

Міністерство освіти і науки України

Системні технології

System technologies

4 (147) 2023

Регіональний міжвузівський збірник наукових праць

Засновано у січні 1997 року.

У випуску:

- ПРОГРЕСИВНІ ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ
ТА ОРГАНІЗАЦІЯ СУЧАСНОГО ВИРОБНИЦТВА**
- МАТЕМАТИЧНЕ ТА ПРОГРАМНЕ ЗАБЕЗПЕЧЕННЯ
ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ**
- СИСТЕМНІ ТЕХНОЛОГІЇ ОБРОБКИ ІНФОРМАЦІЇ
ТА КІБЕРБЕЗПЕКА**

Системні технології. Регіональний міжвузівський збірник наукових праць. – Випуск 4 (147). - Дніпро, 2023. – 181 с.

ISSN 1562-9945 (Print).

ISSN 2707-7977 (Online).

Редакційна колегія випуску:

Алпатов А.П. - д.т.н., проф. (відп. редактор)
Архипов О.Є. - д.т.н., проф.
Білозьоров В.Є. - д.ф.-м.н., проф.
Бабічев С.А. (Чеська Республіка) - д.т.н., доц.
Єрьомін О.О. - д.т.н., проф.

Прогресивні інформаційні
технології та організація
сучасного виробництва

Гече Ф.Е. - д.т.н., проф., (відп. редактор)
Гуда А.І. - д.т.н., проф.
Гнатушенко Вік.В. - д.т.н., проф.
Скалозуб В.В. - д.т.н., проф.

Математичне
та програмне забезпечення
інтелектуальних систем

Гнатушенко В.В. - д.т.н., проф., (відп. редактор)
Гожий О.П. - д.т.н., проф.
Кіріченко Л.О. - д.т.н., проф.
Светличний Д.С. (Польща) - д.т.н., проф.
Хандецький В.С. - д.т.н., проф.

Системні технології
обробки інформації
та кібербезпека

Збірник друкується за рішенням Вченої Ради
Українського державного університету науки і технологій
від 23.01.2023 р., № 4

Адреса редакції: 49600, Дніпро, пр. Гагаріна, 4
Український державний університет науки і технологій,
ННІ «Інститут промислових та бізнес технологій»
кафедра Інформаційних технологій та систем.
Тел. +38(097)6854525
E-mail: st@nmetau.edu.ua
<https://journals.nmetau.edu.ua/index.php/st>

© Український державний університет науки і технологій,
ННІ «Інститут промислових та бізнес технологій»,
ІБК «Системні технології», 2023

В.В. Скалозуб, В.М. Горячкін, О.В. Мурашов

РЕЛЯЦІЙНО-СЕПАРАБЕЛЬНІ МОДЕЛІ ПРОЦЕСІВ МОНІТОРИНГУ ПРИ ПЕРЕМІННИХ І НЕЧІТКИХ ІНТЕРВАЛАХ СПОСТЕРЕЖЕНЬ

Анотація. Стаття присвячена розвитку комбінованих моделей, методів і засобів, призначених для вирішення актуальних завдань моделювання та аналізу даних процесів моніторингу, які представлені часовими рядами і відрізняються перемінним або нечітким інтервалом спостережень (ЧРПНІ). В наших попередніх дослідженнях для реалізації завдань аналізу і прогнозування характеристик ЧРПНІ була запропонована сепарабельна модель (СПМ), а також удосконалений квантильний алгоритм. При реалізації процесів моніторингу з нечітким кроком інтервалів спостережень застосовувався підхід на основі декомпозиції за допомогою α рівнів. Засобами моделі СПМ і квантильного алгоритму були досліджені дані клінічного моніторингу процесів реабілітації хворих на діабет.

В цій роботі для підвищення точності та ефективності моделювання і аналізу процесів ЧРПНІ запропоновані нові реляційно сепарабельна модель (РСМ) і комбінований квантильний алгоритм. Реляційна модель визначається системою нечітких реляційних відношень першого та другого порядку, отриманих на основі вихідної послідовності даних. У комбінованому алгоритмі результати розрахунків, отримані за СПМ і моделями нечітких реляційних відношень, узгаляювалися при оптимальному виборі вагових коефіцієнтів для окремих складових.

В результаті виконаних досліджень шляхом числового моделювання було встановлено, що за провадження комбінованих моделей процесів при ЧРПНІ являється раціональним та результативним. Приклади аналізу даних моніторингу процесів реабілітації хворих на діабет показали певні можливості забезпечення вимоги до точності результатів аналізу показників та їх короткострокового прогнозування.

Ключові слова: процеси моніторингу, нерівномірний та нечіткий інтервал вибірки, сепарабельна модель, нечіткі реляційні відношення, реляційно сепарабельна модель, комбінований квантильний алгоритм, моніторинг стану хворих.

Вступ та постановка проблеми. Комп'ютерне моделювання в сучасних технологіях являється майже необхідним при збиранні даних про процеси функціонування, при їх аналізі, інтерпретації та прогнозуванні, створенні складних технологічних та автоматизованих процесів та систем. Разом з тим відбувається постійне підвищення вимог до повноти характеристик процесів, що приводить до необхідності урахування все більшого набору властивостей досліджуваних технологій,

процесів і систем. Натепер існують процеси (клінічний моніторинг процесів реабілітації, контроль стану віддалених об'єктів, моніторингу функціонування та параметрів експлуатації систем з великою кількістю незалежно функціонуючих елементів, процесів моніторингу технічного стану систем, які обслуговуються на основі параметрів «поточного стану» тощо), для яких можливо отримати дані, які мають нечітку, розмиту, структуру. Зазначене показує актуальність завдань щодо аналізу, моделювання і прогнозування характеристик недетермінованих процесів, представлених часовими послідовностями при нерівномірних та нечітких моделях інтервалів (ЧРПНІ) між спостереженнями.

У цій статті представлені результати дослідження можливостей та ефективності комплексних алгоритмів із моделювання та прогнозування процесів із нечіткими ознаками. Розглянуті та реалізовані завдання зі створення процедур аналізу та прогнозування характеристик недетермінованих часових рядів на основі FTS алгоритмів, а також їх комбінацій з іншими. Виконана розробка нечітких реляційно-сепарабельних моделей (PCM) першого та другого порядку, призначених для аналізу та прогнозування послідовностей даних, отриманих та формалізованих моделями з «нечітким кроком» між спостереженнями. На основі PCM моделей процесів моніторингу та реляційно-сепарабельних алгоритмів були реалізовані завдання дослідження недетермінованих послідовностей з перемінним і нечітким кроком між спостереженнями.

Аналіз останніх досліджень і публікацій. У сучасних складних системах деякі параметри станів або контрольовані характеристики процесів можуть мати значну ступінь невизначеності. При цьому на практиці процедури моніторингу таких процесів часто дають змогу отримати лише короткі часові послідовності даних з нерівномірною у часі вибіркою. Такі короткі нерівномірні (також нечіткі) у часі послідовності даних (НЧПД) не дозволяють використовувати для аналізу традиційні статистичні моделі [1 – 3]. Для вирішення завдань аналізу та прогнозування НЧПД натепер застосовують кілька підходів [1– 3, 10, 12]. Відзначимо серед них включення часової інформації у метрики відстані, що використовуються для кластеризації часових рядів (ЧР) [12]. Застосування для моделювання і прогнозування нечітких ЧР процедур із встановлення нерівних частин областей універсуму дискурсу, з подальшою оцінкою взаємозв'язків між послідовними рівнями засобами генетичного алгоритму [10]. Створюються нові структури моделювання, що забезпечують економію моделей нечітких часових рядів (FTS) при збереженні певного рівня точності поза вибіркою [11] та інше.

У попередніх дослідженнях нами була запропонована нова математична модель НЧПД, яка для моделювання використовує сепарабельні форми обліку часових інтервалів між рівнями ряду. Модель СПМ [4] передбачає формуванням послідовностей величин показників та інтервалів між спостереженнями окремо, з подальшим їх співставленням на основі параметру часу (кроку процесу). Ефективність методу СПМ залежить від алгоритмів моделювання нечітких часових рядів. СПМ процесів моніторингу у цілому визначає як значення нового моменту виникнення чергової події спостережуваного процесу, так і відповідні кожному моменту характеристики процесу. У виконаних дослідженнях СПМ була застосована для моделювання процесів клінічного моніторингу стану хворих на діабет. При цьому головна мета аналізу полягала у прогнозуванні максимального (нечіткого) періоду до подій, які відповідають встановленим вимогам. В статті [4] також отримано удосконалення алгоритму (FTS) шляхом поєднання квантильної моделі [7, 14] з випадковими процесами спостережень, при тому встановлено числову ефективність цього алгоритму.

Дослідженням проблем і питань моделювання процесів з нерівномірними та нечіткими інтервалами спостереженнями приділяється все більше уваги [8 – 12]. Відзначимо роботу [1], в якій зазначається, що на практиці багато ЧР даних містять спостереження в нерегулярний час, тоді як більшість методів прогнозування обмежені випадком рівних інтервалів часу між точками даних. В зазначеній роботі виконано розширення однокоекспоненціального згладжування та методу Холта для випадку нерівномірно розподілених даних. Також показано його високу обчислювальну ефективність. Новий метод був застосований до шести опублікованих серій, і їх ефективність аналізується через помилки щодо змін параметрів згладжування.

Стаття [2] дозволяє визначити, що більшість моделей часових рядів передбачає, що дані надходять із спостережень, які однаково рознесені в часі. Однак це припущення не відповідає багатьом різноманітним науковим галузям, таким як астрономія, фінанси, кліматологія тощо. Існують деякі методи, які підбирають часові ряди з нерівним інтервалом, наприклад процеси безперервної авторегресії ковзного середнього (CARMA). При тому зазначено що єдині серед відомих в літературі методів прогнозування для даних з нерегулярним інтервалом у часі засновані на фільтрі Калмана. У цій роботі [2] запропонована нова модель для відповідності нерегулярним часовим рядам, яка називається комплексною нерегулярною авторегресійною моделлю (CIAR), яка представлена безпосередньо як дискретний часовий процес. Показано, що модель CIAR слаб-

ко стаціонарна, її можна представити як систему простору станів. Це дозволяє ефективно оцінювати максимальну правдоподібність на основі рекурсії Калмана. Встановлено за допомогою моделювання Монте-Карло, що продуктивність кінцевої вибірки з оцінки параметра є точною. Модель CIAR, через представлення простору станів, дозволяє прогнозувати неспостережувані астрономічні вимірювання.

Усереднення динамічної моделі (DMA) широко використовується для цілей економічного прогнозування. Дослідження [3] розширює рамки DMA, вводячи адаптивне навчання з простору моделі. У структурі DMA всі моделі оцінюються незалежно, інформація в них щодо інших моделей залишається невикористаною. У статті [3] запропоновано усереднення за прогнозами, а також динамічне усереднення за параметрами, що змінюються в часі. Актуальність цього розширення проілюстровано трьома емпіричними прикладами, пов'язаними з прогнозуванням інфляції в США, споживчих витрат тощо. Показано, що адаптивне навчання з простору моделі забезпечує покращення ефективності прогнозування.

Моделювання нерівномірних у часі вибірок з метою удосконалення результатів кластеризації ЧР на основі включення часової інформації у метрики відстані запропоновано у [9]. Визначається, що порядок даних і різні інтервали вибірки важливі, і вони повинні враховуватися при кластеризації ЧР. Шляхом включення нової моделі відстані в стандартну схему нечіткої кластеризації був розроблений удосконалений алгоритм, а також продемонстрована його продуктивність. Концепція моделювання і аналізу НЧР, яка ураховує прогноз тенденції процесу з використанням нечітких відносин третього порядку, була запропонована в [8]. Результати застосування моделей з прогнозуванням тренду показали переваги цього методу за точності результатів прогнозування та складністю процесу моделювання. Аналіз публікацій свідчить про значну актуальність завдань щодо методів аналізу процесів при нерівномірних та нечітких у часі спостереженнях, суттєву обмеженість підходів до моніторингу таких процесів.

Мета дослідження. Метою цього дослідження являється підвищення точності та ефективності моделювання і аналізу процесів, представлених часовими рядами при перемінному та нечіткому інтервалі спостережень, шляхом формування нової реляційно-сепарабельної моделі процесів (PCM) і комбінованого квантильного алгоритму. При тому нечітка реляційна модель (HPM), як складова PCM, визначається системою нечітких реляційних відношень першо-

го та другого порядку, отриманих на основі вихідної послідовності даних. У комбінованому квантильному алгоритмі передбачається узагальнення результатів розрахунків, отриманих окремо за СПМ та НРМ моделями, при оптимальному виборі вагових коефіцієнтів для складових на основі експонентного згладжування.

Також необхідно було розробити та шляхом числового моделювання дослідити процедури аналізу та короткострокового прогнозування процесів при ЧРПНІ. На прикладі аналізу даних моніторингу процесів реабілітації хворих на діабет потрібно показати можливості забезпечення вимоги до точності результатів аналізу показників та оцінок їх короткострокового прогнозування.

Результати та основний матеріал дослідження. Головне завдання, що було вирішене шляхом створення РСМ процесів з ЧРПНІ кроком спостережень, полягає в формуванні комплексних моделей процесів та алгоритмів, що поєднують реляційні алгоритми нечіткого моделювання з сепарабельною моделлю, формою відображення недетермінованих процесів. В рамках РСМ поєднуються результати застосування СПМ [14] з нечіткими реляційними моделями (НРМ) першого та другого порядку, що представлені в роботі далі. НРМ процесу визначається системою реляційних відношень, отриманих на основі вихідної послідовності даних. яку узагальнено можливо представити так: $R(V, T) = \{R_1(V, V); R_2(T, T); R_3(V*V, V); R_4(T*V, T); R_5(T*V, V); R_6(T*T, V*V, V); R_7(T*T, V*V, T)\}$. У відношенні $R(V, T)$ параметр «V» відповідає певному контрольованому показнику процесу (наприклад, рівень цукру в крові), а «T» - показник часу. Відношення $R_6(*)$ та $R_7(*)$ НРМ відповідають моделям другого порядку, враховують кілька попередніх етапів розвитку процесів. Модель $R(V, T)$ формується за первинними даними процесів, які перетворюються до табличних форм на основі нечітких моделей областей варіювання показників (V, T). Таблиці $R(V, T)$ представляють нечіткі «правила» зв'язку між показниками (V, T) на послідовних кроках досліджуваного процесу. Кожна клітина таблиць містить два параметри – «номер нечіткого терму» що має бути на наступному кроці процесу, а також нечітка міра для очікуваного терму. Наприклад, узагальнено деякий крок процедури розрахунків в рамках НРМ можна представити так $R_7(tp* tq, vp* vq, ts/vs)$. Значення (ts/vs) використовуються для розрахунку поточних оцінок показників процесу (V, T), а також для визначення наступних кроків процесу. Результати розрахунків за моделями НРМ та СПМ узагальнюються шляхом оптимального вибору вагових коефіцієнтів за процедурою експонентного згладжування.

В першу чергу відзначимо, що реалізація досліджуваних процесів ЧРПНІ основана на методології FTS [1] та удосконаленій формі квантильної моделі [8], що була запропонована в [4]. В них використовують нечіткий ряд НЧР даних $F(t)$, а також його моделі першого порядку $R(t, t-1)$. Якщо також визначають для $F(t)$ нечіткі множин попередніх етапів $F(t-1)$, $F(t-2)$, ..., $F(t-n)$, у підсумку в якості моделі має місце відношення n-го порядку

$$F(t-n), \dots, F(t-2), F(t-1) \rightarrow F(t). \quad (1)$$

Процедура квантильної регресії третього порядку з трендом для (1) запропонована в [8] і удосконалена в [4, 14]. Для ефективного (по точності результатів) застосування квантильної моделі та модифікованих алгоритмів при реалізації завдань СПМ було виконано подальше удосконалення квантильного алгоритму [8]. Змістовно модифікація нечітких процедур моделювання полягала в узагальненні (агрегуванні) результатів, отриманих за різними схемами квантильного алгоритму. Формально запропоноване поєднання являє зважену згортку результатів, отриманих на основі серединх інтервалів m_j [8] з результатами моделювання за алгоритмами [14]. Підвищення ефективності алгоритмів моделювання на основі схем узагальнення виявилось результативним. Шляхом числових досліджень кількох схем агрегування була визначена краща форма, представлена у наступному вигляді

$$Y_k = \alpha V_k + (1-\alpha)u_k \quad 0 \leq \alpha \leq 1 \quad (2)$$

За (2) результат моделювання на кроці «k» формується на основі відповідного центру інтервалу m_k (складова V_k) та шляхом заміни серединх m_j інтервалів квантування u_j на значення u_k рівнів вхідного ЧР (складова u_k). Величина параметру « α » підбиралась за процедурою мінімізації квадрату середньої похибки. В представленій статті моделі та процедури РСМ використовують схему виду (2) для агрегування результатів СПМ та нечітких реляційних моделей НРМ.

Наступним питанням щодо формування РСМ являється визначення тих особливостей, властивостей процесів ЧРПНІ, що відображає НРМ *додатково* до моделей СПМ. Для пояснення цього нагадаємо твердження роботи [8], в якій зазначено, що змістовною передумовою, «аксіоматичною підставою» щодо можливості існування сепарабельної моделі складних процесів, являється *системна єдність всіх властивостей* досліджуваного процесу. При тому в СПМ контрольовані властивості через високу складність зв'язків подаються *окремими* характеристиками у вигляді часових рядів. Таким чином, *змінні інтервали* спо-

стерезень також є окремою ознакою процесу, його властивістю поряд з іншими характеристиками. Тому сепарабельну модель (СПМ) процесу моніторингу ЧРПНІ $F(t)$ представляють у вигляді

$$SpM(t) = SpM_i(k) \cup \{SpM_i(SpM_i(k))\}, i=1, 2, \dots, q \quad (3)$$

де позначено, як t – координата часу, i – індекси вектору характеристик процесу моніторингу, k – порядковий номер рівня вибірки даних $SpM(k)$. Таким чином СПМ (11) утворюється шляхом поєднання моделей окремих характеристик $F(t)$ розрахованих для моментів, визначених на основі нечіткої моделі часу процесу $SpM_i(k)$.

