описом послуг до перегляду фотогалереї, актуальних акцій, відгуків клієнтів і контактних даних.

Сайт виконує не лише інформаційну, а й комунікативну функцію – через інтерактивну форму зворотного зв'язку користувач має змогу звернутися до адміністрації, залишити запит чи побажання, а також підписатися на розсилку новин компанії. Таким чином, проєкт можна вважати повноцінним сучасним представленням діяльності фірми у цифровому просторі.

Практична цінність розробленого сайту полягає у створенні ефективного комунікаційного каналу, який дозволить фірмі:

- популяризувати власні послуги та підвищити впізнаваність бренду;
- збільшити обсяг замовлень і розширити клієнтську базу;
- підвищити конкурентоспроможність на ринку розважальних послуг;
- оптимізувати рекламні витрати й підвищити рентабельність бізнесу в цілому.

Висновки та перспективи

Результатом виконаної роботи є функціональний, інформативний та естетично привабливий сайт фірми, що спеціалізується на організації святкових заходів для дітей.

Підсумовуючи, можна зазначити, що розроблений сайт є вдалим прикладом сучасного адаптивного вебресурсу, який поєднує в собі функціональність, естетичність і зручність користування. Подальший розвиток проєкту може бути спрямований на вдосконалення інтерфейсу, розширення інтерактивних можливостей і впровадження додаткових сервісів для клієнтів.

Список використаних джерел

1. Исси Коэн Л. Повний довідник по HTML, CSS и JavaScript. – М.: ЭКОМ Паблишерз, 2019. – 938 с.
2. Флэнаган Д. JavaScript: Посібник. – СПб.: Символ Плюс, 2008. – 992 с.
3. Росс В.С. Створення сайтів: HTML, CSS, PHP, MySQL: підручник, ч. 1 – М.: МГДД(Ю)Т, 2020. – 107 с.
4. Welling L., Thomson L. PHP and MySQL Web Development. – Boston: Addison-Wesley Professional, 2016. – 688с.

Шикула О.М., професор ДУІКТ
Марохонько М.О., студент ДУІКТ

МЕРЕЖЕВЕ СХОВИЩЕ NAS (NETWORK ATTACHED STORAGE) З ВИКОРИСТАННЯМ REACT.JS, NODE.JS, MONGODB

Постановка задачі

У сучасному цифровому світі зберігання та управління даними є критично важливим завданням як для приватних користувачів, так і для бізнесу. Традиційні методи зберігання інформації на локальних пристроях мають значні обмеження: відсутність централізованого доступу, складність синхронізації між пристроями, ризики втрати даних та обмежені можливості масштабування. Мережеві сховища (NAS) вирішують ці проблеми, надаючи централізований доступ до файлів через мережу.

Створення власного NAS рішення з використанням сучасних веб-технологій відкриває нові можливості для гнучкого управління файлами. Використання React.js для інтерфейсу користувача забезпечує швидку та інтуїтивно зрозумілу роботу, Node.js на серверній частині гарантує високу продуктивність та масштабованість, а MongoDB як база даних дозволяє ефективно зберігати метадані файлів та управляти структурою каталогів.

Така система має забезпечувати не лише базові функції зберігання та отримання файлів, а й розширені можливості: організацію файлової структури, контроль доступу користувачів, швидкий пошук, попередній перегляд медіа-контенту, управління спільним доступом та синхронізацію даних у режимі реального часу. Інтеграція цих технологій створює потужну платформу для ефективного управління файлами в мережевому середовищі.

Мета дослідження

Метою роботи є розробка повнофункціонального мережевого сховища NAS з використанням сучасного технологічного стеку React.js, Node.js та MongoDB. Система має забезпечувати зручний веб-інтерфейс для управління файлами, підтримку багатокористувацького режиму з розмежуванням прав доступу, можливість організації файлової структури у вигляді вкладених каталогів та ефективний пошук по вмісту сховища.

Розроблене рішення спрямоване на демонстрацію можливостей створення власної альтернативи комерційним хмарним сервісам зберігання даних, таким як Dropbox або Google Drive, з повним контролем над інфраструктурою та даними. Система має бути масштабованою, безпечною та продуктивною, забезпечуючи швидкий доступ до файлів через веб-браузер з будь-якого пристрою.

Додатковою метою є впровадження механізмів спільного доступу до файлів та папок, системи версіонування змін, інтеграції з сокетом для синхронізації в режимі реального часу та створення інтуїтивного користувачького інтерфейсу, який забезпечує комфортну роботу з великими обсягами даних.

Результати дослідження

У результаті дослідження було створено повнофункціональне мережеве сховище NAS з веб-інтерфейсом, яке забезпечує комплексне управління файлами

через мережу. Система складається з клієнтської частини на React.js з використанням TypeScript для типобезпеки коду та серверної частини на Node.js з Express.js фреймворком.

Впроваджено функціонал швидкого пошуку з підтримкою дебаунсингу, який дозволяє знаходити файли та папки як серед власних даних користувача, так і серед спільних ресурсів. Система підтримує основні операції з файлами: завантаження, скачування, переміщення, копіювання, перейменування та видалення з можливістю відновлення з кошика.

Реалізовано систему логуювання з використанням Winston для відстеження всіх операцій у системі, що забезпечує можливість аудиту та діагностики проблем. Додатково впроваджено механізм фонові синхронізації файлів з файловою системою сервера через демон-процес, який відстежує зміни та автоматично оновлює базу даних.

Висновки та перспективи

Розроблене мережеве сховище NAS демонструє ефективність використання сучасних веб-технологій для створення систем управління файлами. Поєднання React.js, Node.js та MongoDB забезпечує високу продуктивність, масштабованість та зручність використання. Система має модульну архітектуру, що дозволяє легко розширювати функціонал та адаптувати рішення під різні потреби.

Додатково планується розробка системи аналітики використання сховища, інтеграція з різними протоколами доступу (WebDAV, FTP, SMB) та створення API для інтеграції з іншими системами. Впровадження технологій машинного навчання для автоматичної класифікації та тегування файлів також може значно покращити зручність роботи з великими обсягами даних.

Список використаних джерел

1. React Documentation. – [Електронний ресурс]. – URL: <https://react.dev/>
2. Node.js Documentation. – [Електронний ресурс]. – URL: <https://nodejs.org/docs/>
3. MongoDB Documentation. – [Електронний ресурс]. – URL: <https://www.mongodb.com/docs/>
4. Express.js Guide. – [Електронний ресурс]. – URL:

Шикун О.М., професор ДУІКТ
Мазуренко А.В., студент ДУІКТ

ВЕБ-ЗАСТОСУНОК ДЛЯ СТВОРЕННЯ РЕЗЮМЕ НА МОВІ ПРОГРАМУВАННЯ JAVASCRIPT З ВИКОРИСТАННЯМ ФРЕЙМВОРКУ VUE.JS

Постановка задачі

У сучасних умовах високої конкуренції на ринку праці формування якісного та структурованого резюме є критично важливим для фахівців різних галузей.

Більшість доступних сервісів для створення резюме або мають обмежений функціонал, або передбачають платний доступ до ключових можливостей. Наявні рішення часто не враховують індивідуальні потреби користувачів, гнучкість налаштування дизайну та зручність експорту документа у сучасні формати. Це зумовлює актуальність розробки веб-застосунку, який дозволив би користувачам швидко та просто створювати персоналізовані резюме, редагувати їх у реальному часі та експортувати у PDF чи інші поширені формати.

Отже, виникає потреба у створенні доступного, адаптивного та інтуїтивно зрозумілого інструменту для формування резюме.

Мета дослідження

Метою роботи є розробка інтерактивного веб-застосунку, що дозволяє користувачеві створювати, редагувати та зберігати резюме без необхідності встановлення додаткового програмного забезпечення. Застосунок має забезпечити зручний інтерфейс, підтримку живого прев'ю документа, вибір шаблонів, формування PDF-версії та можливість збереження даних на стороні користувача або в хмарному сховищі.

Результати дослідження

У ході дослідження виконано детальний огляд існуючих веб-сервісів для створення резюме, визначено їхні сильні та слабкі сторони. На основі аналізу сформульовано ключові функціональні вимоги до майбутнього застосунку: простота введення даних; можливість гнучкої зміни структури резюме; підтримка кількох шаблонів; миттєве оновлення прев'ю; збереження та експорт у формат PDF.

Для реалізації клієнтської частини обрано JavaScript та фреймворк Vue.js, який забезпечує реактивність інтерфейсу та високу продуктивність завдяки компонентному підходу. Використання Vue.js дозволяє легко створити динамічні форми, забезпечити оновлення даних у режимі реального часу та реалізувати масштабовану архітектуру.

Для стилізації застосунку використано HTML5 та CSS3, а також адаптивні підходи для коректного відображення на різних пристроях. Експорт у PDF реалізовано за допомогою бібліотек JavaScript, що дозволяють генерувати документ безпосередньо на стороні клієнта, без участі сервера.