Задача формування РСМ полягає в уведенні до моделей процесів, представлених ЧРПНІ, додаткових об'єктивних властивостей у формі нечітких відношень, НРМ. При тому далі необхідно виконувати агрегування результатів, отриманих з використанням СПМ і НРМ, у формі подібній (2). Відзначимо суттєву відмінність правил, які застосовує структура нечітких реляційних відношень першого та другого порядку НРМ, від відношення n -го порядку (1). Процедури (1) встановлюють зв'язки між рівнями *одного* виміру, на відміну від цього модель НРМ визначає зв'язки між рівнями *різних* вимірів, окремими характеристиками моделей процесів ЧРПНІ, які узгоджені між собою за номером кроку (етапу процесу).

Загальна форма моделювання за РСМ для показника $V(*)$ на основі результатів сепарабельної $V(Y)$ та нечіткої реляційної моделі $V(\Delta)$ процесу, отримана шляхом процедури експонентного згладжування, представлена у вигляді (4), де параметр « w » розраховується методом найменших квадратів

$$V(Y, \Delta, w) = w \cdot V(Y) + (1-w) \cdot V(\Delta), \quad 0 \leq w \leq 1. \quad (4)$$

За (4) комбінуванням не можливо погіршити результати СПМ і НРМ.

Процедура формування і реалізації СПМ для показників $V(Y)$ приведена у [14]. Представимо структуру моделювання за НРМ для процесів з одновимірними показниками. При цьому вихідні дані представляють числові послідовності пар величин рівнів « k » – («інтервал», «показник»)

$$(T, V) = \{(\Delta T(k), \Delta V(k))\}, k = 1, 2, \dots, n. \quad (5)$$

Величини (T, V) можуть бути дійсними або нечіткими, а інтервали нерівномірними. При моделюванні за НРМ передбачається перехід від довільних типів даних (5) до нечітких рис.1, певного вигляду, трикутних рис.2. Для забезпечення цього застосовуються загальновідомі процедури [6]. Щоб спростити виклад графічно представимо лише результати формування моделей.

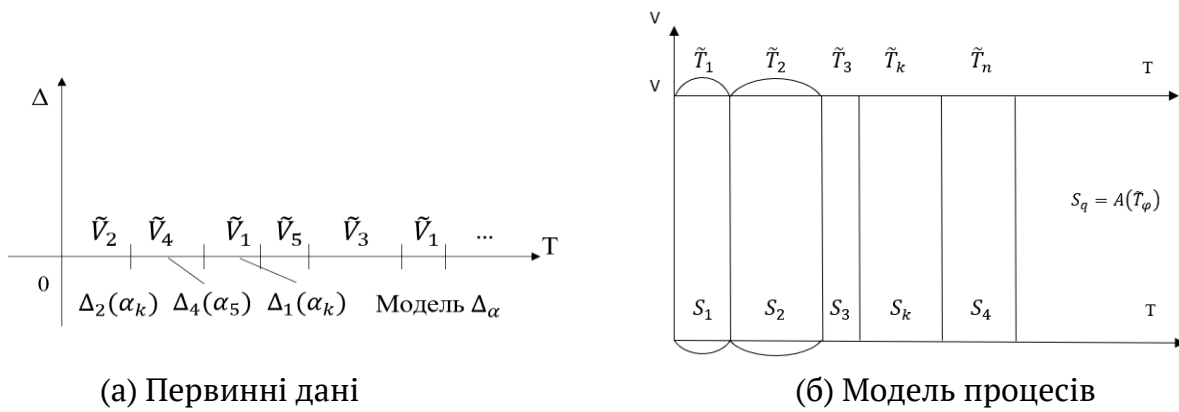


Рисунок 1 - Модель нечітких інтервалів та їх апроксимації

На рис. 1 приведені результати етапу попереднього моделювання даних (13) за допомогою системи нечітких трикутних величин рис. 2, які далі застосовуються для формування нечітких відношень першого та другого порядку, що утворюють НРМ. В результаті застосування апроксимацій S_1, S_2, S_3, S_4, S_5 рис. 2, відображають область можливих значень показників (Т, V), вхідним даним рівнів процесу приписують трійку ознак – («номер S_j », ступінь приналежності « μ_j », код формули трикутної моделі (ліва (значення менше центру, права – більше чи рівні))). Для зменшення числа показників при зберіганні прийнята наступна умова – для значень термів зліва від центру показник « μ_j » записують з мінусом, тобто « $-\mu_j$ », а для значень характеристик більших ніж центр « μ_j ». Саме такі показники використовуються далі при формуванні відношень НРМ і в процедурах моделювання, аналізу та прогнозування. Наприклад, для значення S_2 рис. 2 ознака приналежності до S_2 буде « $\mu_j = -0.7$ », а для S_3 « $\mu_j = 0.9$ ».

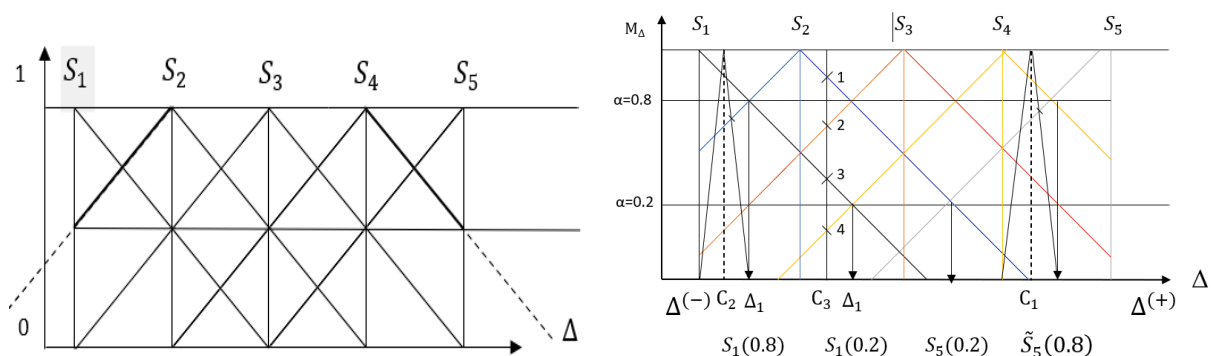


Рисунок 2 - Системи трикутних моделей для апроксимації «розмитой» області змін інтервалів між спостереженнями S_1, S_2, S_3, S_4, S_5 ..

Зауважимо, що наведена вище форма представлення апроксимацій S_1, S_2, S_3, S_4, S_5 фактично реалізує принцип дефазифікації по «максимуму функції

приналежності», відповідно якому відбирається терм S_j з найбільшим показником ступеня приналежності « μ_j ».

Рисунок також демонструє перетворення нечіткої послідовності S_1, S_2, S_3, S_4, S_5 при моделювання на основі α -рівнів [6]. На рисунку величини C_1, C_2, C_3 являються центрами нечітких моделей, для яких необхідно отримати дійсні послідовності величин з перемінним кроком, використовуючи моделі апроксимації S_1, S_2, S_3, S_4, S_5 . На наступних етапах такі послідовності з перемінним кроком обробляються алгоритмами СПМ.

Відповідно до рис. 3 для різних (кожного) α -рівнів будуть отримані дві послідовності нерівномірних інтервалів – для лівого та правого рівняння моделей термів, як трикутних чисел рис. 2. Для прикладу наведені послідовності на рівнях α_1, α_2

$$\begin{aligned} a_1 : R_1^{(1)}(\Delta_1; \Delta_4; \Delta_8; \dots); R_2^{(1)}(\Delta_3; \Delta_7; \Delta_{10}; \dots); \\ a_2 : R_1^{(2)}(\Delta_1; \Delta_5; \Delta_9; \dots); R_2^{(2)}(\Delta_2; \Delta_6; \Delta_{10}; \dots); \end{aligned} \quad (6)$$

Послідовності типу (6) використовуються для відображення і

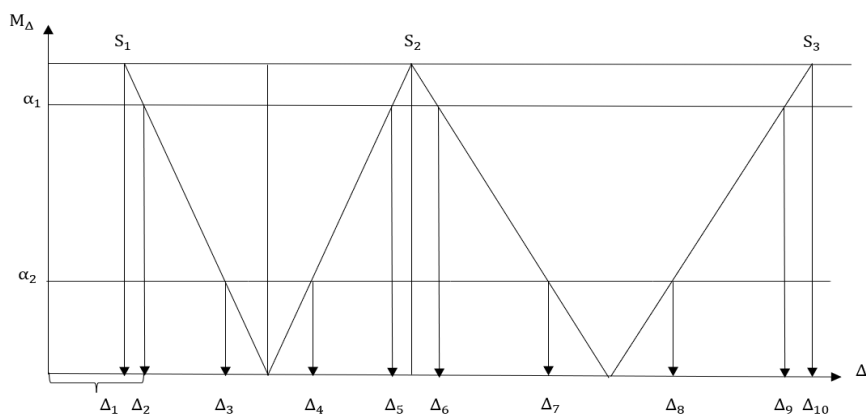


Рисунок 3 - Утворення α -рівневих послідовностей нерівномірних інтервалів

формування складових РСМ, і подальшого аналізу та прогнозування нечітких процесів у вигляді моделей з перемінним кроком між спостереженнями.

Головний зміст представленого підходу формування РСМ полягає в тому, щоб використовувати не лише сепарабельну модель СПМ, а комбінувати її результати з НРМ. На відміну від СПМ, в яких кожний параметр формується незалежно від інших, а вони узгоджуються на основі номерів спостережень (11), в НРМ безпосередньо формується модель зв'язків між нечіткими кроками у часових інтервалах та змінами контрольованого параметру процесів. Модель НРМ змістовно відтворює нечітку регресію, яка дозволяє на основі поточних та попередніх значень (номери термів нечіткої величини S_j) характеристик процесу визначити, до яких нечітких величин процес «перейде» на наступному

своєму кроці, етапі. То ж за попередніми НВ (S_i) будуть відтворені терми (дефазифікація), які визначають НВ що очікувані на наступному кроці процесу.

Як зазначено, НРМ процесу являє систему реляційних відношень виду:

$$R(V, T) = \{R_V(V, V); R_T(T, T); R_{V1}(V^*V, V); R_{TV}(T^*V, T); R_{VT}(T^*V, V); R_{V2}(T^*T, V^*V, V); R_{T2}(T^*T, V^*V, T)\}. \quad (7)$$

Компоненти (7) $R_V(V, V)$, $R_T(T, T)$ не мають зв'язків з попередніми етапами, мають «нульовий» порядок; порядок інших відношень (7) визначається числом попередніх етапів, що враховуються моделлю.

Процедури формування та використання компонентів (7) подібні між собою. Сутність та призначення цієї математичної моделі полягає в представлення процесів перетворення нечітких послідовностей апроксимацій S_1, S_2, S_3, S_4, S_5 для часових характеристик інтервалів (ПТ), а також таких характеристик для контрольованих показників процесів (ПР), у реляційну (табличну) модель як нечіткого відношення. В якості однієї (вхідної) координати відношення виступає або період ПТ (заданий можливими нечіткими величинами), або показник процесів (ПР) (свої нечіткі величини). На перетині рядку та стовпця стоїть значення очікуваного нечіткого показника, як результат прогнозування наступного кроку. Розглянемо основні варіанти (7).

Моделі нечітких відношень $R_{TV}(T^*V, T)$; $R_{VT}(T^*V, V)$ мають метою представлення зв'язків $(\Delta T_k, \Delta V_k \rightarrow \Delta T_{k+1}(I))$ також $(\Delta T_k, \Delta V_k \rightarrow \Delta T_{k+1}(II))$

Структура реляційних таблиць подана на рис. 4, де формування складових РСМ, і подальшого аналізу та прогнозування нечітких процесів у вигляді моделей з перемінним кроком між спостереженнями. На рис. 4 позначено

$$Z_i = \begin{cases} \tilde{T}_q, \text{ для моделі (I)} \\ \tilde{V}_r, \text{ для моделі (II)} \end{cases}$$

$\Delta V \backslash \Delta T$	$\tilde{V}_1$	$\tilde{V}_2$	$\tilde{V}_3$	$\tilde{V}_4$	$\tilde{V}_5$
T_1	Z_1	Z_5	Z_2	Z_5	Z_3
T_2	Z_3	..	..		
T_3					
T_4		..	..		
T_5	Z_3	Z_4	Z_1	Z_5	Z_2

Рисунок 4 - Структура із представлення змішаних відношень $R_{TV}(T^*V, T)$;

$R_{VT}(T^*V, V)$ (період часу / показник)

Структури відношень рис. 4 однакові, але значення Z_i побудовані за даними для ΔT , модель (I) або за даними для ΔV , модель (II). У змішаних відношеннях $R_{TV}(T^*V, T)$; $R_{VT}(T^*V, V)$ значення прогнозованих показників визначається за рахунок перевірки всіх величин вихідних даних для $\{T, V\}$, що відповідають номерам рівня $\langle k \rangle$, як моделі наступного кроку для показників процесу.

Приведемо модель та процедуру аналізу та прогнозування НЧП на основі реляційного відношення другого порядку, за якими наступний елемент (очікувана на наступному кроці нечітка величина) визначається по двох попередніх як для параметру часу (T), так і для контрольованого параметру (V). Змістовно модель реляційного відношення другого порядку формується таким чином, щоб відобразити вхідний процес, поданий через апроксимації S_1, S_2, S_3, S_4, S_5 , у вигляді табл. 1. А саме для всього процесу (подібно до формування таблиць рис. 4) послідовно визначають пари значень для $\{T, V\}$, через номери яких встановлюються рядки та стовпи табл. 1, в котрі необхідно ввести інформацію про наступний крок процесу (окремо для процесу T і V номери апроксимацій S_j , значення ступенів приналежності ін.). Процедура моделювання за табл. 1 наведена далі.

При неможливості застосування реляційної моделі табл.1, що потенційно можливо через, наприклад, відсутність в вихідній послідовності даних всіх передбачених табл.1 варіантів послідовностей (S_1, S_2, S_3, S_4, S_5), процедура рис. 4 використовується двічі. За ці два кроки наближено виконається урахування двох попередніх етапів процесу. При тому використовують сформовані таблиці безпосередньо для відношень першого порядку. Для реалізації моделі другого порядку використовують поточні значення вхідних (номер рядку, наприклад, i_1) та вихідних (номер стовпця, наприклад j_1) показників реляційних таблиць. При цьому на підставі значень (i_1, j_1), використовуючи таблиці для показників T, визначається номер рядку прогнозу, наприклад, i_2 , а за таблицею вихідних характеристик (номер стовпця, наприклад j_2). Після цього на підставі значень (i_1, j_1) визначаються номери величин, що визначають прогноз відповідно до моделі другого порядку, наприклад (i_3, j_3). Саме ці значення (**i_3, j_3**) мали зберігатися в табл. 1 для моделі відношення 2-го порядку, $(i_1, i_2, j_1, j_2) \rightarrow (i_3, j_3)$.

Процедура прогнозування на основі моделей процесу (8) відбувається таким чином. Нехай для поточного кроку процесу $\{T, V\}$ відомі попередньо визначені значення $\tilde{T}_{i_1}, \tilde{T}_{i_2}, \tilde{V}_{j_1}, \tilde{V}_{j_2}$. Тоді на основі продукційних моделей (8) можуть бути визначені НВ для V_{iq}, T_{ip} . За допомогою базових моделей цих НВ можна роз-

рахувати (функція дефазифікації [6]) числові параметри наступного кроку процесу. Це забезпечить можливість перейти до нового $\langle k+1 \rangle$ кроку (та до оцінки значень відповідних рівнів) процедур аналізу або прогнозування.

Таблиця 1

Структура моделі реляційного відношення другого порядку

step	(n-1)	V_1	V_1	V_1	V_2	V_2	V_2	V_3	V_3	V_3	...
(n-1)	(n)	V_1	V_2	V_3	V_1	V_2	V_3	V_1	V_2	V_3	...
T_1	T_1	$\Gamma_{11,11}$									
T_1	T_2			$\Gamma_{13,12}$							
T_1	T_3							$\Gamma_{11i2,j1j2}$			
T_2	T_1			$\Gamma_{21,13}$							
T_2	T_2										
T_2	T_3			$\Gamma_{11i2,j1j2}$			$\Gamma_{23,23}$				
T_3	T_1										
T_3	T_2							$\Gamma_{32,31}$			
T_3	T_3										

Реляційні відношення другого порядку визначають зв'язки процесу

$$\begin{aligned}
 R_v : \tilde{T}_{i1}, \tilde{T}_{i2}, \tilde{V}_{j1}, \tilde{V}_{j2} &\rightarrow V_{iq} (q \in \{1.2.3\}) \\
 R_v : \tilde{T}_{i1}, \tilde{T}_{i2}, \tilde{V}_{j1}, \tilde{V}_{j2} &\rightarrow V_{ip} (p \in \{1.2.3\}) \\
 r_{11i2,j1j2} &= \begin{cases} T_i \text{ для } R_i(*) \\ V_j \text{ для } R_v(*) \end{cases}
 \end{aligned} \tag{8}$$

Необхідно відзначити особливості процедур аналізу та прогнозування процесів НЧР на основі НМР (7). В запропонованих продукційних моделях і алгоритмах (7) – (8), отриманих на основі спостережень процесів, можливі «зупинки», якщо в первинних даних немає необхідних послідовностей зміни кроків (інтервалів/значень). При цьому тут передбачена необхідність використовувати відношення (7) різного рангу, розпочинаючи з другого. Якщо відношення другого рангу не дозволяють визначити очікувані величини наступного кроку процесу, використовують відношення першого або нульового рангу. У разі відсутності необхідних прогнозованих величин, виконується «зупинка» алгоритмічного процесу аналізу чи прогнозування.

При розрахунках контрольованих показників моделей НРМ, представлених нечіткими відношеннями (7), (8) для дискретних НВ таблиць використовується метод «центр ваги» (9)

$$B(S_j)^* = \frac{\sum B_k^* \mu_k}{\sum \mu_k} \quad (9)$$

Рівняння (9) записане для випадку, в якому для прикладу прогнозованим за (15) термом був (S_j) із показником приналежності « μ_j ». Моделі апроксимацій $S_1, S_2, S_3, S_4, S_5..$ – відомі, тому для всіх (S_k) $k = 1, 2, \dots$ при заданому « μ_j » розраховуються їх показники « μ_k » для рівня « α_j » = « μ_j », рис.2, а також відповідні їм значення характеристик (B_k) . На основі цих величин визначаються значення (9).