Також створено систему локального збереження даних (LocalStorage), що дає можливість продовжувати роботу над резюме у будь-який момент. Проведене тестування підтвердило коректність роботи інтерфейсу, стабільну реактивність та відповідність вимогам юзабіліті.

Висновки та перспективи

У результаті виконаної роботи створено функціональний веб-застосунок, який дозволяє користувачеві швидко й зручно створювати персональне резюме, обирати шаблон оформлення, переглядати результат у реальному часі та експортувати його у PDF. Розроблений інструмент може бути використаний як самостійний сервіс або як частина системи кар'єрного консалтингу чи HR-платформи.

Подальше вдосконалення ресурсу може передбачати створення особистих кабінетів; хмарне збереження резюме; генерацію резюме на основі штучного

інтелекту; інтеграцію з професійними соцмережами (LinkedIn тощо); розширення бібліотеки шаблонів і тем оформлення.

Список використаних джерел

1. Флэнаган Д. JavaScript. Посібник. – СПб.: Символ–Плюс, 2018. – 1160 с.
2. You Y. The Vue.js Handbook. – FreeCodeCamp, 2020. – 182 p.
3. Duckett J. HTML & CSS: Design and Build Websites. – Wiley, 2014. – 512 p.
4. Kramnik I. Modern JavaScript Tools. – O'Reilly Media, 2021. – 324 p

Шикула О.М., професор ДУІКТ
 Марохонько М.О., студент ДУІКТ

ДОСЛІДЖЕННЯ МЕТОДІВ ШИФРУВАННЯ ПОВІДОМЛЕНЬ ТА ПРОБЛЕМИ ЗАХИСТУ ДАНИХ У СУЧАСНИХ ЧАТ СЕРВІСАХ

Постановка задачі

З розвитком цифрових технологій чат-сервіси стали невід'ємною частиною повсякденного спілкування. Однак широке використання таких платформ супроводжується значними ризиками для приватності користувачів. Виникають проблеми витоку персональних даних, несанкціонованого доступу до повідомлень та неналежного регулювання з боку провайдерів. Постає необхідність дослідження сучасних викликів у сфері захищеності приватності в чат-сервісах з метою пошуку ефективних шляхів їх вирішення.

Мета дослідження

Метою дослідження є аналіз актуальних проблем захищеності приватності в чат-сервісах, визначення ключових викликів та розробка рекомендацій щодо підвищення рівня безпеки персональних даних користувачів.

Результати дослідження

1. Аналіз основних загроз приватності:
 - Перехоплення даних третіми сторонами (man-in-the-middle атаки).
 - Незашифровані або слабко зашифровані повідомлення.
 - Збирання та монетизація персональних даних платформами.
2. Юридичні аспекти захисту приватності:
 - Дотримання міжнародних стандартів (GDPR, CCPA).
 - Політики конфіденційності провайдерів.
3. Технологічні шляхи захисту:
 - Впровадження наскрізного шифрування (end-to-end encryption).
 - Автентифікація користувачів за допомогою багатофакторних методів.
 - Децентралізовані системи зберігання даних.
4. Огляд популярних платформ:
 - Порівняння рівнів захисту (WhatsApp, Telegram, Signal, Viber).
 - Виявлення сильних та слабких сторін кожного сервісу.

Висновки та перспективи

В результаті дослідження встановлено, що найбільш ефективними методами захисту приватності в чат-сервісах є наскрізне шифрування, багатофакторна автентифікація та прозорі політики конфіденційності. Проте значна кількість платформ все ще не відповідає високим стандартам безпеки, що вимагає подальшого вдосконалення технологій та нормативно-правового регулювання. Перспективними напрямками є розвиток децентралізованих систем обміну даними та підвищення цифрової грамотності користувачів у питаннях захисту особистої інформації.

Список використаних джерел

1. Stallings, W. (2017). *Cryptography and Network Security: Principles and Practice*. 7th Edition. Pearson.
2. Diffie, W., & Hellman, M. (1976). New Directions in Cryptography. *IEEE Transactions on Information Theory*, 22(6), 644–654.
3. Moxie Marlinspike (2014). Signal Protocol: Cryptographic Protocol for Encrypted Communication. <https://signal.org/docs/>
4. Perrin, J. (2013). Telegram: Secrets of the MTProto Protocol. <https://core.telegram.org/mtproto>
5. Kaufman, C., Perlman, R., & Speciner, M. (2002). *Network Security: Private Communication in a Public World*. Prentice Hall.

Гніденко М.П., доцент ДУІКТ
Сітко Д.О., аспіран ДУІКТ

ІНТЕЛЕКТУАЛЬНІ МЕТОДИ ПОБУДОВИ ВЕКТОРНИХ КОНТУРІВ НА ОСНОВІ ТОПОЛОГІЧНОЇ ІНТЕРПРЕТАЦІЇ ПІКСЕЛЬНИХ СТРУКТУР

Сучасні системи обробки графічної інформації дедалі частіше потребують високоточної векторизації растрових зображень для подальшого аналізу, редагування чи інтеграції у складні графічні середовища. Проте багато традиційних алгоритмів мають обмеження, пов'язані з їх чутливістю до шумів, неоднорідності піксельної структури та локальних спотворень. У цьому контексті особливо перспективним є застосування топологічної інтерпретації даних, що дозволяє аналізувати зображення не лише як набір пікселів, а як складну структуру взаємопов'язаних елементів.

Постановка проблеми. Під час автоматичної векторизації часто виникають труднощі зі збереженням цілісності контурів, особливо у випадках, коли зображення має нерівномірну текстуру, фрагментацію або значний рівень шуму. Стандартні методи, які ґрунтуються на градієнтних детекторах або локальній сегментації, схильні до втрати зв'язності та помилкового об'єднання областей. Топологічні методи дають змогу уникнути цих проблем, але питання створення

інтелектуальних алгоритмів, що поєднують топологічний аналіз та адаптивну генерацію векторних контурів, залишається відкритим.

Мета дослідження. Метою цієї роботи є створення інтелектуального підходу до побудови векторних контурів на основі глибокого аналізу топологічної структури піксельних даних. Основний акцент робиться на інтеграції алгоритмів топологічної сегментації та структурного аналізу з адаптивними методами реконструкції кривих, реалізованими у середовищі C#.

Результати дослідження. У ході роботи сформовано інтелектуальний метод, що включає такі етапи:

1) Топологічна інтерпретація піксельного поля.

Виконано аналіз структури зображення через компоненти зв'язності, критичні точки та локальні топологічні інваріанти. Застосовано підхід топологічних фільтрацій, що дозволяє виділяти стійкі інваріантні елементи незалежно від шуму.

2) Ідентифікація структурних одиниць зображення.

Запропоновано механізм автоматичного визначення лінійних, контурних та циклічних фрагментів, що формують основу майбутнього векторного представлення.

3) Побудова топологічного графа.

Піксельна структура інтерпретується як граф, вершини якого відповідають ключовим топологічним особливостям, а ребра - шляхам між ними. Така модель забезпечує відновлення зв'язності навіть у зашумлених та фрагментованих регіонах.

4) Генерація адаптивних векторних кривих. Використано комбінований підхід:

- апроксимація сплайнів з урахуванням топологічних обмежень;
- автоматичний вибір рівня згладжування залежно від локальної структури;
- уникнення перетинів і помилкових замикань контурів.

Завдяки цьому вдається зберегти як плавність, так і структурну точність.

5) Програмна реалізація мовою C#.

Створено модульну бібліотеку у середовищі .NET, що включає:

- модуль обробки піксельних структур;
- модуль побудови топологічного графа;
- алгоритми трасування та оптимізації контурів;
- експорт у формат SVG;
- оптимізацію лінійних обчислень та ефективне оперування піксельними структурами.

Експериментальні дослідження демонструють, що запропонований метод істотно підвищує якість відтворення контурів у складних умовах, зменшує кількість розривів та забезпечує точну передачу циклічних та дрібних елементів.

Висновки та перспективи. Розроблений інтелектуальний підхід до побудови векторних контурів на основі топологічної інтерпретації піксельних структур має великий потенціал у задачах точної векторизації. Використання топологічної моделі дозволяє забезпечити високу стійкість до шумів та локальних

спотворень, а адаптивні методи генерації кривих забезпечують збереження деталей та правильну структуру контурів.

Подальші перспективи роботи включають:

- інтеграцію елементів машинного навчання для покращеної класифікації структур;
- розширення методу для кольорових та багатопланових зображень;
- оптимізацію алгоритму для обробки у реальному часі;
- розроблення комплексного пакету інструментів для промислових систем векторизації.