В результаті розробок, моделювання та досліджень було встановлено, що запропонований підхід до процесів при ЧРПНІ, за яким для моделей СПМ послідовно виконується процедури переходу від «ЧРПНІ довільного типу НВ» - «ЧРПНІ з трикутними моделями НЧР (апроксимація НЧР)» - «ЧР з перемінним кроком спостережень для системи α -рівнів» - «аналіз та прогнозування дійсних ЧР з перемінним кроком для сукупності α -рівнів» - «узагальнення результатів дослідження ЧР з перемінним кроком на α -рівнях методами скаляризації, «центр тяжіння», являється раціональним та результативним. Для реляційних моделей НРМ процесів ЧРПНІ використовуються процедури формування нечітких відношень (7), (8), відповідно них визначаються характеристики цих процесів з використанням (9). Загальна форма моделювання за РСМ процесу ЧРПНІ визначається комбінованою моделлю (4), що записана відповідно для показника $V(*)$ на основі результатів СПМ сепарабельної $V(Y)$ та НРМ, нечіткої реляційної моделі $V(\Delta)$.

На основі (7) було розроблене програмне забезпечення, загальна структура якого наведена на рис. 5. На рис. 5 блоки визначають наступне: Б1 – Комплекс функцій з управління даними у базі даних спостережень; Б1.1 – пошук моделей процесів у БД; Б1.2 – збереження, контроль, формування даних у БД; Б1.3 – видалення даних моделей процесів з БД; Б2 – функції формування даних вихідних процесів для дослідження; Б2.1 – уведення даних процесів для нечіткого моделювання, які представлені дійсними числами (H, T) , де – послідовність «Н» величини показника, упорядкованими за номерами спостережень, а «Т» - інтервали між спостереженнями; Б2.2 – відображення даних (H, T) ; Б2.3 – збереження моделі у БД; Б3 – Комплекс функцій моделювання процесів (H, T) нечіткими величинами, детермінованими структурами, відношеннями; Б3.1 – аналіз та розмноження даних (H, T) методами «бутстреп» [13];

Б3.2 – структурування окремих послідовностей «Н» та « $\Delta T(T)$ », визначення системи нечітких лінгвістичних змінних та їх термів для відображення послідовностей «Н» та « $\Delta T(T)$ » нечіткими трикутними НВ; Б3.3 – Дослідження властивостей послідовностей «Н» та « $\Delta T(T)$ »; Б3.3.1 – формування і дослідження результатів нечіткої апроксимації послідовності « $\Delta T(T)$ » ; Б3.3.2 – моделювання даних нечіткої апроксимації, вибір структури та оптимізація форм нечітких термів, представлення і відображення результатів нечіткої апроксимації послідовності « $\Delta T(T)$ » ; Б3.3.3 – формування реляційної моделі процесу для («Н» та « $\Delta T(T)$ » та ін.); Б4. Формування моделей для аналізу процесів («Н» та « $\Delta T(T)$ »), розрахунків та прогнозування результатів нечіткого моделювання; Б4.1 – моделювання нечітких апроксимацій процесів («Н» та « $\Delta T(T)$ ») ; Б4.1.1 – формування α -рівневих моделей нечітких апроксимацій процесів («Н» та « $\Delta T(T)$ »); Б4.1.2 – формування нерівномірних послідовностей моделей (Н, Т); Б4.2 – процедури прогнозування нечітких апроксимацій процесів («Н» та « $\Delta T(T)$ ») ; Б4.2.1 – формування результатів α -рівневих моделей нечітких апроксимацій процесів («Н» та « $\Delta T(T)$ ») ; Б4.2.2 – скаляризація результатів α -рівневих моделей; Б4.2.3 – розрахунок і представлення результатів прогнозування на встановлений період часу.

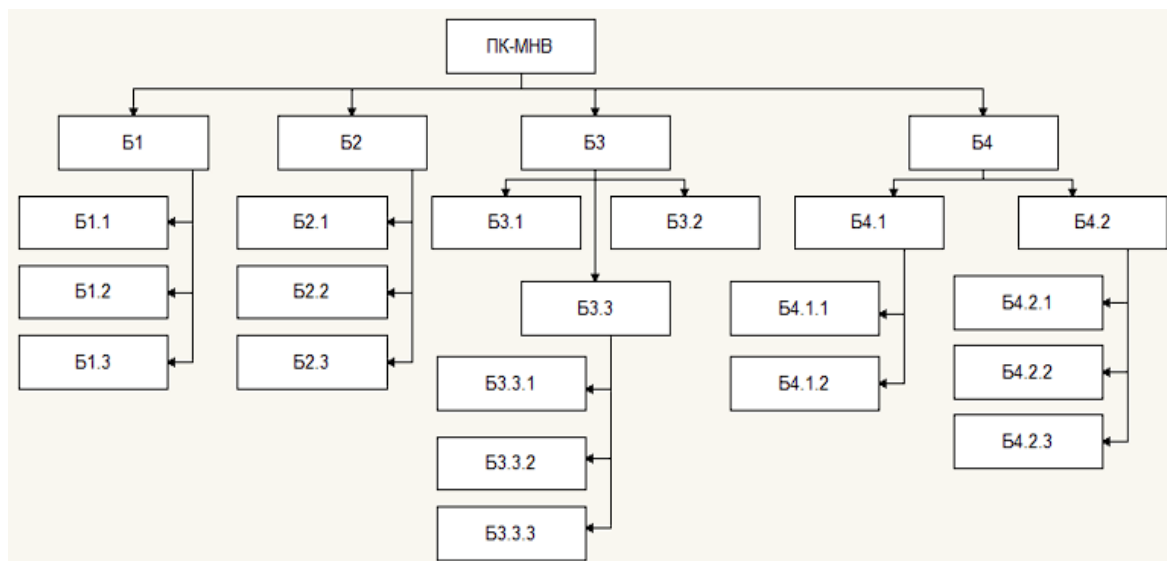


Рисунок 5 – Структура програмного комплексу з аналізу та моделювання процесів ЧРПНІ на основі РСМ

Моделювання процесів клінічного моніторингу при нечітких інтервалах (в термінах нечіткого кроку) [4] в цьому дослідженні було застосовано через скла-

дну структуру опису послідовностей інтервалів між контрольними спостереженнями процесів реабілітації.

Результати виконаних досліджень щодо моделювання ЧРПНІ далі були використані при формуванні програмного забезпечення щодо моделювання та прогнозування характеристик моніторингу процесів реабілітації хворих на діабет, з урахуванням їх особливостей. До головних відмінностей моніторингу процесів реабілітації хворих були віднесені наступні - урахування нечітких інтервалів між спостереженнями, а також індивідуальність моделей (ІМ) процесів для кожного клієнта. Застосування ІМ дозволяло у певному наближенні вирішувати головні завдання процесів реабілітації. А саме – перше, прогнозувати максимальні періоди (у дійсній або нечіткій формі) до стану/подій, що визначені певними вимогами; – друге, визначати можливість небезпечних оцінок показників моніторингу процесів реабілітації хворих.

Висновки. В статті для підвищення точності та ефективності моделювання і аналізу процесів, представлених часовими рядами з перемінними та нечіткими інтервалами, ЧРПНІ запропонована нова реляційно-сепарабельна модель (PCM) і комбінований квантильний алгоритм. Реляційна модель визначається системою нечітких реляційних відношень першого та другого порядку (НРМ), отриманих на основі вихідної послідовності даних. У комбінованому алгоритмі результати розрахунків, отримані за СПМ і моделями нечітких реляційних відношень, узагальнювалися при оптимальному виборі вагових коефіцієнтів для окремих складових. У дослідженнях шляхом числового моделювання було встановлено, що запровадження комбінованих PCM моделей процесів при ЧРПНІ являється раціональним та результативним. Приклади аналізу даних моніторингу процесів реабілітації хворих на діабет показали певні можливості забезпечення вимоги до точності результатів аналізу показників та їх короткострокового прогнозування. У статті встановлено певні можливості щодо підвищення точності та числової ефективності алгоритмів моделювання процесів моніторингу з перемінним та нечітким кроком спостережень, які були отримані на основі PCM моделювання, та при їх порівнянні з відомими розрахунками.

ЛІТЕРАТУРА

1. David J. Forecasting Data Published at Irregular Time Intervals Using an Extension of Holt's Method: Published by: INFORMS Stable URL: Accessed: 26-05-2015 UTC REFERENCES: http://www.jstor.org/stable/2631582?seq=1&cid=pdf-reference#references_tab_contents

2. Elorrieta F., Eyheramendy S., Palma W. Discrete-time autoregressive model for unequally spaced time-series observations. A&A 627, A120 (2019) <https://doi.org/10.1051/0004-6361/201935560>, ESO 2019 Astronomy & Astrophysics.
3. Jan Prüser. Adaptive learning from model space. DOI: 10.1002/for.2549.
4. Скалозуб В.В., Мурашов О.В. (2021) Моделювання даних процесів моніторингу при нерівномірних і нечітких інтервалах спостережень. «Системні технології». Випуск 4 (135). 135-148.
5. Q. Song, and B. S. Chissom, “Forecasting enrollments with fuzzy time series — Part I,” Fuzzy Sets and Systems, vol. 54, issue 1, 1993a, pp. 1–9.
6. Piegat A. Fuzzy modelling and control [Text] / Pfysica-Verlag, Heidelberg, 2001. – 698 pp.
7. Koenker R. “Quantile Regression”, Cambridge University Press, NY- 2005. pp. 137 – 143.
8. Tahseen A., Aqil S., Burney Cemal A. A New Quantile Based Fuzzy Time Series Forecasting Model [Електронний ресурс] – Режим доступу: <https://publications.waset.org/14214/pdf>
9. E. Bas, U. Yolcu and E. Egrioglu, “Intuitionistic fuzzy time series functions approach for time series forecasting”, Granular Computing, 2020.
10. Pal S.S., Kar S. “Fuzzy Time Series Model for Unequal Interval Length Using Genetic Algorithm”. Advances in Intelligent Systems and Computing, vol. 699 2019.
11. Bernal J.L., Cummins S., Gasparrini A. “Interrupted time series regression for the evaluation of public health interventions: a tutorial”. International Journal of Epidemiology, vol. 46, Issue 1, 2016 pp. 348–355.
12. Carla S. M`oller-Levet, F. Klawonn, Kwang-Hyun Cho and O. Wolkenhauer, “Fuzzy Clustering of Short Time-Series and Unevenly Distributed Sampling Points”, Advances in Intelligent Data Analysis V, 2003 pp. 330–340.
13. Геєць В.М. Моделі і методи соціально-економічного прогнозування /Геєць В.М., Клебанова Т.С., Черняк О.І. – Харків: ВД «ІНЖЕК», 205. – 396 с.
14. Скалозуб В.В. Методи інтелектуального моделювання процесів з перемінним інтервалом спостережень та конструктивного упорядкування «з вагою» / Скалозуб В.В., Білий Б.Б., Галабут О.О., Мурашов О.В.. // Системні технології. – 2020. – Випуск 5 (132). – С. 83-98.

REFERENCES

1. David J. Forecasting Data Published at Irregular Time Intervals Using an Extension of Holt's Method: Published by: INFORMS Stable URL: Accessed: 26-05-

2015 UTC REFERENCES: http://www.jstor.org/stable/2631582?seq=1&cid=pdf-reference#references_tab_contents

2. Elorrieta F., Eyheramendy S., Palma W. Discrete-time autoregressive model for unequally spaced time-series observations. A&A 627, A120 (2019) <https://doi.org/10.1051/0004-6361/201935560>, ESO 2019 Astronomy & Astrophysics.
3. Jan Prüser. Adaptive learning from model space. DOI: 10.1002/for.2549.
4. Skalozub V.V., Murashov O.V. (2021) Modeling of data monitoring processes with uneven and unclear observation intervals. "System technologies". Issue 4 (135). 135-148.
5. Q. Song, and B. S. Chissom, "Forecasting enrollments with fuzzy time series — Part I," Fuzzy Sets and Systems, vol. 54, issue 1, 1993a, pp. 1-9.
6. Piegat A. Fuzzy modelling and control [Text] / Pfysica-Verlag, Heidelberg, 2001. – 698 pp.
7. Koenker R. "Quantile Regression", Cambridge University Press, NY- 2005. pp. 137 – 143.
8. Tahseen A., Aqil S., Burney Cemal A. A New Quantile Based Fuzzy Time Series Forecasting Model. <https://publications.waset.org/14214/pdf>
9. E. Bas, U. Yolcu and E. Egrioglu, "Intuitionistic fuzzy time series functions approach for time series forecasting", Granular Computing, 2020.
10. Pal S.S., Kar S. "Fuzzy Time Series Model for Unequal Interval Length Using Genetic Algorithm". Advances in Intelligent Systems and Computing, vol. 699 2019.
11. Bernal J.L., Cummins S., Gasparrini A. "Interrupted time series regression for the evaluation of public health interventions: a tutorial". International Journal of Epidemiology, vol. 46, Issue 1, 2016 pp. 348–355.
12. Carla S. M'oller-Levet, F. Klawonn, Kwang-Hyun Cho and O. Wolkenhauer, "Fuzzy Clustering of Short Time-Series and Unevenly Distributed Sampling Points", Advances in Intelligent Data Analysis V, 2003 pp. 330–340.
13. Models and methods of socio-economic forecasting / Geets V.M., Kleba-nova T.S., Chernyak O.I. - Kharkiv: VD "INZHEK", 205. - 396 p.
14. Skalozub V.V. Methods of intellectual modeling of processes with a variable interval of observations and constructive ordering "with weight" / Skalozub V.V., White B.B., Galabut O.O., Murashov O.V. // System technologies. – 2020. – Issue 5 (132). - P. 83-98.

Received 17.04.2023.
Accepted 19.04.2023.

***Relational-separable models of monitoring processes at variable
and unclear observation intervals***

The article is devoted to the development of combined models, methods and tools designed to solve the current problems of modeling and analysis of monitoring process data, which are represented by time series and differ in variable or fuzzy observation intervals (CHRPNI). In the article, a new relational separable model (RSM) and a combined quantile algorithm are proposed to increase the accuracy and efficiency of modeling and analysis of the processes of CHRPNI. The relational model is defined by a system of fuzzy relational relations of the first and second order obtained on the basis of the original sequence of data. In the combined algorithm, the results of calculations obtained by SPM and models of fuzzy relational relationships were generalized with the optimal selection of weighting factors for individual components.

As a result of the conducted research by means of numerical modeling, it was established that the introduction of combined process models in the case of PNEU is rational and effective. Examples of data analysis of monitoring processes of rehabilitation of diabetic patients showed certain possibilities of ensuring the accuracy of the results of the analysis of indicators and their short-term forecasting.

Keywords: monitoring processes, uneven and fuzzy sampling interval, separable model, fuzzy relational relations, relational-separable model, combined quantile algorithm, monitoring of patients' condition.

Скалозуб Владислав Васильович - професор, каф. «Комп'ютерні інформаційні технології», Український державний університет науки і технологій, УДУНТ.

Горячкін Вадим Миколайович – зав. каф. «Комп'ютерні інформаційні технології», Український державний університет науки і технологій, УДУНТ.

Мурашов Олег Вячеславович - аспірант, каф. «Комп'ютерні інформаційні технології», Український державний університет науки і технологій, УДУНТ.

Skalozub Vladyslav - professor, dep. “Computer and Information Technology”, Ukrainian State University of Science and Technology, USUST.

Horiachkin Vadim – Head of dep. “Computer and Information Technology”, Ukrainian State University of Science and Technology, USUST.

Murashov Oleg - post-graduate student, Dep. “Computer and Information Technology”, Ukrainian State University of Science and Technology, USUST.

ОГЛЯД МАТЕМАТИЧНИХ МОДЕЛЕЙ ТА ІНФОРМАЦІЙНИХ ТЕХНОЛОГІЙ БІЗНЕС АНАЛІЗУ ВЕЛИКИХ WEB-ДАНИХ

Анотація. У сучасному бізнесі, особливо у сфері інтернет технологій, існує величезна кількість даних, які постійно надходять та накопичуються. Для ухвалення ефективних рішень необхідно проводити аналіз цих даних. Однак, обробка великих обсягів даних потребує спеціальних математичних моделей та інформаційних технологій. У зв'язку з цим дослідження математичних моделей та інформаційних технологій для аналізу великих web даних є актуальною темою.

Проблема полягає в тому, що для аналізу великих обсягів даних, отриманих зі сфери інтернет технологій, потрібні ефективні методи обробки та аналізу. Існуючі методи аналізу не завжди дозволяють отримати актуальну та точну інформацію, яка може бути використана для прийняття рішень.

Метою дослідження є огляд та аналіз існуючих математичних моделей та інформаційних технологій, що використовуються для аналізу великих web даних. Для досягнення цієї мети були використані методи аналізу літератури, порівняльний аналіз методів та засобів аналізу даних.

В результаті дослідження були виявлені основні математичні моделі та інформаційні технології, які широко використовуються для аналізу великих web даних. Було проведено аналіз та порівняння існуючих методів, виявлено їх переваги та недоліки.

Аналіз даних є важливим інструментом прийняття ефективних рішень у сфері інтернет технологій. Використання ефективних математичних моделей та інформаційних технологій дозволяє отримати точну та актуальну інформацію з великих web даних. Результати дослідження можуть бути використані для розробки нових методів та засобів аналізу даних, що дозволить покращити якість прийнятих рішень.

Ключові слова: web, бізнес аналіз, нейронні мережі, кластеризація, регресійний аналіз, NLP, машинне навчання, e commerce.

В даний час великі веб-дані є основним джерелом інформації для бізнесу. Ці дані містять у собі величезну кількість інформації про поведінку споживачів, тенденції в галузях, маркетингових кампаніях та багато іншого. Але для того, щоб використовувати ці дані в бізнес-аналізі, потрібні математичні моделі та

інформаційні технології. У цій статті буде розглянуто огляд математичних моделей та інформаційних технологій бізнес-аналізу великих web-даних.