Список використаної література

1. Lin X., Peng S., Xie X., Zhu J., Zhou Y., Gao L. Clair Obscur: An Illumination-Aware Method for Real-World Image Vectorization. 2025. arXiv. <https://arxiv.org/abs/2511.20034>
2. He R. Y., Kang S. H., Morel J.-M. A Formalization of Image Vectorization by Region Merging. 2024. arXiv. <https://arxiv.org/abs/2409.15940>
3. Image vectorization using a sparse patch layout. 2024. Graphical Models. ScienceDirect. <https://www.sciencedirect.com/science/article/pii/S1524070324000171>
4. Yoshizawa S., Michikawa T., Yokota H. Topological Delaunay Graph for Efficient 3D Binary Image Analysis. 2024. International Journal of Automation Technology. <https://www.fujipress.jp/ijat/au/ijate001800050632/>
5. Eremeev S. V., Abakumov A. V., Andrianov D. E., Shirabakina T. A. Vectorization Method of Satellite Images Based on Their Decomposition by Topological Features. 2023. Informatics and Automation. <https://journals.rcsi.science/2713-3192/article/view/265798>
6. Image Vectorization: A Review. 2023. EmergentMind. <https://www.emergentmind.com/papers/2306.06441>
7. Sitko D. O., Hnidenko M. P. Топологічний підхід до векторизації растрової графіки на основі симпліціальних комплексів. 2025. Наукові записки ДУІКТ. <https://journals.dut.edu.ua/index.php/sciencenotes/article/view/3270>
8. Shaxislam Joldasov, Saida Beknazarova. Methods and Algorithms for Vectorization of Raster Images. 2022. IEJRD. <https://www.iejrd.com/index.php/article/view/2909>
9. Bazelyuk V. M., Svyatiy I. R. Алгоритм векторизації растрових зображень для створення інфографіки освітнього процесу. 2021. Тези доповідей Міжнародної науково-технічної конференції. <https://repository.kpi.kharkov.ua/items/22d8976b-653c-42be-9636-1742741c8a20>
10. Karas I. R., Bayram B., Batuk F., Akay A. E., Baz I. Multidirectional Scanning Model (MUSCLE) to Vectorize Raster Images with Straight Lines. 2008. Sensors (Basel). <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC3673439/>

Кисіль Т. М., старший викладач ДУІКТ
Кривий А.В., студент ДУІКТ

AIOPS: МЕТОДОЛОГІЯ ЗАСТОСУВАННЯ МАШИННОГО НАВЧАННЯ ДЛЯ КОРЕЛЯЦІЇ ПОДІЙ ТА ОПТИМІЗАЦІЇ МОНІТОРИНГУ

У сучасних високодинамічних інформаційно-технічних системах (ІТ-системах) обсяг телеметричних даних зростає експоненційно, що критично ускладнює оперативне реагування на інциденти та підтримання стійкої стабільності сервісів. Традиційні методи моніторингу, які ґрунтуються на статичних та фіксованих порогових правилах і вимагають значного ручного налаштування, виявляються малоефективними в умовах непередбачуваних навантажень та складних розподілених архітектур. З огляду на це, дослідження можливостей застосування штучного інтелекту (ШІ) у межах *DevOps-процесів* набуває високої актуальності та значної практичної значущості для підвищення операційної ефективності.

Постановка задачі. У контексті експоненційного зростання обсягів метрик виникає необхідність заміни реактивних ручних методів моніторингу на проактивні автоматизовані підходи. Актуальна задача полягає у дослідженні та валідації концепції *AIOps (Artificial Intelligence for IT Operations)* як інтелектуального інструменту. Необхідно визначити найбільш ефективні методи ШІ для автоматичного аналізу часових рядів метрик, виявлення прихованих аномалій, кореляції подій та автоматизації визначення першопричин збоїв з метою значного зниження часу відновлення після інцидентів (*MTTR*) та підвищення загальної стійкості сервісів.

Мета дослідження. Метою роботи є аналіз підходів *AIOps* та визначення оптимальних способів використання методів штучного інтелекту для автоматизованого аналізу метрик, прогнозування інцидентів та оптимізації моніторингу в умовах сучасних ІТ-інфраструктур. Ключовий акцент робиться на порівнянні адаптивних можливостей *AI* з обмеженнями традиційних порогових систем.

Результати дослідження. *AIOps* забезпечує фундаментальний перехід від реактивного до проактивного управління інфраструктурою. Саме це досягається шляхом аналізу багатовимірних часових рядів метрик і виявлення прихованих аномалій, що не фіксуються статичними правилами. На відміну від традиційного моніторингу, ШІ адаптується до динамічних змін навантажень і враховує контекст нормальної поведінки сервісів. Дана адаптивність дозволяє прогнозувати інциденти до їхнього фактичного настання та суттєво знижувати кількість хибних сповіщень (*alert fatigue*).

Застосування алгоритмів машинного навчання, зокрема методів виявлення аномалій та прогнозування часових рядів, забезпечує високу точність аналізу великих масивів телеметрії. *AIOps* автоматизує визначення першопричин збоїв завдяки автоматичній кореляції подій та метрик між різними компонентами системи, що значно скорочує час реагування. Практика провідних компаній

підтверджує, що інтеграція ШІ у моніторинг знижує MTTR (Mean Time To Recovery) та підвищує якість обслуговування.

Наприклад, модель аномалій, навчена на історичних даних *CPU*, *latency* та кількості запитів, здатна ідентифікувати нетипову активність ще до появи критичних сповіщень. На основі прогнозу зростання навантаження система може автоматично активувати *autoscaling* або перерозподілити ресурси. У випадку несподіваних затримок на рівні бази даних, *AIOps* визначає взаємозв'язок між деградацією API та збільшенням навантаження на DB-шар, генеруючи рекомендацію щодо усунення проблеми. Така взаємодія підтверджує перевагу *AIOps* як інтелектуального інструменту запобігання інцидентам, а не лише їх фіксації.

Концептуальна основа дослідження полягає у розробці та емпіричній апробації *гібридної AIOps-моделі*, здатної здійснювати автоматичну корекцію прогностичних порогових значень та патернів адекватної поведінки метрик. Така корекція має відбуватися не лише на підставі аналізу історичних часових рядів, а й із залученням контекстних метаданих про поточні інфраструктурні зміни (*Infrastructure Changes*).

Традиційні алгоритми виявлення аномалій, такі як *Isolation Forest*, *LSTM-кодеру*, навчаються виключно на історичних даних. Як наслідок, вони часто генерують хибно-позитивні результати (*false positives*) під час запланованих або значущих змін у системі. Типовими прикладами є: раптове зростання використання CPU/пам'яті під час розгортання нового коду (*Deployment*), що є нормою, але класифікується як аномалія; різке падіння навантаження на один сервер через масштабування (*Autoscaling*); або тимчасові затримки, пов'язані з плановим обслуговуванням (*Maintenance*) бази даних. Для вирішення цієї проблеми пропонується впровадження спеціалізованого механізму адаптації.

Механізм функціонує на основі контекстного входу (*Contextual Input*). Модель отримує не лише традиційну телеметрію (показники *CPU*, *Latency*), але й метадані від зовнішніх систем: CI/CD-системи, систем оркестрації (*Kubernetes*, *Terraform*) та систем планування (*Jira/ServiceNow*). Метадані включають: тип події (*Deployment*, *Scaling Up*, *Database Upgrade*), час початку/завершення зміни та компоненти, яких стосується зміна (*Affected Services*).

На основі цього контексту здійснюється динамічне зважування (*Dynamic Weighting*). У момент інфраструктурної зміни гібридна модель автоматично знижує свою чутливість до прогнозованих типів аномалій, що можуть бути викликані зміною. Наприклад, вона ігнорує піки CPU, якщо відомо, що саме в цей час відбувається рестарт контейнера.

Після завершення планової зміни активується процес автоматичного перенавчання (*Online Retraining*). Модель оперативно використовує дані, зібрані під час цієї зміни, для оновлення патернів та інтеграції нового рівня базової лінії продуктивності.

Реалізація даної концепції дозволить суттєво знизити кількість хибних результатів *AIOps*-систем, які наразі є основною причиною недовіри операційних команд до ШІ-інструментів. Саме це забезпечить більш точне виявлення несанкціонованих аномалій (несанкціоноване використання ресурсів або

приховані помилки, що виникають під час розгортання), підвищуючи надійність та ефективність автоматизованого моніторингу.

Розглянута гібридна AIOps-модель, яка використовує контекст інфраструктурних змін, забезпечує перехід до якісно нового рівня проактивного моніторингу. Динамічна адаптація чутливості моделі дозволить суттєво мінімізувати кількість хибно-позитивних результатів, підвищуючи довіру операційних команд до ШІ-інструментів, що забезпечує точне виявлення інцидентів та створює основу для більш автономних і стійких ІТ-систем.

Висновки та Перспективи. Проведене дослідження підтвердило, що застосування ШІ в аналізі метрик суттєво підвищує точність виявлення аномалій та мінімізує кількість хибних сповіщень. AIOps прискорює реагування на інциденти завдяки автоматичній кореляції подій. Практична цінність підходу полягає у можливості інтеграції прогностичних моделей безпосередньо у DevOps-процеси з метою створення більш автономних та ефективних систем моніторингу.

Впровадження AIOps є невід'ємним кроком у розвитку сучасних ІТ-операцій, оскільки воно забезпечує новий рівень автоматизації та стійкості систем. Використання штучного інтелекту в моніторингу дозволяє зменшити операційне навантаження на DevOps-команди, підвищити якість кінцевих сервісів та оптимізувати використання ресурсів інфраструктури. Отримані результати підтверджують як значний практичний, так і теоретичний потенціал AIOps.

Список використаних джерел

1. Gartner. Market Guide for AIOps Platforms. Gartner Research, 2023, с. 12–34.
2. IBM. AIOps: The Future of IT Operations. IBM Whitepaper, 2022, с. 5–28.
3. Google SRE Team. Site Reliability Engineering: How Google Runs Production Systems. O'Reilly Media, 2020, с. 50–90.
4. Villamizar M., Garcés O. Machine Learning for Anomaly Detection in Cloud and Distributed Systems. ACM Publications, 2021, с. 15–40.
5. Microsoft. Intelligent Monitoring with Artificial Intelligence. Microsoft Research, 2022, с. 22–48.

Кисіль Т.М., старший викладач ДУІКТ

Марченко Б.С., студент ДУІКТ

Сущенко В.В., студент ДУІКТ

АРХІТЕКТУРНИЙ ПОРІВНЯЛЬНИЙ АНАЛІЗ АСИНХРОННИХ МОДЕЛЕЙ ОБРОБКИ ПОДІЙ ПРИ ПРОЕКТУВАННІ ЧАТ-БОТІВ

У сучасному цифровому просторі чат-боти є незамінним інструментом для автоматизації обслуговування користувачів, ефективної модерації онлайн-спільнот та реалізації широкого спектру інтерактивних сервісів. Створення цих систем ґрунтується на використанні спеціалізованих програмних бібліотек, які

забезпечують надійну та зручну взаємодію з офіційними API відповідних платформ. Серед найпопулярніших платформ для розробки та розміщення ботів чітко виділяються Discord та Telegram.

Постановка задачі. Оскільки Discord і Telegram використовують принципово різні архітектурні підходи до обміну даними та структуру API, виникає необхідність у проведенні комплексного порівняльного аналізу бібліотек, що застосовуються для створення ботів у цих середовищах. Дослідження фокусується на вивченні бібліотек *discord.py* та *aiogram*, які є провідними у своїх екосистемах та реалізують асинхронну взаємодію. Такий аналіз дозволяє визначити ключові архітектурні особливості, оцінити ефективність реалізації асинхронних процесів та дослідити можливості інтеграції проміжних шарів обробки подій, відомих як *middleware*.

Метою дослідження є проведення порівняльного аналізу архітектурних рішень бібліотек *discord.py* та *aiogram*. Особливу увагу приділено аналізу принципів асинхронної обробки подій, відмінностям у підходах до взаємодії з API та механізмам реалізації *middleware* як інструменту керування потоком даних між обробниками подій. Результати аналізу мають визначити ключові переваги та функціональні обмеження кожної бібліотеки у контексті ефективності, масштабованості та простоти розроблення складних бот-систем.

Результати дослідження. Для проектування ботів у середовищі Discord застосовуються такі бібліотеки, як *discord.py* та *nextcord*, які реалізують подійно-орієнтовану модель взаємодії з використанням постійного *WebSocket*-з'єднання. Дана архітектура забезпечує миттєве реагування бота на події в реальному часі, такі як надходження повідомлень, зміни статусів користувачів або керування голосовими каналами. На противагу цьому, в екосистемі Telegram найпоширенішими є бібліотеки *aiogram* та *Telebot*, які взаємодіють із сервером через HTTP-запити з використанням технологій *long polling* або *webhooks*. Такий підхід передбачає переважно послідовну обробку подій, що є оптимальним для створення інформаційних і сервісних ботів із чітко визначеною логікою діалогу. Основна архітектурна відмінність полягає у типі взаємодії з сервером: Discord орієнтується на обмін даними в реальному часі, тоді як Telegram використовує періодичні запити до API.

Для поглибленого дослідження асинхронної взаємодії та механізмів *middleware* був проведений детальний порівняльний аналіз *discord.py* та *aiogram*. Обидві бібліотеки функціонують на основі асинхронної моделі виконання (*asyncio*), що дозволяє їм ефективно обробляти велику кількість одночасних подій без блокування основного потоку. Однак їхня внутрішня архітектура має суттєві розбіжності. *Discord.py* обробляє події через цикли обробки подій у реальному часі, проте не має власного вбудованого прошарку *middleware*, дозволяючи користувачеві реалізовувати перехоплювачі подій через декоратори. Натомість *aiogram* пропонує структуровану систему *middleware*, що забезпечує гнучке перехоплення, зміну або фільтрацію вхідних та вихідних повідомлень. Це робить *discord.py* високоефективним при роботі з великим потоком подій у реальному часі (реакції, статуси), а *aiogram* — оптимізованим для послідовної обробки діалогів, команд і повідомлень, необхідних для бізнес-логіки. Продуктивність

discord.py у багатокористувацьких сценаріях є високою завдяки постійному з'єднанню та мінімальним затримкам, тоді як продуктивність *aiogram* може залежати від частоти опитування (*polling rate*) або стабільності *webhook*-з'єднання.

Висновки та перспективи. Проведений аналіз показує, що вибір бібліотеки залежить від архітектурних особливостей цільової платформи та поставленої задачі. Бібліотека *discord.py* демонструє вищу ефективність у розробленні інтерактивних ботів для *Discord*, оскільки платформа сама є подійно-орієнтованою і працює через постійне WebSocket-з'єднання. Це дозволяє боту отримувати події миттєво і швидко реагувати на динамічну активність користувачів. З іншого боку, *aiogram* забезпечує більш гнучку й структуровану архітектуру для Telegram-ботів, де взаємодія побудована на long polling або *webhook*. Саме це робить *aiogram* особливо зручною для створення сценарно-орієнтованих ботів із чіткою бізнес-логікою та послідовними етапами діалогу. Таким чином, *discord.py* краще підходить для подійно-орієнтованих середовищ із високим потоком даних, а *aiogram* — для структурованих Telegram-ботів зі складними сценарними взаємодіями.

Список використаних джерел:

1. Discord Inc. Discord Developer Portal: Bot Accounts [Електронний ресурс] // Documentation. Режим доступу: <https://discord.com/developers/docs/topics/oauth2#bot-accounts> (дата звернення: 8.11.2025).
2. Dewangan D., Soni N. AI-Powered Discord Companion Bot: Enhancing Job Search and Information Access through Intelligent Automation [Електронний ресурс] // International Journal of Research Publication and Reviews. – Vol. 6, Issue 4. – P. 16942–16949. – April 2025. Режим доступу: <https://ijrpr.com/uploads/V6ISSUE4/IJRPR44056.pdf>
3. Nachankar A. P., Tamgade A., Bhawalkar H., Tidke G., Dhale G., Shinde S., Jawale N., Chouhan S. DISCORD BOT AI [Електронний ресурс] // International Research Journal of Modernization in Engineering Technology and Science. – Issue 12, December 2024. – P. 1312–1316. Режим доступу: https://www.irjmets.com/uploadedfiles/paper/issue_12_december_2024/65116/final/final_irjmets1734021319.pdf
4. Srivastava V., Singh P., Balodhi S., Mishra S., Gahlot P. Discord Bot Using AI [Електронний ресурс] // International Journal of Research Publication and Reviews. – Vol. 5, No. 5. – P. 10195–10198. – May 2024. Режим доступу: <https://ijrpr.com/uploads/V5ISSUE5/IJRPR28558.pdf>

Кисіль Т.М., ст. викладач ДУІКТ
Товстенко М.В., студент ДУІКТ

MODEL CONTEXT PROTOCOL (MCP): АРХІТЕКТУРА ТА ІНТЕГРАЦІЯ LLM З КОРПОРАТИВНИМИ БАЗАМИ ДАНИХ

Стрімкий розвиток технологій штучного інтелекту (ШІ) генерує нагальну потребу у створенні надійних та ефективних механізмів для взаємодії інтелектуальних систем із зовнішніми сервісами, ресурсами та різноманітними джерелами даних. Одним із найбільш перспективних рішень у цій галузі є технологія Model Context Protocol (MCP), розроблена компанією Anthropic. MCP являє собою стандартизований протокол, призначений для безпечного підключення великих мовних моделей (LLM) до зовнішніх інструментів та ресурсів, включно з базами даних, що відкриває нові можливості для автоматизації та інтелектуального аналізу даних у широкому спектрі галузей.