У сучасному світі величезна кількість даних генерується та зберігається в інтернеті. Великі компанії, у тому числі й ті, що займаються веб-дизайном, маркетингом та рекламою, використовують ці дані для ухвалення бізнес-рішень. Однак, обсяги цих даних стають настільки великими, що їх аналіз та обробка стають трудомісткими та дорогими. Крім того, деякі параметри, наприклад, відгуки клієнтів можуть бути неструктурованими, що додатково збільшує складність їх аналізу.

Аналіз останніх досліджень і публікацій. В останні роки було проведено безліч досліджень у галузі аналізу великих веб-даних. Деякі були присвячені аналізу тексту, наприклад, аналізу тональності тексту і класифікації тексту. У роботі "A Comparative Study on Sentiment Analysis Techniques" [1] було проведено огляд різних методів аналізу тональності тексту та проведено їх порівняння з урахуванням метрик якості. У роботі "A review of text mining techniques and applications" [2] були розглянуті різні методи та технології, що використовуються для аналізу тексту, включаючи тематичне моделювання, кластеризацію, аналіз настроїв та інші.

Інші дослідження були присвячені аналізу зображень та відео, наприклад, розпізнаванню об'єктів та класифікації зображень. У роботі "Image Classification with Deep Convolutional Neural Networks" [3] була проведена класифікація зображень з використанням глибоких нейронних мереж. У роботі "Video classification with deep convolutional neural networks" [4] була представлена модель глибокої нейронної мережі для класифікації відео.

Крім того, були проведені дослідження в галузі аналізу соціальних мереж та аналізу мережових структур. У роботі "Social media analytics: a survey of techniques, tools and platforms" [5] було проведено огляд різних методів та технологій, що використовуються для аналізу соціальних мереж. У роботі "Graph Analysis: An Overview" [6] було проведено огляд методів та технологій, що використовуються для аналізу графових структур.

Також існують дослідження безпосередньо аналізу великих даних:

У роботі "Big Data Analytics in E-commerce: A Review" [7] автори оглядово розглядають різні методи аналізу великих даних в e-commerce, включаючи аналіз асоціативних правил, кластеризацію, класифікацію, рекомендаційні системи тощо. Також автори відзначають важливість використання глибокого навчання та нейронних мереж для обробки та аналізу великих даних.

У роботі "Big Data Analytics in Supply Chain Management: A Review" [8] автори розглядають застосування аналітичних методів управління ланцюгами поставок. Автори відзначають, що в даний час більшість компаній мають величезну кількість даних, і використання методів аналізу великих даних може допомогти оптимізувати та покращити ефективність управління ланцюгами постачання.

У роботі "Data Mining for Business Applications" [9] автори представляють огляд різних методів та алгоритмів для аналізу даних у бізнесі. Зокрема, автори розглядають методи кластеризації, класифікації, регресії, пошуку асоціативних правил тощо. Крім того, автори наголошують на важливості правильного вибору методу аналізу даних, враховуючи специфіку конкретного бізнесу та цілей аналізу.

В цілому, всі ці дослідження та публікації показують, що існує безліч методів та технологій, які можуть бути використані для аналізу великих веб-даних. Кожен метод та технологія мають свої переваги та недоліки, і вибір найбільш відповідного методу залежить від конкретної задачі аналізу даних.

Метою даного дослідження є огляд математичних моделей та технологій бізнес-аналізу великих web-даних. В рамках даного огляду буде розглянуто декілька методів та технологій, таких як аналіз асоціативних правил, кластеризація, класифікація, регресія та нейронні мережі. Також буде проведено порівняння ефективності даних методів та технологій на прикладі реальних даних.

Викладення основного матеріалу дослідження. Машинне навчання - це область штучного інтелекту, де комп'ютери навчаються з урахуванням даних. Машинне навчання використовується для передбачення поведінки користувачів, виявлення патернів у даних та створення моделей машинного навчання для певних завдань. Одним із найпоширеніших методів машинного навчання є метод наївного Байєса.

Формула методу наївного Байєса:

$$P(A|B) = \frac{\{P(B|A) * P(A)\}}{\{P(B)\}}$$

Метод наївного Байєса використовується для класифікації даних. Він заснований на теоремі Байєса, яка говорить про те, як ймовірність того, що подія A станеться за наявності події B залежить від ймовірності події B за умови, що подія A відбудеться. Метод наївного Байєса дозволяє створювати моделі ма-

шинного навчання, що використовуються для класифікації текстів, фотографій, звукових файлів та інших типів даних.

Кластерний аналіз - це метод, який використовується для виявлення угруповань даних на основі їхньої подібності. Кластерний аналіз може використовуватися для виявлення груп користувачів, схожих з поведінці чи інтересам.

Одним із найпоширеніших методів кластерного аналізу є метод k-середніх.

Формула методу k-середніх:

$$V = \sum_{i=1}^k \sum_{x \in S_i} (x - \mu_i)^2$$

Метод k-середніх полягає в розбитті даних на k груп, де k - це задане число кількості груп, а потім обчисленні центру мас для кожної групи. Дані потім перерозподіляються, щоб мінімізувати відстань між кожним елементом та центром мас групи. Цей процес повторюється, поки дані будуть розбиті на k груп.

Кластерний аналіз може бути використаний для аналізу поведінки користувачів на сайті, виявлення схожих груп клієнтів або виявлення схожих товарів чи послуг.

Регресійний аналіз - це метод, який використовується для визначення зв'язку між двома чи більше змінними. Регресійний аналіз може використовуватися для передбачення тенденцій даних і для визначення, які змінні можуть бути пов'язані з певним результатом.

Одним із найпоширеніших методів регресійного аналізу є метод лінійної регресії.

Формула методу лінійної регресії:

$$y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_n X_n + \epsilon$$

Метод лінійної регресії полягає у визначенні лінійного зв'язку між залежною змінною Y і однією або декількома незалежними змінними X. Коефіцієнти використовуються для визначення сили зв'язку між змінними, а помилка ϵ використовується для вимірювання точності моделі.

Регресійний аналіз може бути використаний для передбачення продажів, визначення факторів, що впливають на утримання клієнтів, і визначення того, які маркетингові кампанії найбільш ефективні.

В таблиці 1 наведені результати порівняльного аналізу основних математичних методів аналізу даних.

Порівняння математичних методів та їх ефективності для аналізу даних

Метод	Опис	Переваги	Недоліки
Машинне навчання	Навчання комп'ютерів на основі даних	Висока точність передбачень, здатність обробляти великі обсяги даних	Вимагає великої кількості даних для навчання
Кластерний аналіз	Метод групування даних на основі подібності	Дозволяє виявляти приховані зв'язки між даними, може використовуватися для створення персоналізованих рекомендацій	Вимагає вибору кількості груп, може бути неефективним при високих обсягах даних
Регресійний аналіз	Визначення зв'язку між змінними	Дозволяє прогнозувати результати на основі даних, визначати вплив факторів на результат	Потребує точних даних, не враховує можливих невідомих факторів

Додаткові методи аналізу даних, які можуть бути використані для аналізу великих web-даних (таблиця 2), включають:

1. Аналіз часових рядів: Цей метод використовується для аналізу даних, що змінюються з часом, таких як дані про продаж, відвідування сайту і т.д. Він може допомогти виявити тренди та сезонні зміни, а також передбачити майбутні значення. Для цього методу можуть бути використані такі інструменти, як ARIMA (авторегресійна інтегрована ковзна модель) та Prophet.

2. Аналіз текстів: Цей метод використовується для аналізу текстової інформації, як-от відгуки клієнтів, коментарі на соціальних мережах і т.д. Він може допомогти виявити теми, настрої та думки. Для цього можуть бути використані такі інструменти, як Natural Language Processing (NLP) і Sentiment Analysis.

3. Аналіз соціальних мереж: Цей метод використовується для аналізу соціальних мереж, таких як Twitter, Facebook і т.д. Він може допомогти виявити тенденції, думки та взаємозв'язки між користувачами. Для цього можуть бути використані такі інструменти, як Social Network Analysis і Graph Theory.

Методи аналізу web-даних

Метод аналізу	Опис	Переваги	Недоліки
Аналіз часових рядів	Аналіз даних, що змінюються з часом	Допомагає виявити тренди та сезонні зміни, передбачити майбутні значення	Вимагає точних даних, може бути неефективним при великих обсягах даних
Аналіз текстів	Аналіз текстової інформації, такий як відгуки клієнтів	Допомагає виявити теми, настрої та думки	Точність аналізу може бути обмежена через різні варіанти вираження сенсу
Аналіз соціальних мереж	Аналіз соціальних мереж, таких як Twitter, Facebook	Допомагає виявити тенденції, думки та взаємозв'язки між користувачами	Не всі дані доступні для аналізу, може бути неефективним при великих обсягах даних

Висновки. У цій статті було проведено огляд математичних моделей та технологій, що використовуються для аналізу великих web-даних. Машинне навчання, кластерний аналіз та регресійний аналіз були розглянуті як основні методи.

Машинне навчання є гнучким та ефективним методом, який може використовуватися для створення різних моделей, у тому числі для класифікації та кластеризації даних. Кластерний аналіз дозволяє виявляти приховані зв'язки між даними, але може бути неефективним при великих обсягах даних. Регресійний аналіз дозволяє прогнозувати результати на основі даних, але потребує точних даних і не враховує можливих невідомих факторів.

Залежно від цілей бізнес-аналізу та доступних даних, вибір конкретного методу може змінюватись. Важливо розуміти, що кожен метод має свої переваги та недоліки, і правильний вибір методу може суттєво вплинути на результати аналізу.

Інформаційні технології, такі як Hadoop та Apache Spark, можуть бути використані для обробки великих обсягів даних та прискорення обчислювальних процесів. Крім того, інструменти для візуалізації даних, такі як Tableau та Power BI, можуть допомогти з легким розумінням результатів аналізу та прийняттям ефективних бізнес-рішень.

Загалом аналіз великих web-даних є важливим компонентом сучасної бізнес-стратегії. Правильний вибір математичних моделей та інформаційних технологій може допомогти компаніям приймати ефективні бізнес-рішення, оптимізувати процеси та збільшити прибуток.

ЛІТЕРАТУРА / REFERENCES

1. Goyal, P. & Kumar, V. (2017). A Comparative Study on Sentiment Analysis Techniques. In Proceedings of the 2nd International Conference on Telecommunications and Networks (TEL-NET) (pp. 73-76). IEEE.
2. Islam, M. R., Islam, M. M., Islam, M. N., Islam, M. R., & Rahman, M. A. (2020). A review of text mining techniques and applications. Journal of King Saud University-Computer and Information Sciences, 32 (3), 257-269.
3. Krizhevsky, A., Sutskever, I., & Hinton, G. E. (2012). Image Classification with Deep Convolutional Neural Networks. In Advances in Neural Information Processing Systems (pp. 1097-1105).
4. Feichtenhofer, C., Pinz, A., & Zisserman, A. (2015). Video classification with deep convolutional neural networks. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (pp. 4475-4483).
5. Salehan, M., & Kim, D. J. (2015). Social media analytics: A survey of techniques, tools and platforms. Telematics and Informatics, 32 (4), 575-591.
6. Al Hasan, M., Chaoji, V., Salem, S., & Zaki, M. (2016). Graph Analysis: An Overview. Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery, 6 (1), 1-15.
7. Choudhury, P., & Shankar, R. (2018). Big Data Analytics in E-commerce: A Review. Journal of King Saud University-Computer and Information Sciences, 30 (4), 431-448.
8. Mishra, D., & Gunasekaran, A. (2018). Big Data Analytics in Supply Chain Management: A Review. Journal of Business Research, 83, 36-46.
9. Fayyad, U., Piatetsky-Shapiro, G., & Smyth, P. (1996). Data Mining for Business Applications. AI Magazine, 17 (3), 37-55.

Received 24.04.2023.

Accepted 27.04.2023.

Review of mathematical models and information technologies for business analysis of the big web data

The article provides a comprehensive review of mathematical models and information technologies used for analyzing large amounts of data in web applications. The latest research and publications in the field are analyzed, including a comparative analysis of machine learning methods, text, image, video analysis, social network analysis, and graph algorithms. The goal of this research is to analyze the effectiveness and applicability of mathematical models and information technologies in business analysis of large

web data. The article presents the results of the research and a comparative analysis of the efficiency of methods, which will help business analysts choose the optimal tools for processing and analyzing large amounts of data in web applications.

The article begins with an overview of the problem and the latest research and publications in the field. The article provides a detailed description of various mathematical models and information technologies, including their strengths and weaknesses. A comparative analysis of these methods is presented, with a focus on their effectiveness and applicability in business analysis.

The article also provides a detailed description of the applications of mathematical models and information technologies in various industries, such as e-commerce and supply chain management. The article analyzes the challenges and opportunities associated with the use of these technologies in business analysis and provides recommendations for businesses that want to take advantage of these technologies.

Overall, the article provides a comprehensive overview of mathematical models and information technologies used in business analysis of large web data. The article is a valuable resource for business analysts, data scientists, and researchers who want to learn more about the latest developments in this field.

Keywords: mathematical models, information technologies, business analysis, web data, machine learning, social network analysis, graph algorithms, big data analytics

Малієнко Станіслав Євгенович – аспірант кафедри інформаційних технологій і систем, факультет прикладних комп'ютерних технологій, Український державний університет науки і технологій.

Селівьорова Тетяна Віталіївна – кандидат технічних наук, доцент, доцент кафедри інформаційних технологій і систем, факультет прикладних комп'ютерних технологій, Український державний університет науки і технологій.

Maliienko Stanislav – postgraduate student of the Department of Information Technologies and Systems, Faculty of Applied Computer Technologies at Ukrainian State University of Science and Technologies.

Selivorstova Tatyana – candidate of technical sciences, associate professor, associate professor department of Information Technologies and Systems, Faculty of Applied Computer Technologies at Ukrainian State University of Science and Technologies.

РОЗРОБКА ПРОГРАМНОГО МОДУЛЮ ДЛЯ ІДЕНТИФІКАЦІЇ ЕМОЦІЙНОГО СТАНУ КОРИСТУВАЧА

Анотація. Дослідження ідентифікації емоцій у текстовому спілкуванні є актуальним на прямом досліджень в галузі обробки природної мови та машинного навчання. Основна мета роботи полягає в розробці програмного модулю, який реалізує алгоритми та моделі автоматичної ідентифікації емоцій людини у текстових повідомленнях. В роботі емоційні слова знаходилися за допомогою аналізу семантики речення та розглянуто два алгоритми для визначення емоцій.

Ключові слова: ідентифікація емоцій, дерево дискурсу, семантична модель, булевий алгоритм, векторний алгоритм, програмний модуль.

Величезна кількість сфер діяльності людини зумовлює появу інформаційних ресурсів, що відображають соціальні спілкування. Повідомлення користувачів, коментаторів стрічок тощо. Їх висловлювання містять інформацію про особисті стосунки до того, що відбувається у суспільному житті. Внаслідок чого виникає завдання обробки інформації з метою визначити та проаналізувати їх спектр дій [1].

Дослідження ідентифікації емоцій у текстовому спілкуванні є актуальним напрямком досліджень в галузі обробки природної мови та машинного навчання. Основна мета роботи полягає в розробці програмного модулю, який реалізує алгоритми та моделі, які можуть автоматично визначати емоційний стан людини на основі текстових повідомлень. Дана робота присвячена огляду деяких моделей та алгоритму для поліпшення обробки даних в середині текстового спілкування користувачів [2].

Основною метою роботи є створення програмного модулю для спільного аналізу емоції та ідентифікація емоційного стану користувача. За допомогою даного модулю, який є елементом соціальної мережі, користувачі мають змогу виконати пошук емоцій при роботі з текстами у ній. За допомогою реалізованих алгоритмів семантичної фільтрації та індексації, програма

аналізує наявні текстові данні та визначає ступінь їх приналежності до заданих тем.

Один з методів, який використовується у роботі – метод фільтрації. Метод фільтрації визначає дискусії тексту, які записує у вигляді ієрархічної деревоподібної структури. Також будує семантичну модель, дані які отримані з текстового спілкування користувачів. За допомогою описаних структур метод фільтрації знаходить емоційні слова, які записані в базі. Пошук відбувається за рахунок ключових слів. У свою чергу ключові слова визначаються відмінком.

Відмінок - граматична словозмінна категорія, характерна для іменних частин мови (іменника, прикметника, числівника, займенника). Відмінок як категорія також притаманний дієприкметнику (особлива форма дієслова, яка поєднує ознаки і дієслова, і прикметника).

Закінчення іменника залежить не тільки від того, в якому відмінку та числі він стоїть, але і від його скасування та групи (дивись рисунок 1), а іноді і від значення. Тому в іменниках української мови багато різних закінчень.

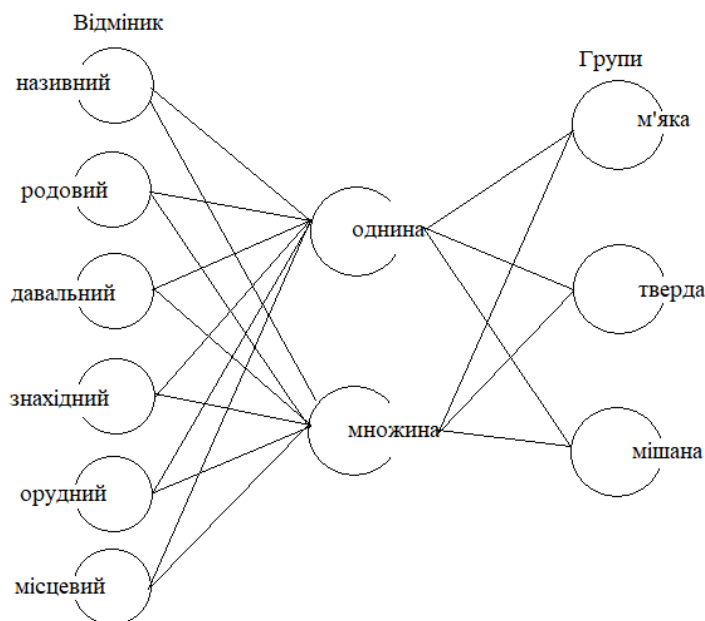


Рисунок 1- Відмінки та їх групи

Позначення відмінків у словах, значно зменшує навантаження на базу даних, у якій зберігаються слова для фільтрації. Слова в базі даних зберігаються в однині та твердій групі.