Постановка задачі. Традиційні підходи до роботи з корпоративними базами даних часто вимагають глибоких технічних знань, необхідності ручного написання складних SQL-запитів та детального розуміння внутрішньої структури бази даних. Поява Model Context Protocol створює можливість кардинального спрощення цього процесу, дозволяючи системам штучного інтелекту безпосередньо взаємодіяти з базами даних через стандартизований інтерфейс. Це призводить до зменшення фрагментації системи та суттєвого спрощення доступу до даних. У цьому контексті актуальним є дослідження архітектурних особливостей та практичних способів використання MCP для роботи з базами даних, аналіз його переваг порівняно з класичними методами, оцінка потенціалу для оптимізації робочих процесів, а також виявлення та мінімізація потенційних ризиків безпеки та підтримки.

Метою дослідження є комплексний аналіз можливостей застосування Model Context Protocol для доступу та роботи з базами даних у різних секторах економіки. Дослідження включає визначення ключових практичних сценаріїв використання, порівняння ефективності з традиційними методами доступу до даних та окреслення перспектив розвитку цього напрямку для глибокої та безпечної інтеграції ШІ з корпоративними інформаційними системами.

Результати дослідження. *Model Context Protocol* реалізований як відкритий стандарт, що дозволяє великим мовним моделям безпечно встановлювати з'єднання з різноманітними джерелами даних, охоплюючи як реляційні, так і нереляційні бази даних. Архітектура MCP побудована на клієнт-серверній моделі, де клієнт (зазвичай ШІ-асистент) взаємодіє з MCP-сервером за допомогою стандартизованих повідомлень JSON-RPC 2.0. Дана дворівнева структура забезпечує критично важливий додатковий рівень безпеки та контролю доступу на відміну від прямого доступу LLM до бази даних.

Протокол підтримує три основні типи компонентів: *Resources* (ресурси даних, такі як таблиці бази даних), *Tools* (інструменти для виконання специфічних операцій) та *Prompts* (шаблони для типових, повторюваних запитів). Кожен MCP-сервер може експортувати визначений набір своїх можливостей, які автоматично

стають доступними для ШІ-моделі після відповідного налаштування сервера, при цьому не вимагаючи від LLM прямого програмування.

Особлива перевага MCP полягає в реалізації механізму контролю доступу на основі можливостей (*capability-based access control*). Сервер має можливість обмежити доступ моделі до певних таблиць, колонок або операцій, використовуючи гнучкі налаштування через конфігураційні файли. Завдяки цьому забезпечується принцип найменших привілеїв, при якому модель отримує доступ виключно до тих даних, які необхідні для виконання конкретної задачі. Можливість налаштування детального аудиту всіх операцій також легко імплементується на MCP-сервері, що є критично важливим для дотримання регуляторних та корпоративних стандартів.

Для комунікації MCP використовує двосторонній транспорт через механізми stdio (стандартні потоки введення/виведення) або HTTP з підтримкою Server-Sent Events. Дана двонаправлена комунікація дозволяє серверу ініціювати оновлення даних або оперативно повідомляти клієнта про зміни в базі даних у режимі реального часу.

На поточний момент існує широкий спектр MCP-серверів для роботи з різними типами баз даних, включаючи PostgreSQL, SQLite, MongoDB, Redis, Supabase та BigQuery. Універсальність та легкість розширення протоколу підтверджується наявністю офіційних SDK для розробки власних MCP-серверів на мовах Java, Python, TypeScript та інших.

У медичній галузі MCP відкриває шлях для інтелектуального аналізу масивів медичних записів пацієнтів. Лікарі можуть формувати запити природною мовою для аналізу ефективності лікування, виявлення закономірностей у медичних даних, прогнозування тенденцій захворювань та навантаження на медичні заклади, при цьому MCP забезпечує дотримання вимог конфіденційності та захисту персональних даних. В освітній галузі протокол може бути використаний для аналізу успішності студентів, формування персоналізованих навчальних планів та ефективного управління навчальними ресурсами. Для компаній у сфері E-commerce MCP надає можливість аналізу поведінки покупців, управління товарними запасами та тонкої персоналізації рекомендацій без прямого залучення технічних спеціалістів.

Висновки та перспективи. Проведений аналіз демонструє, що Model Context Protocol має значний потенціал для трансформації способів взаємодії систем штучного інтелекту з даними. MCP не лише стандартизує цю взаємодію, дозволяючи використовувати один сервер з різними ШІ-асистентами, але й значно спрощує доступ нетехнічних користувачів до складних баз даних для аналітики та прийняття обґрунтованих рішень.

Варіантами подальшої модернізації MCP є розробка підтримки розподілених транзакцій для одночасної роботи з множинними базами даних, інтеграція механізмів машинного навчання для оптимізації SQL-запитів, створення ефективних систем кешування для підвищення продуктивності, а також розробка спеціалізованих серверів для окремих галузей із попередньо налаштованими схемами та правилами доступу. Сценарії застосування MCP продовжують розширюватися, охоплюючи аналіз експериментальних даних у

науці, роботу з відкритими даними у державному управлінні, оптимізацію ланцюгів постачання у логістиці та моніторинг мережевої інфраструктури у телекомунікаціях. Очікується подальший розвиток екосистеми MCP-серверів для специфічних бізнес-процесів різних галузей економіки.

Список використаних джерел

1. Partha P. R., A Survey on Model Context Protocol: Architecture, State-of-the-art, Challenges and Future Directions. (techrxiv.org) , Дата публікації: 18.04.2025 р.
2. Xinyi H., Yanjie Z., Shenao W., Haoyu W., Model Context Protocol (MCP): Landscape, Security Threats, and Future Research Directions. arXiv:2503.23278 (arxiv.org), Дата публікації: 30.03.2025 р.
3. Mohammed M. H., Hao L., Emad F., Gopi K. R., Bram A., Ahmed E. H., Model Context Protocol (MCP) at First Glance: Studying the Security and Maintainability of MCP Servers. arXiv:2506.13538 (arxiv.org), Дата публікації: 16.06.2025 р.
4. Ziyang L., Zhiqi S., Wenzhuo Y., Zirui Z., Prathyusha J., Amrita S., Doyen S., Silvio S., Caiming X., Junnan L., MCP-Universe: Benchmarking Large Language Models with Real-World Model Context Protocol Servers. arXiv:2508.14704 (arxiv.org), Дата публікації: 20.08.2025 р.

Худік Б.О., доцент ДУІКТ
Гайшук А.О., студент ДУІКТ

РОЗРОБКА МЕТОДИКИ ВИЯВЛЕННЯ ФІШИНГОВИХ ЕЛЕКТРОННИХ ЛИСТІВ ЗА ДОПОМОГОЮ АЛГОРИТМІВ МАШИННОГО НАВЧАННЯ

Постановка задачі. Електронна пошта залишається одним із основних каналів комунікації в інформаційних системах, що робить її привабливою ціллю для фішингових атак. Фішингові електронні листи використовуються для викрадення конфіденційних даних, компрометації облікових записів та поширення шкідливого програмного забезпечення. Сучасні фішингові повідомлення часто маскуються під легітимні листи та використовують складні текстові конструкції, що значно ускладнює їх виявлення традиційними методами фільтрації. Більшість існуючих систем виявлення фішингу ґрунтуються на сигнатурних або rule-based підходах, які не забезпечують достатньої адаптивності до нових та модифікованих атак. Такі системи мають обмежену здатність до узагальнення та часто демонструють підвищену кількість хибнопозитивних і хибнонегативних спрацьовувань. Це негативно впливає як на рівень безпеки, так і на зручність користування електронною поштою. У зв'язку з цим актуальним є наукове завдання розробки ефективного методу виявлення фішингових електронних листів на основі алгоритмів машинного навчання, здатного

адаптуватися до змінних шаблонів атак та забезпечувати стабільну якість класифікації. Додатковою проблемою є інтеграція таких методів у реальні системи електронної пошти, що потребує стандартизованого механізму взаємодії між моделями аналізу та системними компонентами.

Мета дослідження. Підвищення ефективності виявлення фішингових електронних листів шляхом розробки та дослідження методу, що поєднує алгоритми машинного навчання з використанням уніфікованого інтерфейсу АМСІ. Досліджується вплив застосування різних моделей машинного навчання на якість класифікації електронних листів за стандартними метриками оцінювання, зокрема Accuracy, Precision, Recall та F1-score.

Основна ідея запропонованого підходу полягає у використанні алгоритмів машинного навчання для аналізу текстового вмісту електронних листів з подальшою інтеграцією результатів класифікації через інтерфейс АМСІ. Такий підхід дозволяє стандартизувати процес аналізу електронних повідомлень, зменшити залежність від статичних правил та забезпечити можливість масштабування і впровадження методу в реальні системи електронної пошти.

Результати дослідження. У ході дослідження було встановлено, що при використанні традиційних методів фільтрації електронної пошти ефективність виявлення фішингових листів знижується в умовах змінних шаблонів атак та складних текстових структур. Це зумовлено обмеженими можливостями таких методів щодо аналізу контексту та прихованих закономірностей у даних.

Для вирішення зазначеної проблеми було застосовано низку алгоритмів машинного навчання, зокрема Multinomial Naive Bayes, Logistic Regression, Random Forest, Decision Trees, k-nearest neighbors, Support Vector Classifier та Multilayer Perceptron. Оцінювання ефективності моделей проводилося з використанням стандартних метрик класифікації, що дозволило комплексно проаналізувати якість прогнозування та порівняти результати між різними підходами. Результати експериментального дослідження показали, що використання підходу ML у поєднанні з уніфікованим інтерфейсом АМСІ забезпечує більш збалансовані показники точності та повноти виявлення фішингових електронних листів.

Зокрема, спостерігалось зменшення кількості хибнопозитивних і хибнонегативних спрацьовувань у порівнянні з окремими класичними моделями машинного навчання, що підтверджується значеннями метрик Precision, Recall та F1-score. Інтеграція моделей машинного навчання з АМСІ дозволила стандартизувати процес передачі результатів аналізу та спростити впровадження розробленого методу у системи електронної пошти без необхідності зміни їх базової архітектури. Крім того, результати дослідження підтверджують, що запропонований підхід є перспективним рішенням для підвищення рівня безпеки електронної пошти та може бути використаний при розробці сучасних систем захисту від фішингових атак.

Список використаних джерел

1. Alshamrani A., Myneni S., Chowdhary A., Huang D. A survey on phishing attacks and countermeasures for email security // *IEEE Access*. – 2024. – Vol. 12. – P. 112345–112370.
2. Nguyen T., Kang J., Kim H. Phishing email detection using transformer-based language models // *Expert Systems with Applications*. – 2024. – Vol. 236. – Article 121435.
3. Zhang Y., Liu Q., Wang S. Context-aware phishing email detection based on BERT and attention mechanisms // *Knowledge-Based Systems*. – 2024. – Vol. 284. – Article 111320.
4. Kumar S., Jain A., Bansal A. Machine learning approaches for phishing email detection: A comparative study // *Applied Soft Computing*. – 2024. – Vol. 151. – Article 111095.
5. Aljabri M., Alqahtani A. Deep learning-based phishing email detection in enterprise environments // *Computers & Security*. – 2024. – Vol. 136. – Article 103438.
6. Chen L., Zhao Y., Sun J. Hybrid phishing detection using NLP and ensemble learning techniques // *Information Sciences*. – 2024. – Vol. 646. – P. 119–134.

Кисіль Т.М., старший викладач ДУІКТ

CNN-LSTM ПІДХІД ДО РОЗПІЗНАВАННЯ ЕМОЦІЙНИХ СТАНІВ ЗА ПОЗОЮ ОСІБ В РЕЖИМІ РЕАЛЬНОГО ЧАСУ

Традиційні методи розпізнавання емоцій (РЕМ) покладаються на аналіз міміки та вимагають великих обсягів попередньо розмічених даних. Проте поза та кінетика тіла є критичними, але часто ігнорованими, індикаторами прихованих емоційних станів. Крім того, системи, що базуються на сирих пікселях, часто нестійкі до зміни освітлення та масштабу. Проблема полягає у необхідності розробки ефективного, стійкого до шуму та масштабу методу РЕМ за позою, здатного працювати у відеопотоці в реальному часі та мінімізувати залежність від дорогих маркованих наборів даних шляхом використання некерованого навчання.

Метою дослідження являється розробка та експериментальна апробація гібридної архітектури CNN-LSTM, яка використовує принцип *самоконтрольованого контрастного навчання (SSCL)* для некерованого вивчення дискримінантних ознак пози та забезпечення надійної класифікацію базових емоційних станів (радість, сум, злість, нейтральність, тощо) за динамікою тіла у відеоконтенті в режимі реального часу.

Результати дослідження. Для забезпечення надійного розпізнавання емоцій за позою в умовах обмеженості розмічених даних критично важливим є розроблення ефективних механізмів вивчення високоякісних, інваріантних ознак. Представлений гібридна система (рис. 1) вирішує цю проблему шляхом інтеграції некерованого попереднього навчання для вилучення просторових ознак та

рекурентних мереж для моделювання часової кінетики рухів, що забезпечує високу точність класифікації емоційних станів.

Отримані результати підтверджують високу ефективність запропонованого двофазного підходу, який розділений на офлайн-навчання ознак та онлайн-класифікацію.

На початковому етапі *CNN-кодувальник* (ResNet Backbone) навчається некеровано на значному обсязі немаркованих даних надходженої пози. Даний процес вивчення інваріантних ознак оптимізується за допомогою контрастної ентропійної втрати з нормалізованою температурою (*NT-Xent Loss Calculation*). Вхідні дані проходять через *Data Augmentation* для генерації двох різних, але пов'язаних версій (*Positive Pair*), які незалежно подаються на кодувальник. Його вихід проходить через *Projection* (*MLP Backbone*). Механізм *NT-Xent* мінімізує відстань у ознаковому просторі між представленнями позитивних пар, одночасно максимізуючи відстань між представленнями негативних пар. Це дозволяє *CNN-кодувальнику* ефективно вилучати дискримінантні просторові ознаки пози, критично важливі для подальшої класифікації емоцій.

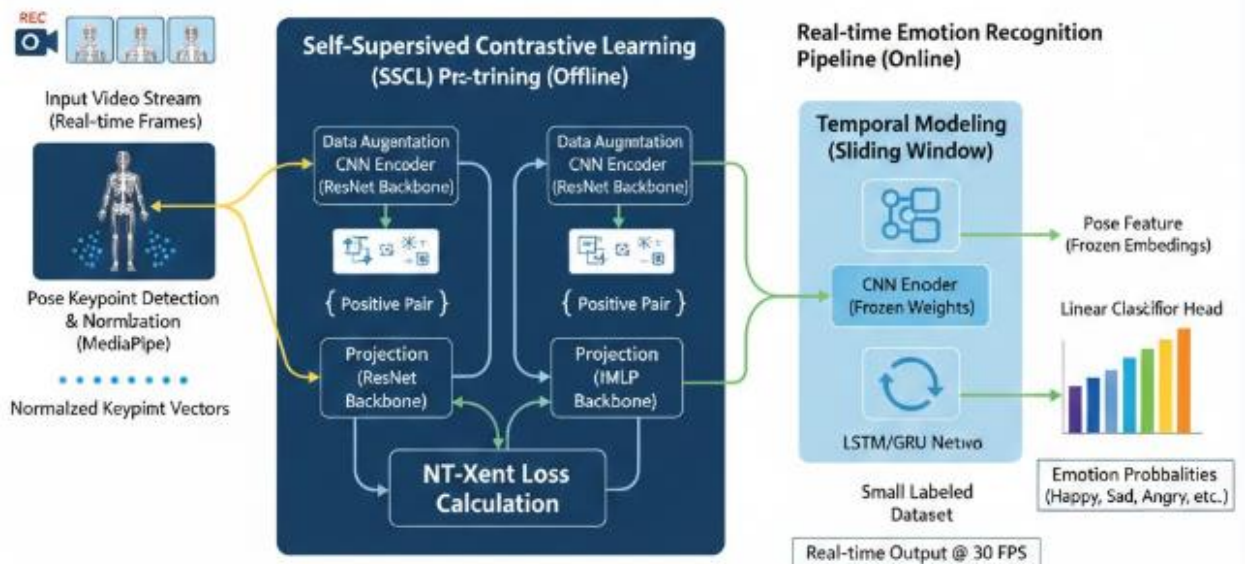


Рис. 1. Гібридна архітектура SSCL-LSTM для розпізнавання емоцій за позою в режимі реального часу

Наступний етап починається з підготовки вхідних даних, де вхідне відео (*Input Video Stream*) подається на блок виявлення та нормалізації ключових точок пози (*MediaPipe*) для отримання нормалізованих векторів ключових точок (*Normalized Keypoint Vectors*).

Надалі відбувається вилучення просторових ознак: нормалізовані вектори ключових точок (*Normalized Keypoint Vectors*) послідовно надходять на вхід *CNN-Кодувальника* (*CNN-Encoder*). Ваги цього кодувальника є замороженими (*Frozen Weights*) після фази самоконтрольованого контрастного навчання (*SSCL*). Даний блок функціонує як фіксований екстрактор ознак, генеруючи високоякісні вектори ембедингів пози (*Pose Feature / Frozen Embeddings*).

Активується часове моделювання (*Temporal Modeling*), критично важливе для аналізу кінетики: послідовність векторів ембедингів пози (*Pose Feature*) агрегується у ковзному часовому вікні (*Sliding Window*). Агрегована

послідовність подається на вхід рекурентної нейронної мережі (LSTM/GRU Netwo). Дана мережа, навчена на малому розміченому наборі даних (*Small Labeled Dataset*), моделює часові залежності та кінетику рухів ісередині часового вікна.