Дискурс - завдання побудови ієрархічної деревоподібної структури над пропозицією, листям якої є пропозиції. Дерево дискурсу представляється як набір складових $R[i; j]$, де $i < j$. У цьому поданні R відноситься до риторично-

го відношенню, яке має місце між одиницею дискурсу, від i до j , та одиниця, що містить від $i + 1$ до j .

Розглянемо на прикладі: «Я ходив сьогодні в кіно. Мені дуже сподобалося – раджу».

Дискурс описується переважно як сукупність текстів, присвячених конкретній темі. Тема цього обговорення - кіно.

У нашому прикладі: хто (я), я що робив (ходив), ходив куди (у кіно).

Побудуємо ієрархічно деревоподібну структуру над пропозицією (рисунок 2).

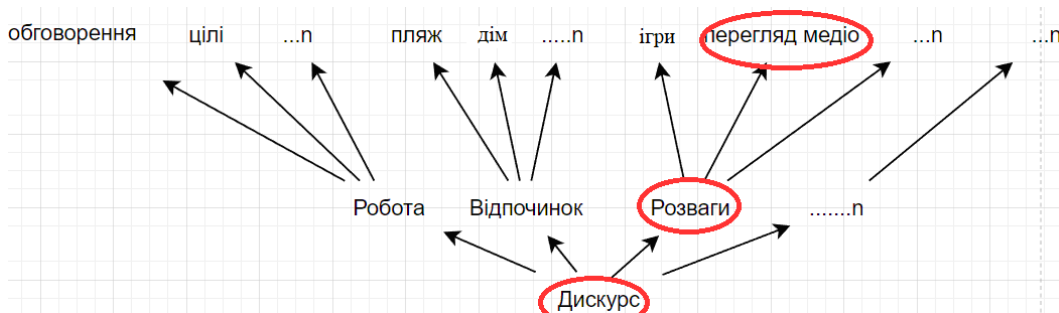


Рисунок 2 - Ієрархічна деревоподібна структура

Дискурс значно спростить і точніше визначить емоцію в тексті.

Також в роботі використано семантичну модель представлення даних. Семантична модель може бути проілюстрована графічно за допомогою ієрархічної діаграми абстракцій. Це робиться ієрархічно, щоб типи, на які посилаються інші типи, завжди були перелічені вище того, на що вони посилаються. Це полегшує читання та розуміння. Переглянемо на нашому прикладі: «Я ходив сьогодні в кіно. Мені дуже сподобалося – раджу».

Абстракції, що використовуються в семантичній моделі даних прикладу представлено на рисунку 3.

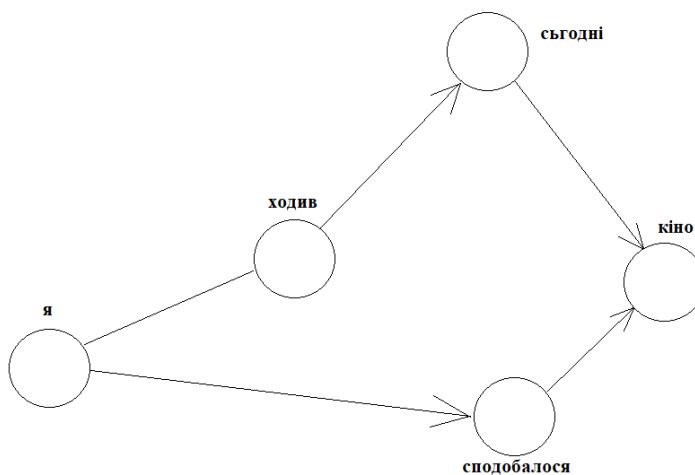


Рисунок 3 - Семантична модель даних прикладу

За цією моделлю визначаємо, що «я ходив сьогодні в кіно» і «сподобалося кіно». Семантична модель значно спростить і точніше допоможе визначити емоцію.

Алгоритм пошуку — алгоритм, який вирішує задачу пошуку, тобто, знаходить інформацію, яка зберігається в певній структурі даних. Структури даних можуть бути реалізовані за допомогою зв'язаних списків, масивів, дерев пошуку, хеш-таблиць чи інших методів зберігання інформації. Алгоритм пошуку на пряму залежить від структури даних, для якої він реалізований.

Найпростішим алгоритмом інформаційного пошуку є Булев пошук. У ньому запити будуються на основі елементарних термінів документів (слів або словоформ), пов'язаних між собою логічними операціями. Знайдені дані визначаються внаслідок описаних запитом логічних операцій над безліччю пошукових образів документів. Користувач отримує лише дані, чиї набори термінів точно збігаються з відповідними комбінаціями термінів запиту.

Для оцінки роботи системи та її окремих елементів використовувалася база листування співрозмовників в обсязі 100 повідомлень.

Середній показник точності знаходження емоційних слів булевим алгоритмом на даній базі повідомлень склав 76%.

Також розглядався векторний алгоритм пошуку.

Багато відомих інформаційно-пошукових систем базуються на векторному алгоритмі пошуку. Цей алгоритм є класичним алгебраїчним. Він заснований на векторній моделі інформаційного масиву, в якій для визначення міри близькості результатів і запитів використовується значення косинуса кута між векторами в багатовимірному просторі інформаційного масиву. Опис запиту, відповідного необхідної користувачеві тематики, також є вектором у тому ж евклидовому просторі термів.

Середній показник точності знаходження емоційних слів векторним алгоритмом на даній базі повідомлень 74.6%.

Також для порівняння результатів та оцінки роботи алгоритмів в роботі проведено пошук емоційних слів експертом. Середня точність ідентифікації емоцій експертом дорівнює 92,5%.

Отже, навіть людина не завжди в змозі ідентифікувати емоційні слова в тексті у повному обсязі, тому в результатах роботи можна вважати роботу алгоритмів достатньо адекватною на даному етапі.

Оскільки текстова обробка додатку призначена для роботи із текстом в інтерактивному режимі, то вона містить також додаткову функціональність, покликану автоматизувати дії для визначення виразів та емоційних слів.

Висновки. В роботі розглянуто питання щодо пошуку емоцій у текстових повідомленнях та розробки програмного модулю для її реалізації. Розглянуто два алгоритми для визначення емоцій - векторний та бульовий. У ході дослідження визначилося, що бульовий алгоритм найбільше підходить для пошуку емоційних слів. В роботі емоційні слова знаходилися за допомогою виявлення та аналізу семантики речення.

Результати проведеної роботи показали, що у розробленому програмному модулі ефективність ідентифікації емоційного стану співрозмовника дорівнює 85.7%.

ЛІТЕРАТУРА

1. Дмитрієва І.С., Кириченко Я.С. Аналіз та класифікація токсичних повідомлень у текстових повідомленнях. The 8th International scientific and practical conference "Trends, theories and ways of improving science" (February 28 – March 03, 2023) Madrid, Spain. International Science Group. 2023. p.529-531
2. Гнезділова Я.В. Емоційність та емотивність сучасного англomовного дискурсу: структурний, семантичний і прагматичний аспекти: дис. ... канд. філолог. наук. – К.: КНЛУ, 2007 – 269 с.

REFERENCE

1. Dmytriieva I.S., Kyrychenko Ya.S. Analiz ta klasyfikatsiia toksychnykh povidomlen u tekstovykh povidomlenniakh. The 8th International scientific and practical conference "Trends, theories and ways of improving science" (February 28 – March 03, 2023) Madrid, Spain. International Science Group. 2023. p.529-531
2. Hnezdilova Ya.V. Emotsiinist ta emotyvnist suchasnoho anhlomovnoho dyskursu: strukturnyi, semantychnyi i prahmatychnyi aspekty: dys. ... kand. filoloh. nauk. – K.: KNLU, 2007 – 269 s.

Received 26.04.2023.
Accepted 28.04.2023.

Development of a software module for the identification of the emotional state of the user

A huge number of spheres of human activity leads to the emergence of information resources that reflect social communication.

The study of the identification of emotions in text communication is an actual direction of research in the field of natural language processing and machine learning. The main goal of the work is to develop a software module that implements algorithms and

models that can automatically determine a person's emotional state based on text messages. This work is devoted to the review of some models and an algorithm for improving data processing in the middle of text communication of users.

One of the methods used in the work is the filtering method. The filtering method determines the discussions of the text, which it records in the form of a hierarchical tree-like structure. Discourse greatly simplifies the work and allows you to more accurately determine the emotion in the text.

It also builds a semantic model, the data of which is obtained from the text communication of users. Using the described structures, the filtering method finds emotional words recorded in the database. The search is based on keywords. In turn, keywords are defined by case.

The work deals with the issue of finding emotions in text messages and the development of a software module for its implementation. Two algorithms for determining emotions are considered - vector and Boolean. During the research, it was determined that the Boolean algorithm is most suitable for searching for emotional words. In the work, emotional words were found by identifying and analyzing the semantics of the sentence.

Keywords: emotion identification, discourse tree, semantic model, Boolean algorithm, vector algorithm, software module.

Дмитрієва Ірина Сергіївна - к.т.н., доцент, доцент кафедри ІТС Інститут промислових і бізнес технологій, Український державний університет науки і технологій, Україна.

Бімалов Дмитро Вікторович – магістр кафедри ІТС Інституту промислових і бізнес технологій, Український державний університет науки і технологій, Україна.

Dmytriieva Iryna Serhiivna - Ph.D., Associate Professor, Associate Professor of the ITS Department of the Institute of Industrial and Business Technologies, Ukrainian State University of Science and Technology, Ukraine.

Bimalov Dmytro Viktorovych - Department of the Institute of Industrial and Business Technologies, Ukrainian State University of Science and Technology, Ukraine.

В.У. Ігнаткін, В.С. Дудніков, Т.Р. Лучишин, С.В. Алексеєнко,
О.П. Юшкевич, Т.П. Карпова, Т.С. Хохлова, Ю.С. Хомош, В.А. Тіхонов
**АЛЬТЕРНАТИВА МЕТОДАМ СЕРЕДНІХ ТА НАЙМЕНШИХ КВАДРАТІВ,
ЯКІ ВИКОРИСТОВУЮТЬСЯ ПРИ ОБРОБЦІ РЕЗУЛЬТАТІВ
НАУКОВО-ТЕХНІЧНИХ ЕКСПЕРИМЕНТІВ**

Анотація. Збільшення складності і розмірів систем різної природи вимагає постійного удосконалення моделювання і перевірки отриманих результатів шляхом експерименту. Чітко провести кожний експеримент, об'єктивно оцінити відомості про досліджуваний процес і поширити матеріал, отриманий в одному дослідженні, на серію інших досліджень можна тільки при правильній їхній постановці й обробці. На основі експериментальних даних підбирають алгебраїчні вирази, які називають емпіричними формулами, що використовують якщо аналітичний вираз деякої функції складній, або не існує на даному етапі опису об'єкту, системи або явища. При підборі емпіричних формул широко використовують поліноми вигляду: $y = A_0 + A_1x + A_2x^2 + A_3x^3 + \dots + A_nx^n$, якими можна апроксимувати будь-які результати вимірів, якщо вони виражаються безперервними функціями. Особливо цінним є те, що навіть при невідомому точному виразі рішення (поліному) можна визначити значення коефіцієнтів A_n за допомогою методів середніх й найменших квадратів. Але у методі найменших квадратів спостерігається зсув оцінок при збільшенні шумів у знайдених даних так як сказується вплив шумів попередніх етапів обробки інформації. Тому для процедур обробки інформації у реальному масштабі часу пропонується операція псевдозвороту, яка виконується за допомогою рекурентних формул. Ця процедура є процедурою послідовного оновлення (зі зсувом) по стовбцям матриці заданих розмірів та і псевдозвороту на кожному кроці зміни інформації. Цей підхід є прямим та використовує переваги, властиві методу облямівки.

При псевдозвороті мається можливість контролювати правильність обчислень на кожному кроці, використовуючи умови Пенроуза. Необхідність псевдозвороту може виникнути при оптимізації, прогнозуванні тих чи інших параметрів та характеристик систем різного призначення, в різноманітних задачах лінійної алгебри, статистики, представленні структури одержаних рішень, зрозуміти зміст некоректності рішення, що виникає, в сенсі Адомара Тихонова і побачити шляхи регуляризації таких рішень.

Ключові слова: метод середніх і найменших квадратів; емпіричні формули; псевдозворот; експериментальні дані; апроксимуючі поліноми.

Постановка задачі. Узагальнити та обґрунтувати недоліки визначення коефіцієнтів апроксимуючих поліномів для будь-яких вимірів і безперервних функцій, зокрема для методів середніх і найменших квадратів, а також сформулювати альтернативу цим методам для реального масштабу часу, при цьому не повинна накопичуватися похибка від експерименту до експерименту.

Викладення основного матеріалу. При підборі емпіричних формул широко використовуються поліноми

$$y=A_0+A_1x+A_2x^2+A_3x^3+\dots+A_nx^n, \quad (1)$$

де $A_0, A_1, \dots, A_n$ — сталі коефіцієнти. Поліномами можна апроксимувати будь-які результати вимірів, якщо вони виражаються безперервними функціями. Особливо цінним є те, що навіть при невідомому точному виразі функції (1) можна визначити значення коефіцієнтів A . Для визначення коефіцієнтів застосовують методи середніх і найменших квадратів [1-5].

Метод середніх квадратів заснований на наступному положенні. По експериментальних точках можна побудувати декілька плавних кривих. Найкращою буде та крива, у якої різниці відхилення виявляються найменшими, тобто $\sum \varepsilon^2 \approx 0$. Порядок розрахунку коефіцієнтів полінома зводиться до наступного. Визначається число членів ряду (1), що звичайно приймають не більш 3...4. У прийняте вираження послідовно підставляють координати x і y декількох (m) експериментальних точок і одержують систему з m рівнянь. Кожне рівняння дорівнює відповідному відхиленню:

$$\begin{aligned} A_0 + A_1x_1 + A_2x_1^2 + \dots + A_nx_1^n - y_1 &= \varepsilon_1; \\ A_0 + A_1x_2 + A_2x_2^2 + \dots + A_nx_2^n - y_2 &= \varepsilon_2; \\ A_0 + A_1x_m + A_2x_m^2 + \dots + A_nx_m^n - y_m &= \varepsilon_m. \end{aligned} \quad (2)$$

Число точок, тобто число рівнянь, повинне бути не менше числа коефіцієнтів A , що дозволить їх обчислити шляхом рішення системи (2).

Розбивають систему початкових рівнянь (2) послідовно зверху вниз на групи, число яких повинно бути дорівнює кількості коефіцієнтів A . У кожній групі складають рівняння й одержують нову систему рівнянь, рівну кількості груп (звичайно 2...3). Вирішуючи систему, обчислюють коефіцієнти A .

Метод середніх квадратів має високу точність, якщо число крапок досить велике (не менш 3...4). Однак ступінь точності можна підвищити, якщо початкові умови згрупувати по 2...3 варіанта й обчислити для кожного варіанта емпіричну формулу. Перевага варто віддати тій формулі, у якій $\sum \varepsilon^2 = \min$. Нехай, наприклад, виконано сім вимірів:

4	5	6	7	8	9	10
10,2	6,7	4,8	3,6	2,7	2,1	1,7

Для підбора емпіричної формули можна вибрати поліном

$$y=A_0+A_1 x+A_2x^2$$

Шляхом підстановки в це рівняння значень вимірів систему початкових рівнянь можна розділити на три групи 1...2, 3, 4, 5, .. 7 у виді:

$$1. A_0+4A_1+16A_2-10,2= \varepsilon_1$$

$$2. A_0+5A_1+25A_2-6,7= \varepsilon_2;$$

$$3. A_0+6A_1+36A_2-4,8= \varepsilon_3;$$

$$4. A_0+7A_1+49A_2-3,6= \varepsilon_4;$$

$$5. A_0+8A_1+64A_2-2,7= \varepsilon_5;$$

$$6. A_0+9A_1+81A_2-2,1= \varepsilon_6;$$

$$7. A_0+10A_1+100A_2-1,7= \varepsilon_7;$$

Додавання рівнянь у кожній підгрупі дає

$$\text{I-я група } 2A_0+9A_1+4A_2= 16,9;$$

$$\text{II-я група } 2A_0+13A_1+85A_2=8,4;$$

$$\text{III-я група } 3A_0+27A_1+245A_2=6,5.$$

Визначення з цих виражень коефіцієнтів A_0 , A_1 і A_2 приводить до емпіричної формули $ax= 26,168 - 5,2168x + 0,2811 x^2$.

Метод середніх квадратів може бути застосований для різних кривих після їхнього вирівнювання. Нехай, наприклад, мається вісім вимірів:

3	6	9	12	15	18	21	24
57,6	41,9	31,0	22,7	16,6	12,2	8,9	6,5

Аналіз кривої в системі прямокутних координат дає можливість застосувати формулу:

$$y=ae^{-bx}.$$

Зробимо вирівнювання шляхом заміни перемінних $Y=\lg y$,

$X=x/2,303$. Тоді $Y=A+BX$, де $A=\lg a$, $Y = b$.

Тому що необхідно визначити два параметри, те усі виміри поділяються на двох груп по чотирьох вимірах. Це приводить до рівнянь:

$$1,7601 = A + \frac{3}{2,303} B ; 1,2201 = A + \frac{15}{2,303} B ;$$

$$1,6222 = A + \frac{6}{2,303} B ; 1,0864 = A + \frac{18}{2,303} B ;$$

$$1,4914 = A + \frac{9}{2,303} B ; 0,9494 = A + \frac{21}{2,303} B ;$$

$$1,3560 = A + \frac{12}{2,303} B ; 0,8129 = A + \frac{24}{2,303} B ;$$

$$6,2300 = 4A + \frac{30}{2,303} B ; 4,0688 = 4A + \frac{78}{2,303} B .$$

Після підсумовування по групах можна одержати систему двох рівнянь із двома невідомими А і В, рішення яких дає: $A = 1,8952$; $a = 78,56$; $Y = -0,1037$; $b = -0,1037$. Остаточна $B = 78,56 e^{-0,1037x}$

Гарні результати при визначенні параметрів рівняння дає використання методу найменших квадратів. Як видно суть цього методу полягає в тому, що якщо усі виміри функцій $y_1, y_2, \dots, y_n$ зроблені з однаковою точністю і розподілені величини помилок виміру відповідають нормальному закону, то параметри досліджуваного рівняння визначається з умови, при якому сума квадратів відхилень обмірюваних значень від розрахункових приймає найменше значення. Для існування невідомих параметрів ($a_1, a_2, \dots, a_n$) необхідно вирішити систему лінійних рівнянь:

$$\begin{aligned} y_1 &= a_1 x_1 + a_2 u_1 + \dots + a_n z_1 ; \\ y_2 &= a_1 x_2 + a_2 u_2 + \dots + a_n z_2 ; \\ y_n &= a_1 x_m + a_2 u_m + \dots + a_n z_m , \end{aligned} \quad (3)$$

де $y_1, \dots, y_n$ - приватні значення обмірюваних величин функції b ; x, u, z — змінні величини.