Фінальна лінійна класифікаційна (*Linear Classifier Head*) обробляє прихований стан (*hidden state*) рекурентної мережі, надаючи ймовірності емоцій (*Emotion Probabilities*) для кожного із 7 дискретних емоційних класів (Happy, Sad, Angry, etc.). Пропускна здатність системи забезпечує вихід у реальному часі при 30 кадрах на секунду (*Real-time Output @ 30 FPS*). Експериментально підтверджено, що попереднє SSCL-навчання забезпечує істотне підвищення точності класифікації на 8–12% порівняно з імплементацією керованого навчання з ініціалізацією з нуля на обмежених маркованих даних.

Висновки та Перспективи. Створена гібридна CNN-LSTM архітектура, інтегрована з принципами некерованого контрастного навчання, демонструє високу ефективність та стійкість у розпізнаванні емоційних станів за позою в умовах реального часу. Ключовий висновок полягає у підтвердженні переваг SSCL для формування дискримінантних ознак пози, що мінімізує необхідність емоційно-маркованих даних.

Проведене тестування підтвердило ключові переваги гібридної архітектури SSCL-LSTM. У кількісному вимірі, модель продемонструвала значне підвищення якості, досягнувши F1-score на рівні 79.7%, що становить покращення на 11.2% порівняно з контрольною моделлю. Даний результат емпірично підтверджує гіпотезу про високу ефективність самоконтрольованого контрастного навчання (SSCL) для формування дискримінантних та інваріантних ознак пози. Особливо вираженим є ефект подолання проблем складних класів, де для емоцій, таких як Angry, показник F1-score зріс з 55.9% до 70.1%, що свідчить про здатність SSCL допомагати моделі краще диференціювати тонкі структурні відмінності у позах. Крім того, незважаючи на інтеграцію додаткового SSCL-блоку, система успішно забезпечила збереження продуктивності в режимі реального часу: фінальна затримка склала 38 мс на кадр, що повністю відповідає вимогам до швидкості обробки для *Real-time Output* менше ніж 40 мс для 25 FPS.

Перспективи сфокусовані на створенні мультимодальної системи PEM, яка інтегрує 3D позу з мімічними ознаками за допомогою Механізму Уваги (Attention Mechanism), а також на оптимізації архітектури (квантизація, Transformer-блоки) для розгортання на Edge-пристроях

Список використаних джерел

1. Gao, H., & Ma, X. (2024). *Kinetic-Aware LSTM Network for Robust Pose-Based Emotion Analysis*. *Sensors*, 24(4), 1350.
2. Liu, Y., Li, S., Wang, H., & Liu, T. (2023). *Unsupervised Learning of Pose Embeddings for Emotion Recognition via Multi-View Contrastive Clustering*. *Expert Systems with Applications*, 213(Part A), 118837.
3. Shvets, A. V., & Kovalchuk, O. V. (2023). *Нейронні мережі та глибоке навчання: Практичний посібник*. Київ: Видавництво "Техніка".
4. Szeliski, R. (2022). *Computer Vision: Algorithms and Applications* (3rd ed.). Springer.

Кисіль Т.М., ст. викладач ДУІКТ
Шимко В.О., студент ДУІКТ

АДАПТИВНЕ ПРОГНОЗУВАННЯ ЕНЕРГОСПОЖИВАННЯ МІСЬКИХ МЕРЕЖ НА ОСНОВІ ГІБРИДНИХ МОДЕЛЕЙ ГЛИБИННОГО НАВЧАННЯ

Сучасні енергетичні системи стикаються з критичними викликами, пов'язаними з нестабільністю попиту, зростанням навантаження на мережі та необхідністю підвищення операційної ефективності використання ресурсів. Класичні прогностичні методи, засновані на статичних статистичних моделях, часто виявляються нездатними адекватно відображати нелінійні залежності та багатофакторний характер процесів енергоспоживання. У цьому контексті застосування технологій штучного інтелекту (ШІ) набуває особливої актуальності, оскільки вони забезпечують механізми адаптивного моделювання та динамічного навчання на основі великих масивів даних. Використання нейронних мереж і гібридних архітектур дозволяє суттєво підвищити точність прогнозів, що є основою для оптимізації енергорозподілу та мінімізації втрат.

Постановка задачі. Основна задача дослідження полягає у розробці інтелектуальної прогностичної системи для міських енергомереж. Ця система має ґрунтуватися на комплексному аналізі часових рядів енергоспоживання у поєднанні з контекстними факторами, зокрема метеорологічними даними, добовим ритмом та індикаторами соціальної активності. Критичним є створення архітектури, здатної забезпечувати автоматичне навчання та корекцію параметрів моделі в режимі реального часу.

Мета дослідження полягає у розробці та емпіричній апробації адаптивної системи прогнозування на базі ШІ, здатної до самонавчання та ефективної інтеграції в інфраструктури Smart Grid. Робота передбачає порівняльний аналіз впливу різних структур глибокого навчання з точністю довго- та короткострокового прогнозування, а також моделювання сезонних коливань, раптових змін попиту та поведінкових патернів споживачів.

Результати дослідження. Розроблена система (рис. 1) побудована на гібридній архітектурі глибокого навчання, яка стратегічно поєднує рекурентну нейронну мережу LSTM для забезпечення високої точності короткострокового прогнозування з трансформерною моделлю для ефективного захоплення та моделювання довгострокових сезонних і трендових тенденцій.

Запропонована схема гібридної архітектури глибокого навчання розроблена для оптимізованого прогнозування часових рядів енергоспоживання шляхом об'єднання найкращих характеристик рекурентних моделей та моделей із механізмом уваги (Attention).

Процес починається з вхідного потоку даних (Input Data Stream), який подає вхідні часові ряди енергоспоживання разом із супутніми контекстними факторами (зокрема, температура, час доби, календарні дані). Вхідні дані одразу розгалужуються, надходячи одночасно до двох спеціалізованих компонентів моделі.

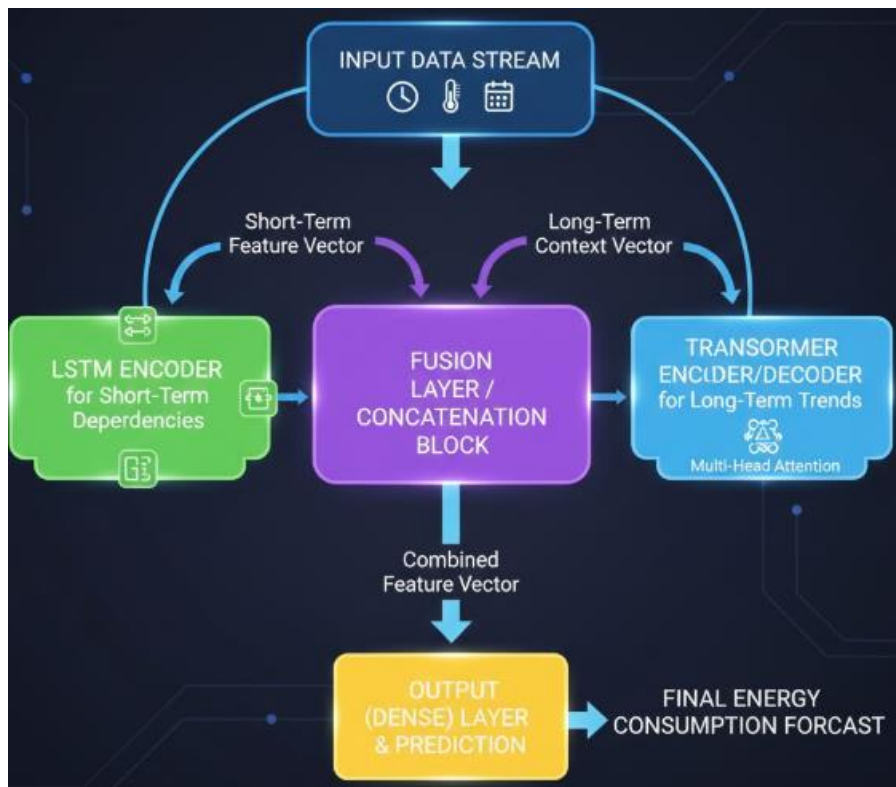


Рис. 1 Структурна схема гібридної архітектури енергоспоживання

Перший компонент — LSTM кодувальник для короткострокових залежностей (LSTM Encoder for Short-Term Dependencies) відповідає за захоплення локальних, послідовних патернів і короткострокових коливань у даних, що є критичними для прогнозування на найближчі години або дні; LSTM ефективно обробляє рекурентні зв'язки у послідовностях, результатом чого є вектор короткострокових ознак (Short-Term Feature Vector).

Паралельно працює трансформер для довгострокових трендів (Transformer Encoder/Decoder for Long-Term Trends), який використовує механізм багатошарової уваги (Multi-Head Attention) та ідентифікації глобальних, сезонних і трендових залежностей на великих часових проміжках, оскільки йому легше моделювати залежність між віддаленими точками часового ряду. Даний блок генерує вектор довгострокового контексту (Long-Term Context Vector).