Цю систему приводять до системи лінійних рівнянь шляхом множення кожного рівняння відповідно на $x_1 \dots x_m$ і наступного їхнього додавання, потім множення відповідно на $i_1 \dots i_m$ рішення якої і дає шукані коефіцієнти.

$$\begin{aligned} \sum_{i=1}^m yx &= a_1 \sum_{i=1}^m xx + a_2 \sum_{i=1}^m xu + \dots + a_n \sum_{i=1}^m xz ; \\ \sum_{i=1}^m yu &= a_1 \sum_{i=1}^m ux + a_2 \sum_{i=1}^m uu + \dots + a_n \sum_{i=1}^m uz ; \\ \sum_{i=1}^m yz &= a_1 \sum_{i=1}^m zx + a_2 \sum_{i=1}^m zu + \dots + a_n \sum_{i=1}^m zz , \end{aligned} \quad (4)$$

Нехай, наприклад, необхідно визначити коефіцієнти a_1 і a_2 у рівнянні $k_p = a_1 + a_2 u$. Тому що потрібно визначити два параметри, то система рівнянь може бути представлена у виді двох рівнянь $y = a_1 x_1 + a_2 u_2$ і $y u_2 = a_2 x_1 u_2 + a_2 u_2^2$, де $v = k_p$; $x_1 = 1$; $x_2 = u$.

Тому що рівняння лінійні, можна обмежитися чотирма серіями дослідів. Якщо вони зведені в табл. 1, то систему нормальних рівнянь можна записати у виді $5,48 = 4a_1 + 1100a_2$; $1519 = 1100a_1 + 307350a_2$, рішення яких дає $a_1 = 0,78$; $a_2 = 0,0025$. Отже, емпірична формула одержить вид $k_p = 0,78 + 0,0025 u$.

Таблиця 1

Результати дослідів

$u_1=u$	$y=k_p$	u^2	yu
230	1,26	52900	289,8
255	1,32	65025	336,6
295	1,40	87025	413,0
320	1,50	102400	480,0
1100	5,48	307350	1519,4

Метод найменших квадратів забезпечує досить надійні результати. При цьому ступінь точності коефіцієнтів A у (1) повинна бути такою, щоб обчислені значення b збігалися зі значеннями у вихідних табличних значеннях. Це вимагає обчислювати значення A тем точніше, чим вище індекс A . тобто A_1 повинний бути точніше (більше число десяткових знаків), чим A_2 ; A_{23} — точніше, ніж A_2 , і т.д. Для обчислення коефіцієнтів A методом найменших квадратів необхідно користатися типовими програмами для ЕОМ.

Недоліком методу найменших квадратів є зміщення оцінок зі збільшенням шуму наступних вихідних даних, які надходять, що впливає на наступні етапи обробки інформації. Тому для процедур обробки інформації в реальному масштабі часу запропоновано операцію псевдозвороту. Алгоритм псевдозвороту реалізується з допомогою нижче поданих рекурентних формул, тобто, процедура послідовної обробки інформації фактично представляється як процедура послідовного оновлення (зі зсувом) по стовпчиках матриці заданих розмірів та її псевдозвороту на кожному кроці зміни інформації. Пропонований алгоритм (метод) є прямим і використовує переваги, властиві методу оздоблення. При псевдозвороті даним методом є можливість контролювати правильність обчислень на кожному кроці, використовуючи умови Пенроуза [7-11]. Цей метод особливо ефективний у процедурах обробки інформації у реальному масштабі часу і суттєво зменшується запізнення порівняно з методом найменших квадратів.

Алгоритм псевдозвороту матриць у завданнях обробки інформації при моделюванні та управлінні об'єктами різної природи. Необхідність псевдозвороту виникає у найрізноманітніших прикладних завданнях, наприклад, при вирішенні як самої проблеми автоматизації системи метрологічного обслуговування засобів вимірювань (ЗВ) промислового підприємства, так і в ході подальшого її функціонування. Така необхідність, наприклад, може виникнути під час оптимізації чи прогнозування тих чи інших параметрів чи характеристик системи метрологічного обслуговування ЗВ промислового підприємства. Незважаючи на те, що проблема псевдозвороту матриць була сформульована ще в 20-х р. минулого століття, інтерес до неї виник порівняно нещодавно [6]. І не випадково: виявилось, що застосування псевдозворотних матриць у найрізноманітніших задачах лінійної алгебри і статистики, що виникають в самих різних додатках, в тому числі при автоматизації метрологічних служб, можна отримати узагальнення класичних рішень, візуалізувати структуру отриманих рішень, зрозуміти сенс часто виникає некоректності рішення в сенсі Адамара-Тихонова, і побачити шляхи регуляризації таких рішень. література, в якій описується значна кількість методів псевдоінверсії реальної матриці довільних розмірів. Особливий інтерес викликають прикладні завдання, призначені для використання комп'ютерів, являють собою прямі псевдоінверсійні методи, такі як метод Гревеля. Ця програма описує метод прямого псевдорозвороту, заснований, на відміну від інших відомих методів, на використанні допоміжної матриці, для якої відомий її псевдозворот. Метод псевдорозвороту полягає в послідовному витісненні допоміжної матриці на скінченну кількість кроків, рівне числу стовпців або рядків (для визначеності в подальшому - стовпців).

При викладенні рекурентного алгоритму псевдорозвороту прийняти наступні позначення: 1) великою літерою позначено прізвище матриці, наприклад, матриці X, Z, а також вектори-рядки невідомих параметрів, відповідні елементи котрих формуються за визначеними законами; 2) малими літерами позначені відповідні стовпці або рядки матриць.

Рекурентний алгоритм має наступний вигляд:

$$A_i^+ = \varphi \begin{cases} S_i^+ - (\lambda - S_i^+ a_i) \lambda' S_i^+ I (1 - a_i' S_i^+ a_i), \text{ якщо } (\lambda - S_i^+ a_i \neq 0) \\ S_i^+ - S_i^+ S_i^+ S_i^+ a_i \lambda' S_i^+ I a_i' S_i^+ S_i^+ S_i^+ S_i^+ a_i - \text{в інших випадках} \end{cases} \quad (5)$$

$$A_i^+ = \varphi \begin{cases} S_i^+ - (\lambda - S_i^+ a_i) \lambda' S_i^+ I (1 - a_i' S_i^+ a_i), \text{ якщо } (\lambda - S_i^+ a_i \neq 0) \\ S_i^+ - S_i^+ S_i^+ S_i^+ a_i \lambda' S_i^+ I a_i' S_i^+ S_i^+ S_i^+ S_i^+ a_i - \text{в інших випадках} \end{cases} \quad (6)$$

$$b_{i+1} = \begin{cases} (I - A_i A_i^+) x_i I x_i' (I - A_i A_i^+) x_i, \text{ якщо } o(I - A_i A_i^+) x \neq 0 \\ A_i^+ A_i^+ x_i I (1 + x_i' A_i^+ A_i^+ x_i) - \text{в інших випадках;} \end{cases} \quad (7)$$

$$B_{i+1} = A_i^+ - A_i^+ x_i b_{i+1}' \quad i=0,1,2,\dots, m-1,$$

де A_i - матриця розмірів $[n \times (m-1)]$, отримано шляхом відкидання першого стовпця a_i в допоміжній матриці S_i ; $|^+$ — операція псевдозвороту; $|'$ — операція транспонування, ϕ — оператор відкидання верхнього нульового рядку; $\lambda'' - m''$ — розмірний вектор, перший елемент якого дорівнює одиниці, а інші елементи дорівнюють нулю; I — одинична матриця; b_{i+1}' - нижній рядок матриці S_{i+1}^+ розміру $(1 \times n)$; x_i - це стовпець "i" заданої матриці, яка підлягає псевдозвороту, x_i ; B_{i+1} - матриця розміру $[(m-1) \times n]$, яка складає спільно з рядком b_{i+1}' матрицю S_{i+1}^+ .

Використання рекурентних співвідношень дозволяє за m кроків замінити допоміжну матрицю S_0 розміру $(n \times m)$ на задану матрицю X однакового розміру і визначити її псевдозвортню. Процес починається з відкидання першого стовпця a_0 допоміжної матриці S_0 і доповнення її праворуч першим стовпцем x_1 матриці X , яка підлягає псевдозвороту. Застосування рекурентних формул (5-7) до матриці, яка засвоєна таким чином, дозволяє легко обчислити її псевдозвортню і для наступного кроку використовувати її вже як допоміжну з відомою псевдозвортною. Послідовне витіснення і використання рекурентних формул (5-7) дозволяє за m кроків повністю замінити початкову допоміжну матрицю S_0 на задану X і отримати її псевдозвортню..

Необхідність априорі мати допоміжну матрицю відповідного розміру з відомою псевдозвортною для реалізації описуваного методу псевдозвороту не є "обтяжливою". Зокрема, при необхідності псевдозвороту квадратної матриці X (виродженої, погано обумовленої, невивродженої або добре обумовленої невивродженої) в якості допоміжної матриці S_0 , можна використати одиничну матрицю, для якої псевдозвортна буде одиничною. При необхідності псевдозвороту прямокутної матриці X , в загальному випадку розміру $(n \times m)$, наступна методика може бути використана для формування допоміжної матриці S_0 розміром $(n \times m)$ і визначення псевдозвортної матриці. Потрібно взяти одиничну матрицю розміру $(n \times m)$ (для визначеності вважаємо, що $n > m$) і від неї відкинути праворуч або зліва $(n-m)$ стовпців. В результаті буде отримано прямокутну матрицю розміру $(n \times m)$, для якої, як легко переконатися, транспонована до неї матриця буде її псевдооборотною,

оскільки задовольнятиме всім умовам Пенроуза. Отже, отримувану таким прийомом матрицю можна використовувати як допоміжну в описаній в цьому додатку методі псевдозвороту.

Метод псевдозвороту за допомогою допоміжної матриці може бути використаний у випадках, коли математична модель досліджуваної системи або процесу, що потікає в ній, може бути приведена до лінійної моделі та представлена у вигляді наступного лінійного рівняння:

$$GX=Z, \quad (8)$$

де: G – вектор рядок невідомих параметрів розміру $(1 \times n)$; X і Z матриці розмірів $(n \times m)$ і $(1 \times m)$, відповідні елементи яких формуються за певними законами, що визначаються типом і сутністю конкретної задачі. Наприклад, під час проведення оптимізації параметрів системи метрологічного обслуговування складних засобів вимірювальної техніки (ЗВТ) промислового підприємства виникає необхідність розв'язання рівняння типу (8)

Однак слід мати на увазі, що описаний метод псевдозвороту за допомогою допоміжної матриці найбільш ефективний в порівнянні з іншими методами псевдозвороту в процедурах обробки інформації в реальному масштабі часу або в масштабі часу, близькому до реального, особливо коли потік даних утворює тимчасові послідовності. У цьому випадку обробка інформації послідовним способом у темпі її надходження забезпечує відомі переваги: по-перше, дозволяє не запам'ятовувати весь обсяг інформації і отримувати цікаві оцінки досліджуваних процесів у міру надходження даних, а по-друге, значно підвищує оперативність синтезованих алгоритмів, що використовуються в автоматизованих системах управління метрологічним обслуговуванням ЗВТ промислового підприємства, та суттєво знижує потрібний машинний час та пам'ять ЕОМ.

Для завдань цього класу математична модель (8) стає залежною від часу, тобто, стає, кажучи мовою теорії управління, нестационарною і набуває вигляду:

$$G_i X_i = Z_i \quad (9)$$

де X_i , Z_i – поточні значення параметрів матриці розміром $(n \times m)$ і $(1 \times m)$ відповідно, формуються з потоків інформації, що надходять, за певними законами; G_i – поточна вектор-рядок оцінюваних параметрів; i – поточний індекс часу, що підкреслює зміну інформації у матрицях X_i та Z_i .

Якщо в матриці X_i зміну інформації здійснювати по стовпцях (зліва направо) шляхом відкидання першого стовпця a_i (застаріла інформація) і додаванням до частини A_l , що залишилася, справа нового стовпця x (знову надійшла інформація), то така постановка дозволяє природним чином для вирішення рівняння (5) застосувати метод псевдозвороту за допомогою допоміжної матриці, який використовує формули (5) ... (7). При цьому слід зазначити, що принцип зміни інформації в матриці Z_i на відміну від матриці X_i може бути довільним.

Наприклад, його можна реалізувати шляхом поелементної заміни (зліва направо) шляхом відкидання першого елемента Z_i (застаріла інформація) і додавання до частини матриці Z_i , що залишилася, праворуч нового елемента Z_{m+i-1} (ново надійшла інформація).

Отже, поточна оцінка вектора рядка $\hat{G}_i$ параметрів моделі (9) визначається наступним рекурентним співвідношенням:

$$\hat{G}_i = Z_i X_i^+, \quad (10)$$

де операція псевдозвороту виконується за рекурентними формулами (5)... (7), тобто процедура послідовної обробки інформації фактично представлена як процедура послідовного оновлення (зі зміщенням) над колонками матриці заданих розмірів і її псевдорозвороту на кожному кроці зміни інформації.

Описаний метод є прямим і використовує переваги, властиві методу оздоблення. При псевдоконвертації за допомогою цього методу можна контролювати правильність розрахунків на кожному кроці за допомогою умов Пенроуза. Як зазначалося, цей метод особливо ефективний в процедурах обробки інформації в режимі реального часу. Це значно скорочує відставання в порівнянні з відомим і найбільш часто використовуваним послідовним методом найменших квадратів. Алгоритм забезпечує відсутність зміщення оцінок зі збільшенням шуму в даних, що надходять, так як вплив шуму попередніх етапів обробки не впливає і є можливість ефективно придушити випадкові шумові компоненти в оброблюваних послідовностях шляхом забезпечення надмірності ($m \gg n$).

Висновки:

1. Розглянуті та обгрунтовані недоліки визначення коефіцієнтів апроксимуючих поліномів для будь-яких вимірів і безперервних функцій, зокрема недоліки для методів середніх і найменших квадратів.

2. Для процедур обробки інформації у реальному масштабі часу запропоновано операцію псевдозвороту, яка виконується за допомогою рекурентних

формул і дозволяє послідовно оновлювати (зі зсувом) по стовбцям матриці заданих розмірів та її псевдозвороту на кожному кроці змін інформації.

3. Запропонована процедура може використовуватися при оптимізації та прогнозуванні параметрів і характеристик систем різного призначення, а також при вирішенні задач лінійної алгебри та статистики.

ЛІТЕРАТУРА

1. Шершень Д.А. Метод наименьших модулей. Магістерська дисертація.- Херсон: ХДУ, 2020. – 86 с.
2. Метод наименьших квадратов. Математическая статистика та обробка геологічних даних. С. 269-326 [Електронний ресурс]. Режим доступу: www.geo.univ.giev.ua. – Заголовок з екрану.
3. Математическая статистика: навчальний посібник для студентів вузів/ В.К. Гаркавий, В.В. Ярова. – К.: Професіонал, 2004. -379 с.
4. Лінійна алгебра та аналітична геометрія: навчальний посібник/ В.В. Булдігін, І.В. Алексеева, В.О. Гайдей та інші; за ред. проф. В.В. Булдігіна. – К.: ТВ і МС, 2011. – 224 с.
5. Чарін В.С. Лінійна алгебра. – К.: Техніка, 2018. – 416 с.
6. Безущак О.О. Навчальний посібник з лінійної алгебри для студентів механіко-математичних факультетів/О.О. Безущак, О.Г. Ганюшкін, Є.А. Кочубінська. – К.: ВПЦ «Київський університет», 2019. – 224 с
7. Ладієва Л.Р. Оптимізація технологічних процесів. Навчальний посібник. – К.: НМЦ ВО, 2013.- 206 с.
8. Цупак А.А. Псевдообратные матрицы. – Пенза, изд-во Пензенского государственного университета, 2008. – 29 с.
9. Альберт А. Регрессия, псевдоинверсия и рекуррентное оценивание. – М.:Наука, 1977. – 224 с.
10. Роджер Пенроуз. Путь к реальности или Законы управляющие Вселенной. Изд-во «Регулярная и хаотическая динамика» (РХД), 2007. – 912 с. (перевод с англ.).
11. Стивен Хокинг, Роджер Пенроуз. Природа пространства и времени. – Издательство АСТ, 2018. – 171 с.
12. ILAC-G17:2002. Introducing the Concept of Uncertainty of Measurement in Testing in Association with the Application of the Standard ISO/IEC 17025 (Вступ до концепції невизначеності вимірювань при випробуваннях з урахуванням застосування стандарту ISO/IEC 17025).

13. EA-04/02:1999. Expression of the Uncertainty of Measurement in Calibration (Відображення невизначеності вимірювань при калібруванні).
14. EA-04/16:2003. EA guidelines on the expression of uncertainty in quantitative testing (ЄА настанови щодо відображення невизначеності при кількісних випробуваннях).
15. EURACHEM/CITAC Guide QUAM-P1:2000. Quantifying Uncertainty in Analytical Measurement (Розрахунок невизначеності при аналітичних вимірюваннях).
16. RACHEM/CITAC Guide, Quantifying Uncertainty in Analytical Measurement, Second Edition. Laboratory of the Government Chemist, London (2000). ISBN 0-948926-15-5.
17. EURACHEM/CITAC Guide, Measurement uncertainty arising from sampling: A guide to methods and approaches. EURACHEM, (2007). Available from <http://www.eurachem.org>.