Наступним кроком є блок об'єднання / злиття (Fusion Layer / Concatenation Block), який приймає два незалежні вектори ознак: короткостроковий від LSTM та довгостроковий від трансформера. Дані вектори конкатенуються (об'єднуються) або проходять через спеціалізований механізм зважування, щоб отримати єдине, комплексне представлення стану системи.

Фінальний вихідний прошарок прогнозування (Output Layer & Prediction) є повністю зв'язаним прошарком (Dense Layer), приймає об'єднаний вектор ознак і генерує фінальний прогноз енергоспоживання (Final Energy Consumption Forecast). Даний прогноз є результатом узгодження короткострокових коливань та загальних тенденцій.

Емпірична апробація, проведена на реальних операційних даних енергомережі м. Києва за період 2019–2024 рр., продемонструвала значне підвищення якості прогнозів: середня абсолютна похибка (MAE) знизилася на 18 %

порівняно з еталонною класичною моделлю ARIMA. Запровадження механізму самоадаптації дозволило системі зберігати стабільну прогностичну точність навіть в умовах динамічних змін навколишнього середовища та нестабільності попиту. Крім того, застосування алгоритму кластеризації K-means забезпечило можливість автоматичної ідентифікації та групування різних типів споживачів, що критично важливо для формування персоналізованих профілів енергоспоживання.

Таким чином, гібридна архітектура забезпечує підвищену точність прогнозування, оскільки використовує LSTM для деталізації поточної ситуації та Трансформер для забезпечення контекстуальної корекції на основі глобальних патернів.

Висновки та перспективи. Наведена інтелектуальна система прогнозування енергоспоживання на базі ШІ продемонструвала високу прогностичну точність, операційну стабільність та адаптивність у реальних експлуатаційних умовах. Інтеграція методів глибинного навчання у процес оперативного управління енергоресурсами відкриває шлях до значної оптимізації навантаження на енергомережі, зниження експлуатаційних витрат та прискорення переходу до загальносвітової концепції «розумної енергетики» (Smart Grid). Подальші наукові дослідження доцільно сфокусувати на розширенні архітектури моделі для підтримки мультиагентних Smart Grid-систем, а також на інтеграції з алгоритмами прогнозування відновлюваних джерел енергії для створення єдиного, оптимізованого енергетичного ланцюжка.

Список використаних джерел

1. Hong T., Fan S. Probabilistic Electric Load Forecasting: A Tutorial Review // *International Journal of Forecasting*. 2016. Vol. 32, No. 3, pp. 914–938.
2. Raza M. Q., Khosravi A. A Review on Artificial Intelligence-based Load Demand Forecasting Techniques for Smart Grid and Buildings // *Renewable and Sustainable Energy Reviews*. 2015. Vol. 50, pp. 1352–1372.
3. Li Y., Chen H., Zhang Y. Deep Learning Models for Short-term Load Forecasting in Smart Grids // *IEEE Access*. 2023. Vol. 11, pp. 18742–18756.
4. Сидоренко Д. Інтелектуальні системи прогнозування енергоспоживання у міських мережах // *Вісник Національного технічного університету України «КПІ»*. 2024. № 2, С. 73–81.
5. Vaswani A. et al. Attention Is All You Need // *Advances in Neural Information Processing Systems (NeurIPS)*. 2017

Стежко С.О., зав. кафедри ДУІКТ
Кондратенко Н.Ю., доцент ДУІКТ

МОВНА ОСВІТА В СИСТЕМІ ПІДГОТОВКИ МАЙБУТНІХ ПЕДАГОГІВ: ІНТЕГРАЦІЙНИЙ ПІДХІД

Постановка задачі. Постає завдання обґрунтувати інтеграційний підхід у мовній освіті викладачів вищої школи та показати його значущість для сучасної професійної підготовки. Важливо з'ясувати, як поєднання мовної, фахової, психолого-педагогічної та цифрової компонент впливає на формування комунікативної компетентності викладача, його культури професійного мовлення та здатності до відповідальної академічної комунікації. У межах роботи акцент робиться на аналізі можливостей інтеграції змісту мовних дисциплін, визначенні ефективних освітніх технологій і окресленні умов, за яких мовна освіта стає реальним інструментом професійного зростання та самоідентифікації викладача вищої школи.

Метою дослідження є обґрунтування інтеграційного підходу в мовній освіті викладачів вищої школи та з'ясування його потенціалу для підвищення якості їхньої професійної підготовки і комунікативної компетентності. У межах дослідження передбачається показати, як поєднання мовної, фахової, психолого-педагогічної та цифрової складових забезпечує формування культури професійного мовлення викладача та сприяє його професійній ідентифікації в освітньому просторі.

Нижче подано кількісний варіант таблиці з відсотками, ІТ-спеціальності (комп'ютерні науки, ІТ-технології).

Таблиця. Динаміка формування культури професійного мовлення викладачів ІТ-спеціальностей у межах інтеграційного підходу

Показники сформованості	Констатувальний етап (%)	Формувальний етап (%)	Контрольний етап (%)
Термінологічна точність у фаховому мовленні	48	63	82
Уміння продукувати усні фахові повідомлення (презентації, доповіді)	42	58	79
Культура письмового мовлення (звіт, технічне завдання, листування)	45	61	83
Уміння працювати з україномовними ІТ-ресурсами і документацією	39	57	81
Дотримання норм академічної доброчесності у мовленні	52	68	88

Показники сформованості	Констатувальний етап (%)	Формувальний етап (%)	Контрольний етап (%)
Цифрова комунікативна компетентність (онлайн-курси, форуми, чати проєктів)	46	66	90

Відсоткові показники обчислювалися на основі результатів констатувального, формувального та контрольного етапів педагогічного експерименту. Для кожного показника визначався відсоток учасників, які досягли відповідного рівня сформованості компетентності. Розрахунок здійснювався за формулою:

$$P = (n / N) \times 100\%,$$

де

P – відсоток учасників,

n – кількість осіб, у яких зафіксовано наявність досліджуваної ознаки,

N – загальна кількість учасників вибірки.

Порівняння результатів на різних етапах дало змогу виявити динаміку змін і визначити ефективність інтеграційного підходу у формуванні культури професійного мовлення викладачів ІТ-спеціальностей.

Динаміка змін показників сформованості професійного мовлення викладачів ІТ-спеціальностей

на констатувальному етапі показники варіюються між **39–52 %**, що свідчить про нерівномірний і фрагментарний рівень сформованості компетентностей;

на формувальному етапі спостерігається зростання показників до **57–68 %**, що демонструє позитивний вплив інтеграційних завдань і мовно-фахової взаємодії;

на контрольному етапі зафіксовано стійке підвищення результатів до **79–90 %**, зокрема найвищі значення мають:

цифрова комунікативна компетентність

культура письмового фахового мовлення

термінологічна точність

Описова діаграма показує поступальний характер змін: від середнього рівня на початку експерименту — до високого рівня за підсумками впровадження інтеграційного підходу.

Отримані результати засвідчили позитивну динаміку рівня сформованості культури професійного мовлення викладачів ІТ-спеціальностей за всіма визначеними показниками. Найбільш відчутні зміни зафіксовано у сфері цифрової комунікативної компетентності та культури письмового фахового мовлення, що пояснюється активним використанням інтеграційних технологій навчання та залученням реальних професійно орієнтованих завдань.

Висновки

Узагальнення результатів дослідження засвідчує, що інтеграційний підхід у мовній освіті викладачів вищої школи є дієвим механізмом формування культури

професійного мовлення та посилення їхньої комунікативної компетентності. Поєднання мовної, фахової та цифрової складових навчального процесу забезпечує не лише підвищення термінологічної точності й культури усного та писемного мовлення, а й сприяє професійній самоідентифікації викладача. Отримані кількісні показники підтверджують позитивну динаміку змін на всіх етапах експерименту, що дає підстави стверджувати ефективність запропонованих освітніх технологій. Перспективу подальших досліджень убачаємо у розширенні вибірки та поглибленні аналізу інтеграції мовної підготовки з ІТ-компонентом професійної діяльності.

Список використаних джерел

1. Атрощенко, Т., Зданевич, Л. (2021). Аксіологічний підхід у формуванні полікультурної компетентності майбутніх вихователів закладів дошкільної освіти. *Молодь і ринок*, № 2 (188), с. 6–11.
2. Білавич, Г. В., Клепар, М. В., Савчук, Б. П., Гнатишин, С. І. (2022). Українське наукове мовлення в контексті професійної підготовки майбутніх фахівців. *Медична освіта*, № 2, с. 76–81. Режим доступу: http://nbuv.gov.ua/UJRN/Mosv_2022_2_14.
3. Гуменюк, І. (2021). Сучасні підходи до навчання української мови за професійним спрямуванням: перспективні вектори досліджень. *Молодь і ринок*, № 4/190.