REFERENCES

1. Shershen D.A. The method of least modules. Master's thesis. - Kherson: KhSU, 2020. - 86 p.
2. The method of least squares. Mathematical statistics and geological data processing. P. 269-326 [Electronic resource]. Access mode: www.geo.univ.riev.ua. – Title from the screen.
3. Mathematical statistics: a study guide for university students/ V.K. Harkavy, V.V. Yarova. - K.: Professional, 2004. -379 p.
4. Linear algebra and analytic geometry: textbook/ V.V. Buldygin, I.V. Alekseeva, V.O. Gaidei and others; under the editorship Prof. V.V. Buldighina. - K.: TV and MS, 2011. - 224 p.
5. Charin V.S. Linear algebra. - K.: Technika, 2018. - 416 p.
6. Bezuschak O.O. Study guide on linear algebra for students of mechanical and mathematical faculties/O.O. Bezushchak, O.H. Hanyushkin, E.A. Kochubinska. - K.: VOC "Kyiv University", 2019. - 224 p.
7. Ladiyeva L.R. Optimization of technological processes. Tutorial. - K.: NMC VO, 2013. - 206 p.
8. Tsupak A.A. Pseudo-inverse matrices. - Penza, publishing house of Penza State University, 2008. - 29 p.
9. Albert A. Regression, pseudoinversion and recurrent estimation. - M.: Nauka, 1977. - 224 p.

10. Roger Penrose. The way to reality or Laws governing the Universe. Publishing house "Regular and Chaotic Dynamics" (RHD), 2007. - 912 p. (translation from English).
11. Stephen Hawking, Roger Penrose. The nature of space and time. – AST Publishing House, 2018. – 171 p.
12. ILAC-G17:2002. Introducing the Concept of Uncertainty of Measurement in Testing in Association with the Application of the Standard ISO/IEC 17025.
13. EA-04/02:1999. Expression of the Uncertainty of Measurement in Calibration.
14. EA-04/16:2003. EA guidelines on the expression of uncertainty in quantitative testing.
15. EURACHEM/CITAC Guide QUAM-P1:2000. Quantifying Uncertainty in Analytical Measurement (Calculation of uncertainty in analytical measurements).
16. RACHEM/CITAC Guide, Quantifying Uncertainty in Analytical Measurement, Second Edition. Laboratory of the Government Chemist, London (2000). ISBN 0-948926-15-5.
17. EURACHEM/CITAC Guide, Measurement uncertainty arising from sampling: A guide to methods and approaches. EURACHEM, (2007). Available from <http://www.eurachem.org>.

Received 02.05.2023.

Accepted 04.05.2023.

***Alternative to mean and least squares methods used in processing
the results of scientific and technical experiments***

Increasing the complexity and size of systems of various nature requires constant improvement of modeling and verification of the obtained results by experiment. It is possible to clearly conduct each experiment, objectively evaluate the summaries of the researched process, and spread the material obtained in one study to a series of other studies only if they are correctly set up and processed. On the basis of experimental data, algebraic expressions are selected, which are called empirical formulas, which are used if the analytical expression of some function is complex or does not exist at this stage of the description of the object, system or phenomenon. When selecting empirical formulas, polynomials of the form: $y = A_0 + A_1x + A_2x^2 + A_3x^3 + \dots + A_nx^n$ are widely used, which can be used to approximate any measurement results if they are expressed as continuous functions. It is especially valuable that even if the exact expression of the solution (polynomial) is unknown, it is possible to determine the value of the coefficients A_n using the methods of mean and least squares. But in the method of least squares, there is a shift in estimates when the noise in the data is increased, as it is affected by the noise of the previous stages

of information processing. Therefore, for real-time information processing procedures, a pseudo-reverse operation is proposed, which is performed using recurrent formulas. This procedure is a procedure of successive updating (with a shift) along the columns of the matrix of given sizes and pseudo-reversal at each step of information change. This approach is straightforward and takes advantage of the bounding method. With pseudo-inversion, it is possible to control the correctness of calculations at each step, using Penrose conditions. The need for pseudo-inversion may arise during optimization, forecasting of certain parameters and characteristics of systems of various purposes, in various problems of linear algebra, statistics, presentation of the structure of the obtained solutions, to understand the content of the incorrectness of the resulting solution, in the sense of Adomars-Tikhonov, and to see the ways of regularization of such solutions .

Keywords: method of mean and least squares; empirical formulas; pseudo reversal; experimental data; approximating polynomials.

Ігнаткін Валерій Устинович - доктор технічних наук, професор, професор кафедри автоматизації та комп'ютерно- інтегрованих технологій Івано-Франківського національного технічного університету нафти і газу.

Дудніков Володимир Степанович - кандидат технічних наук, доцент, доцент кафедри механотроніки Дніпровського національного університету імені Олеся Гончара.

Лучишин Тарас Романович - кандидат медичних наук, спеціаліст вищої категорії, заступник директора «Української готельної групи».

Алексєєнко Сергій Вікторович - доктор технічних наук, професор, завідувач кафедри механотроніки Дніпровського національного університету імені Олеся Гончара.

Юшкевич Олег Павлович - кандидат технічних наук, доцент, доцент кафедри механотроніки Дніпровського національного університету імені Олеся Гончара.

Карпова Тетяна Петрівна - старший викладач кафедри матеріалознавства і термічної обробки металів УДУНТ.

Хохлова Тетяна Станіславівна - кандидат технічних наук, доцент кафедри матеріалознавства і термічної обробки металів УДУНТ.

Хомош Юрій Степанович - кандидат економічних наук, доцент, директор Дрогобицького фахового коледжу нафти та газу.

Тіхонов Василь Андрійович - директор Дніпровського фахового коледжу радіоелектроніки.

Ignatkin Valery Ustinovich - doctor of technical sciences, professor, professor of the department of automation and computer-integrated technologies of the Ivano-Frankivsk National Technical University of Oil and Gas.

Dudnikov Volodymyr Stepanovich - candidate of technical sciences, associate professor, associate professor of the Department of Mechatronics of Oles Honchar Dnipro National University.

Luchyshyn Taras Romanovich - candidate of medical sciences, specialist of the highest category, deputy director of the "Ukrainian Hotel Group".

Alekseenko Serhii Viktorovich - Doctor of Technical Sciences, Professor, Head of the Department of Mechatronics, Dnipro National University named after Oles Honchar.

Yushkevich Oleh Pavlovich - candidate of technical sciences, associate professor, associate professor of the Department of Mechatronics of Dnipro National University named after Oles Honchar.

Karpova Tetyana Petrivna - senior lecturer of the Department of Materials Science and Heat Treatment of Metals, USUNT.

Khokhlova Tetyana Stanislavivna - candidate of Technical Sciences, Associate Professor of the Department of Materials Science and Heat Treatment of Metals at the Ukrainian National University of Science and Technology.

Khomosh Yuriy Stepanovich - candidate of economic sciences, associate professor, director of the Drohobysk Oil and Gas Professional College.

Tikhonov Vasyl Andriyovich - director of the Dnipro Vocational College of Radio Electronics.

**ГЛОБАЛЬНЕ ПОКРИТТЯ НАВКОЛОЗЕМНОГО ПРОСТОРУ
ЗОНАМИ ВИКОРИСТАННЯ ПРИСТРОЇВ ЙОГО СПОСТЕРЕЖЕННЯ:
КОНЦЕПЦІЯ І АЛГОРИТМИ**

Анотація. Представлені результати дослідження в рамках задачі забезпечення повного по криття заданої області висот над поверхнею Землі (області простору між двома сферами зі спільним центром у центрі Землі) миттєвими зонами можливого застосування пристроїв спостереження орбітального базування, розташованими на космічних апаратах в різновисоких орбітальних угрупованнях на колових орбітах. У загальному випадку вирішення задачі передбачає використання декількох різновисоких орбітальних угруповань на колових квазіполярних орбітах, які в спрощеній постановці задачі прийняті полярними. Миттєва зона можливого застосування пристрою спостереження спрощено прийнята у вигляді конусу. Розглянуті випадки застосування пристроїв спостережень «вверх» (над площиною миттєвого місцевого горизонту космічного апарату, що є носієм пристрою спостереження) і спостережень «вниз» (під цією площиною). Запропонована концепція вирішення задачі, яка базується на виборі (на основі розвинення методів застосування відомих алгоритмів) такої структури кожного орбітального угруповання, яка забезпечить неперервне покриття частини заданого простору спостереження (області гарантованого спостереження), границі якої відсутні від розташування пристроїв спостереження, а після того – заповнення простору цими областями. Робота присвячена космічній тематиці, але, узагальнивши постановку задачі, варіюючи низку умов цієї постановки та змінюючи «масштаб» вхідних даних, можна прийти до різноманіття технічних задач, де будуть доцільні і прийнятні (частково або повністю) запропонована концепція та алгоритми, застосовані при її реалізації.

Ключові слова: неперервне покриття області простору зонами у вигляді конусів, зони застосування пристроїв спостереження, пристрої спостереження орбітального базування, супутникові системи спостереження орбітальних об'єктів, космічні апарати, різновисокі орбітальні угруповання.

Постановка проблеми. Множина об'єктів на навколоземних орбітах стрімко зростає [1-3]. Поки ще не вирішена проблема космічного сміття, яке зараз складає основну частину цієї множини [1,2], а також відбувається стрибкоподібне збільшення кількості космічних апаратів (зараз, в основному, внаслідок створення угруповань багатосупутникових систем [3], а в подальшому –

завдяки розвитку орбітального сервісу і космічного виробництва [4]). Вимоги безпечного використання навколоземного космосу, прагнення знижувати всі види ризиків застосування космічної техніки та підвищувати ефективність космічних систем призводять до необхідності контролю множини орбітальних об'єктів в навколоземному просторі. Зокрема, є необхідним розвиток сучасних систем і технологій спостереження множини орбітальних об'єктів задля підтримання в актуальному стані різноманітної інформації про її проточний стан як динамічної механічної системи, а також про різні аспекти функціонування космічних апаратів [5-8]. При цьому має бути забезпечене «тотальне» спостереження, як з точки зору підвищення можливостей щодо реалізації спостережень об'єктів малого розміру, так і з точки зору наявності ресурсів спостереження у кількісному і якісному складі для повного неперервного «просторового охоплення» множини тих орбітальних об'єктів, які дозволяє спостерігати вибірково здатність пристроїв спостереження.

Ще не так давно задача оновлення інформації про множину орбітальних об'єктів розглядалася в основному для засобів наземних станцій спостереження. Але зараз стає все більш актуальним застосування (та в подальшому зладжене використання з наземними засобами) супутникових систем відповідного призначення [9-13]. Дана робота присвячена одному з аспектів створення глобальної супутникової системи спостереження орбітальних об'єктів, зокрема – «остова» цієї системи, побудованого на різновисоких угрупованнях регулярної структури повного охоплення («огортання» Землі) [7,8,14]. Саме така система має в майбутньому забезпечуватиме, по-перше, повне системне покриття навколоземного простору зонами використання пристроїв спостереження відносно далекої дії застосування (дальність застосування може передбачати можливий діапазон від сотень метрів до десятків тисяч кілометрів) і, по-друге, в подальшому доповнюватися застосуванням невеликих орбітальних угруповань з динамічною, високо адаптивною до вирішення конкретних задач структурою, які наближаються до об'єктів і використовують пристрої спостереження невеликої дальності (в межах від сотень метрів до десятків кілометрів).

Завдання спостереження множини орбітальних об'єктів вирішується на основі синергетичного поєднання застосування різних типів пристроїв спостереження та принципів їх використання (будуть розвиватися далі оптичні і радіолокаційними засоби, видові спостереження та спостереження на основі чергових сеансів, тощо). В даній роботі тип пристрою спостереження не виділений, а взята до розгляду тільки та область простору, в якій на конкретну мить

пристрій спостереження орбітального базування (встановлений на космічному апараті) може бути застосований (миттєва зона його можливого застосування). В роботі розглянута проблема повного покриття заданої області висот над поверхнею Землі миттєвими зонами можливого застосування пристроїв спостереження орбітального базування при спрощеному представленні зони можливого застосування пристрою спостереження у вигляді конусу з вершиною в точці розташування пристрою спостереження (у конкретному випадку – в точці знаходження центру космічного апарату, який є носієм цього пристрою).

Аналіз останніх досліджень і публікацій. Названа вище проблема до недавнього часу не була затребувана саме в такий постановці, але можна назвати проблеми зі схожим формулюванням (саме на застосуванні підходів до їх вирішення базується підхід, представлений в даній роботі). Описана задача «перекликається» з задачею забезпечення безперервного глобального покриття поверхні Землі (що розглядається як сфера), зонами застосування пристроїв супутникової системи зв'язку або системи дистанційного зондування Землі.

На поточний час достатньо розвинута задача балістичного проектування (вибору параметрів орбітального угруповання) супутникової системи з метою забезпечення покриття заданої території обслуговування на поверхні Землі (зокрема – глобального або квазі глобального покриття поверхні Землі) зонами застосування пристроїв, встановлених на космічних апаратах цієї системи (в роботах [15-17] представлені базові, класичні підходи до вирішення питання побудови орбітального угруповання за для забезпечення неперервного покриття поверхні Землі із використанням мінімально можливої кількості космічних апаратів угруповання). При колових орбітах космічних апаратів спрощено і формалізовано названі задачі можна звести до розгляду двох концентричних сфер. Задача при цьому полягає до забезпечення покриття поверхні сфери меншого радіусу (поверхні Землі) із застосуванням пристроїв спостереження з конусовидними зонами можливого миттєвого застосування, вершини яких знаходяться на сфері більшого радіусу (на сфері, якій належать орбіти). «Класичні» алгоритми вирішення цієї задачі залежать від вибору одного з базових принципів побудови орбітального угруповання супутникової системи: 1) на основі «замкнених ланцюжків» («кілець») супутників, коли в кожній номінальній орбітальній площині розташовано декілька супутників [15,16]; 2) побудові угруповання як кінематично кінематично правильної структури [17]. На зараз, коли стає актуальним реалізація міжсупутникового зв'язку (і в системах з головною функцією передачі інформації, і в системах іншого призначення, які бу-

дуть застосовувати міжсупутникові комунікації) зростає актуальність побудови орбітальних угруповань на основі «кілець» супутників (у тому числі це буде актуальним і для систем, що реалізують спостереження). Тому знов представляє інтерес розвиток алгоритмів першої групи. Описана задача покриття сфери заданого радіусу розвивається для випадку використання декількох різновисоких угруповань на колових орбітах при збереженні тієї ж мети – забезпечити покриття поверхні сфери зонами обслуговування, що створюють пристрої, встановлені на космічних апаратах цієї системи.

В даній роботі постановка задач змінюється, ускладнюється. Формулюючи узагальнено, відмінність поставленої задачі в тому, що обирають параметри орбітального угруповання, космічні апарати якого є носіями пристроїв спостереження орбітальних об'єктів, і має бути забезпечено не покриття поверхні сфери зонами застосування цих пристроїв, а покриття простору між двома концентричними сферами (інакше, в більш «вузькій» постановці, – покриття простору в заданій області висот над поверхнею Землі). Але задача неперервного покриття сфери заданого радіусу увійде до запропонованого в даній роботі методу забезпечення покриття простору між двома сферами, стане елементарною часткою повного алгоритму розрахунків щодо вирішення задачі покриття простору миттєвими зонами застосування пристроїв спостереження. У загальному випадку вирішення задачі буде потребувати застосування космічних апаратів на різновисоких орбітальних угрупованнях.

В сучасній літературі вже з'являються дослідження задачі забезпечення покриття простору пристроями цільової апаратури космічних апаратів, розташованих в різновисоких орбітальних угрупованнях. Наприклад, ставиться задача мінімізації структури орбітального угруповання при забезпеченні неперервного покриття із заданою кратністю сферичних шарів в навколоземному просторі зонами застосування пристроїв цільової апаратури, розташованої в різновисоких орбітальних угрупованнях, та пропонується алгоритм комп'ютерного підбору параметрів такої структури при їх варіюванні.

В даній своїй роботі не ставимо задачі мінімізації кількості космічних апаратів орбітального угруповання та забезпечення покриття навколоземного простору в заданій області висот над поверхнею Землі миттєвими зонами застосування пристроїв спостереження з кратністю більшою за одиницю. Але тут визначена базова концепція побудови супутникової системи спостереження орбітальних об'єктів на основі заповнення простору в заданій області висот «полосами» неперервного обзору (більш малими за розміром зонами висот), коли

кожна полоса неперервного обзору забезпечується своїм орбітальним угруповання космічних апаратів. При цьому незалежно розглянуто забезпечення «висотних полос» спостереження як у випадку, коли пристрої спостереження орбітального базування ведуть спостереження «вниз» (спостерігають об'єкти на більш низьких орбітах), так і для випадку, коли такі пристрої застосовуються для спостереження «вверх» (спостерігають орбітальні об'єкти, що знаходяться на більш високих орбітах). Запропонована концепція системи передбачає адитивне застосування обох видів пристроїв спостереження. Запропоновано алгоритм, який надає можливість вибору параметрів описаної системи, що відповідає поставленим умовам.

Проблема покриття заданої області висот множиною пристроїв орбітального базування в схожому до представленого в даній роботі формулюванні вже розглядалася в роботі [14], де були запропоновані два підходи до її вирішення. У запропонованих матеріалах досліджень постановка задачі і стратегія її розв'язку змінюються, більш націлені на вирішення проблеми глобального покриття простору в заданій області висот над поверхнею Землі зонами миттєвого застосування пристроїв спостереження суто в «геометричному» формулюванні.

Мета дослідження. Розвиток основ побудови супутникових систем глобального спостереження заданої області висот в навколоземному просторі. Зокрема – розробка принципів побудови та вибору параметрів орбітального угруповання космічних апаратів (носіїв пристроїв спостереження) на колових орбітах, яке забезпечить постійне, неперервне покриття заданої області висот над поверхнею Землі прийнятими конусовидними зонами миттєвого застосування пристроїв спостереження орбітальних об'єктів.

Викладення основного матеріалу.

1. Постановка задачі. Розглядається супутникова система спостереження орбітальних об'єктів. Пристрої спостереження встановлені на космічних апаратах. Для космічного апарату підтримується кутова орієнтація відносно осей зв'язаної з його центром мас барицентричної орбітальної системи координат. Як зазначалося вище, за розташуванням зони можливого застосування пристрою спостереження виділено два типи пристроїв (рис. 1а). Пристрої першого типу ведуть «спостереження вверх» – над площиною миттєвого місцевого горизонту (над площиною, перпендикулярною до радіус-вектору космічного апарату), а пристрої другого типу здійснюють «спостереження вниз» – під цією площиною.

Зони застосування пристрою спостереження, що належить одному з названих типів, максимально спрощено представлена конусом, розміри і розташування якого у просторі задають такі умови (рис. 1а): 1) вершина конусу знаходиться в точці положення центру мас космічного апарату; 2) вісь симетрії конусу співпадає з прямою, якій належить радіус-вектор точки розташування космічного апарату (інакше – з прямою, перпендикулярною до площини миттєвого місцевого горизонту; 3) значення кута β між напрямком з точки розташування космічного апарату на точку положення центру мас об'єкту спостереження і напрямком уздовж вісі симетрії конусу від вершини до основи (для пристроїв спостереження першого типу – уздовж радіус-вектору космічного апарату, що є носієм пристрою; для пристроїв спостереження другого типу – проти напрямку радіус-вектору) не має перевищувати задане значення $\beta_m - \beta \leq \beta_m$; 4) довжина відрізка поверхні конусу, що поєднує його вершину і основу, має значення L_m (так спрощено врахована умова, що дальність L від пристрою спостереження до об'єкту спостереження має не перевищувати це значення – $L \leq L_m$).

Орбітальне угруповання супутникової системи спостереження побудовано за таким принципом (рис. 1б): 1) є k різновисоких угруповань на колових орбітах (k різновисоких сегментів), на рис. 1б такі угруповання показані повністю, або фрагментарно; 2) вся система побудована на квазіполярних орбітах, але при проектуванні системи будемо вважати, що орбіти полярні; 3) орбіти колові, висота орбіт всіх космічних апаратів k -го сегменту – h_k ; 4) в k -му сегменті угруповання M_k номінальних орбітальних площин; 5) в кожній m -й орбітальній площині k -го сегменту N_{km} космічних апаратів (для всіх них одне й теж саме значення довготи висхідного вузла Ω_{km}), тобто в кожній орбітальній площині представлене «кільце» космічних апаратів; 6) орбітальні площини симетрично рознесені за значенням довготи висхідного вузла.

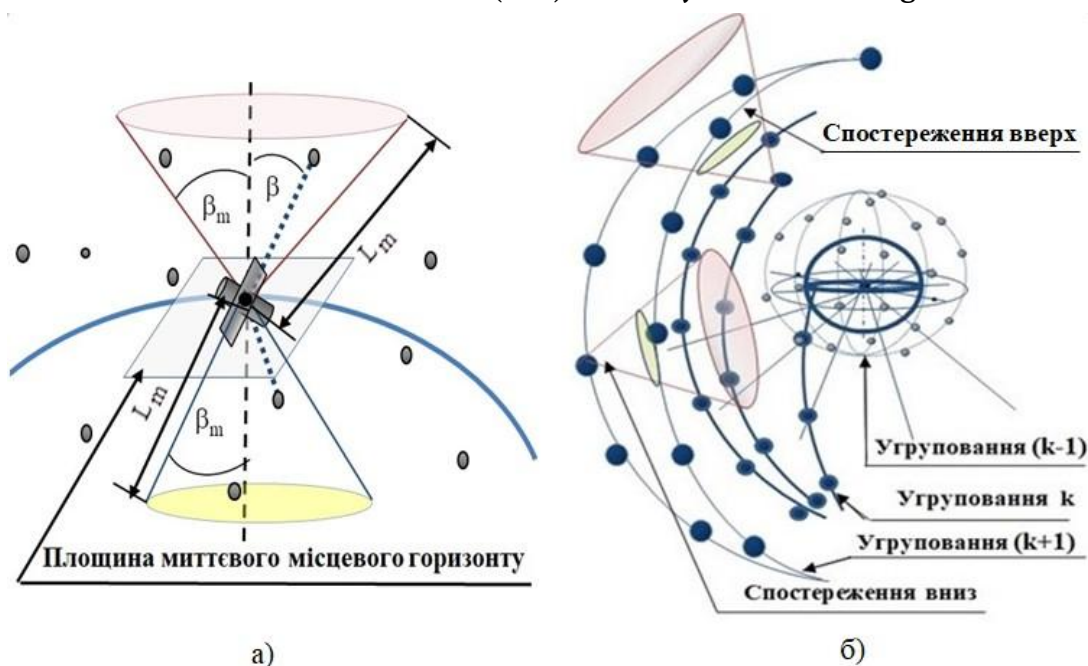


Рисунок 1 – До принципу побудови системи спостереження:
 область миттєвого покриття простору зонами можливого застосування
 пристроїв спостереження з космічного апарату (а);
 фрагменти орбітальних угруповань (б)

Необхідно обрати параметри всього орбітального угруповання, щоб забезпечити кожній миті безперервне покрити зонами спостереження всього простору в заданій області висот над поверхнею Землі (висота h для кожної точки, що належить до області спостереження, знаходиться в діапазоні між заданими значеннями $h_{\min}$ і $h_{\max}$, тобто $h \in [h_{\min}, h_{\max}]$). В даній постановці задачі прийнято, що значення параметру L_m також задане і однакове для всіх пристроїв спостереження. Значення кута β_m має бути однаковим для всіх пристроїв спостереження k -го сегменту супутникової системи (позначається β_{mk}) і в ході реалізації алгоритму обирається на основі умова реалізації спостережень і технічних можливостей апаратної реалізації.

В даній роботі не ставиться завдання оптимізації за якимось критерієм рішення щодо вибору параметрів орбітального угруповання. Завдання більш просте – обрати одну з можливих структур угруповання, яка відповідає поставленим умовам та запропонувати для цього один з можливих розрахункових алгоритмів, який зв'яже вхідні і вихідні данні (на деяких елементарних етапах реалізації алгоритму застосовується підхід, який спрямований на зменшення кількості космічних апаратів в системі спостереження). Запропонований алго-

ритм може бути покладеним в основу ітераційного процесу пошуку квазіоптимальних рішень, або розвинутий до формулювання та розв'язку оптимізаційної задачі із застосуванням аналітичних залежностей.

2. °Стратегія вирішення задачі. Розглянемо як незалежні задачі покриття простору зонами застосування пристроїв, що здійснюють «спостереження вниз», та тих, що здійснюють «спостереження вверх» (не будемо враховувати, що пристрої обох типів раціонально встановлювати на одному з космічних апаратів описаної системи, і це буде давати свій відбиток на пошуку комплексного рішення). Нехай, як зазначалося, неперервне покриття області висот $h \in [h_{\min}, h_{\max}]$, що спостерігається, забезпечують зони миттєвого застосування пристроїв «спостереження вниз», встановлених на космічних апаратах K орбітальних угруповань на колових, прийнятих полярними орбітах. Кожне k -те угруповання ($k = \overline{1, K}$) забезпечує неперервне покриття «полоси висот» між значеннями $h_{\min k}$ і $h_{\max k}$ ($h \in [h_{\min k}, h_{\max k}]$), рис. 2а. Прийнемо, що ці «полоси висот» не перетинаються. Розташування K описаних «полос висот» «встик» (точніше – відповідних частин («шарів») простору між концентричними сферами з центром в центрі Землі і радіусами, які дорівнюють $r_{\min k} = R_E + h_{\min k}$ і $r_{\max k} = R_E + h_{\max k}$, де R_E – радіус Землі) забезпечить неперервне покриття зонами застосування пристроїв спостереження області висот $h \in [h_{\min}, h_{\max}]$ (області між концентричними сферами, центр яких співпадає з центром Землі, а радіуси дорівнюють $r_{\min} = R_E + h_{\min}$ і $r_{\max} = R_E + h_{\max}$).

За описаним нижче підходом до вирішення задачі «Полоса висот» гарантованого повного покриття простору буде «відсунута» від висоти, де розташовані космічні апарати, що є носіями пристроїв спостереження. На рис. 2.б для випадку «спостережень вниз» схематично показано розташування декількох космічних апаратів на одній з орбіт k -го орбітального угруповання та полоса висот, для покриття якої будуть застосовані пристрої, встановлені на космічних апаратах угруповання. Аналогічне схематичне представлення приведено на рис. 2б для випадку реалізації «спостережень вверх». На рис. 2а показаний випадок неперервного заповнення заданої області висот «полосами» висот застосування пристроїв «спостереження вниз», і для k -го сегменту системи зображено розташування космічних апаратів на одній з орбіт цього сегменту (ко-

смічні апарати k -го сегменту знаходяться в полосі висот більш високого сегменту спостереження, в даному випадку – $(k + 1)$ -го сегменту).

При такому підході задача забезпечення покриття заданої області висот $[h_{\min}, h_{\max}]$ розкладається на k незалежні задач (елементарних задач) вибору структур орбітальних угруповань, кожної k -е з яких забезпечить повне неперервне покриття своєї полоси висот $[h_{\min k}, h_{\max k}]$.

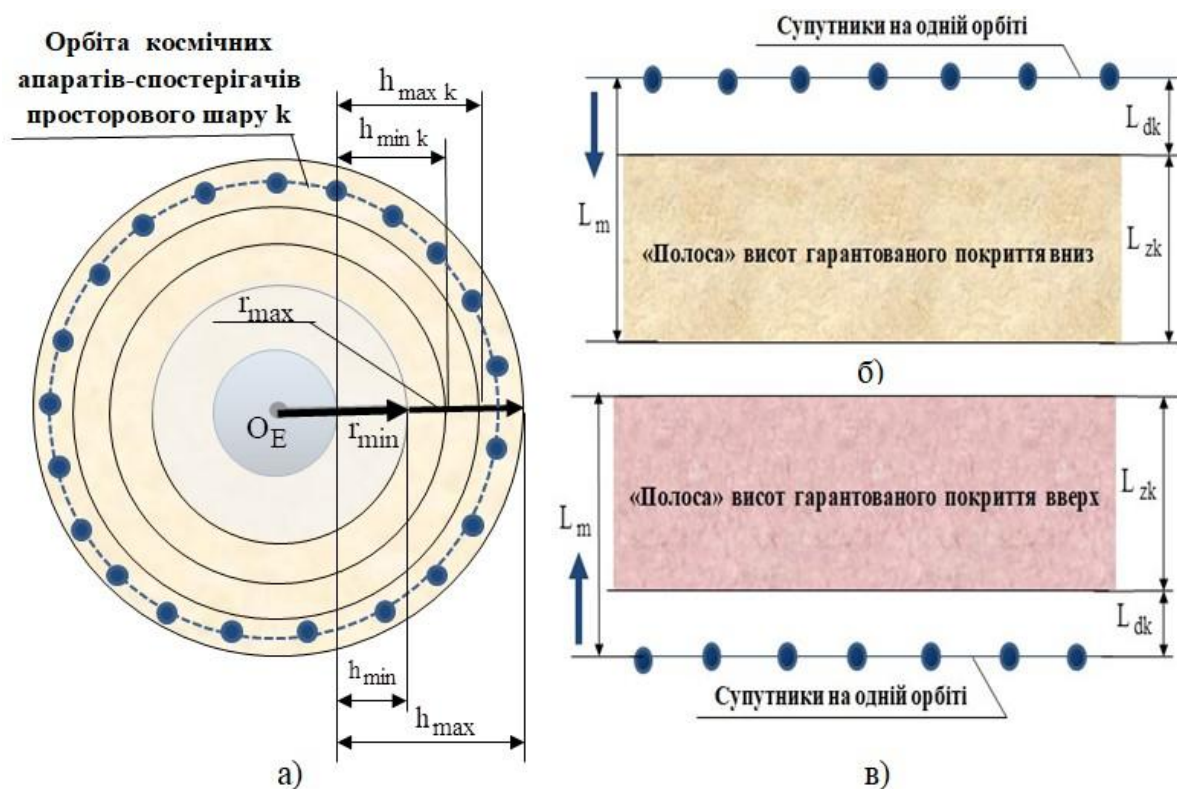


Рисунок 2 – До принципу неперервного заповнення області висот спостереження: покриття області висот «шарами» неперервного спостереження вниз (заповнення «полосами» висот) (а); «полоса» висот неперервного спостереження вниз (б); «полоса» висот неперервного спостереження вверх (в)

Аналогічний підхід реалізується при незалежному вирішенні задачі із застосуванням пристроїв «спостереження вверх», встановлених на космічних апаратах k угруповань на колових орбітах.

3. Елементарна складова задачі. Під елементарною складовою задачі будемо розуміти визначення параметрів k -го сегменту супутникової системи, побудованої на колових полярних орбітах, що містить m_k орбітальних площин, в кожній з яких знаходяться n_{km} космічних апаратів – носіїв пристроїв спосте-

реження, симетрично рознесених у орбітальній площині («уздовж номінальної орбіти»). Обираються такі параметри: 1) висота орбіти h_k , на якій розташовано орбітальне угруповання; 2) параметри m_k і n_{km} , що визначають кількість N_k космічних апаратів в k -му сегменті системи ($N_k = m_k \cdot n_k$) та їх розташування в орбітальних площинах.

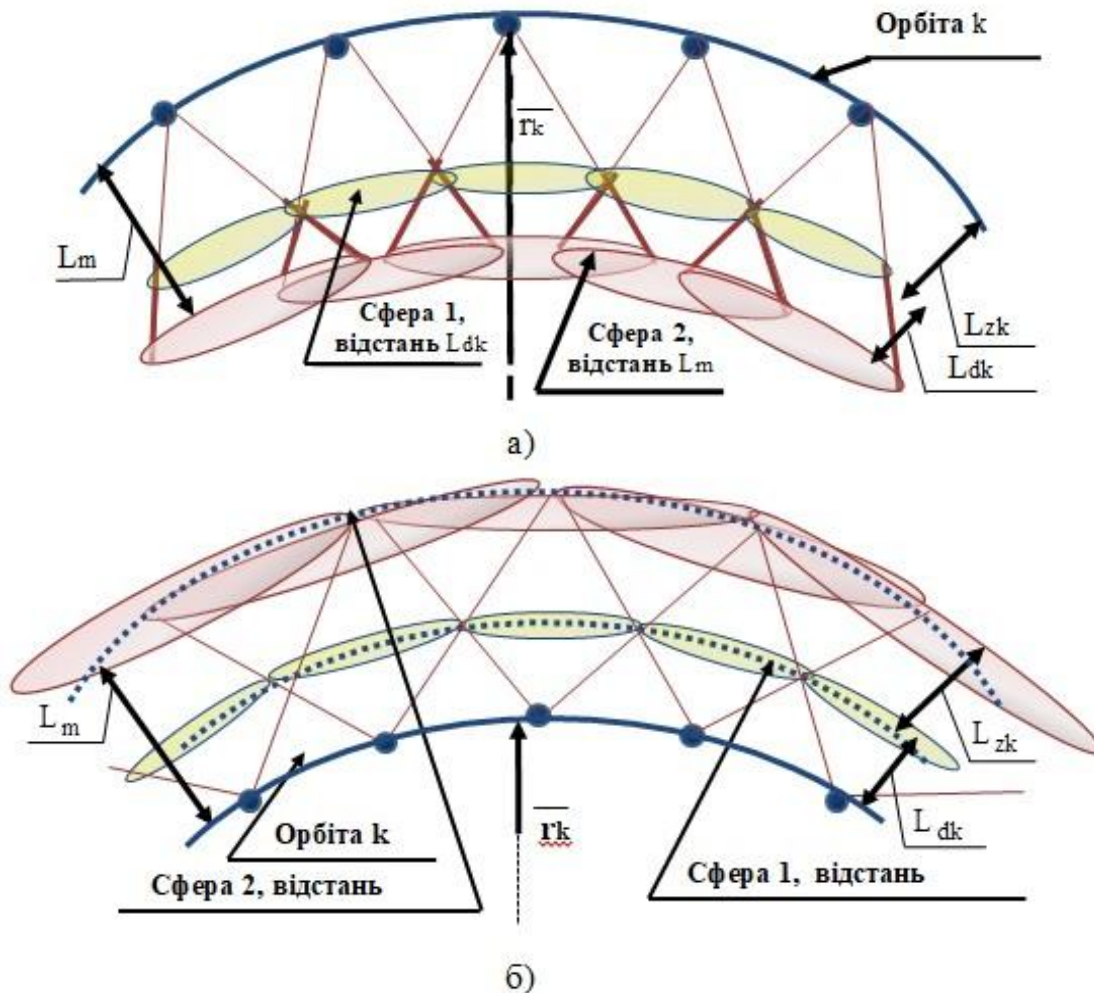


Рисунок 3 – Фрагмент створення частини полоси неперервного спостереження пристроями на одній орбіті: при «спостереженні вниз» (а); при «спостереженні верх» (б)

Розпочнемо з представлення елементарної задачі для випадку, коли застосовуються пристрої «спостереження вниз». Спочатку розглянемо «ланцюжок» супутників в одній орбітальній площині (рис. 2б, рис. 3а). Позначимо радіус орбіт космічних апаратів, що належать угрупованню k -го сегменту системи, – r_k . Тоді за дальністю область максимальної досяжності дії пристроїв «спостереження вниз» – це область між сферами радіусів r_k і $r_{mk} = r_k - L_m$.

Оберемо ще одну сферу з радіусом $r_{dk} = r_k - L_d$, де $L_d < L_m$ (тобто сфера з радіусом r_{kd} ближча до орбітального угруповання, що несе пристрої спостереження, розташована між сферами з радіусами r_{km} і r_k). При вирішенні задачі вважатиме, що пристроями спостереження, встановленими на космічних апаратах k -го угруповання системи, має бути забезпечено гарантоване безперервне покриття області простору між сферами з радіусами r_{mk} і r_{dk} (тобто полоси висот між значеннями $h_{\min k} = r_k - L_d - R_E$ і $h_{\max k} = r_k - L_m - R_E$ шириною $L_z = L_m - L_d$).

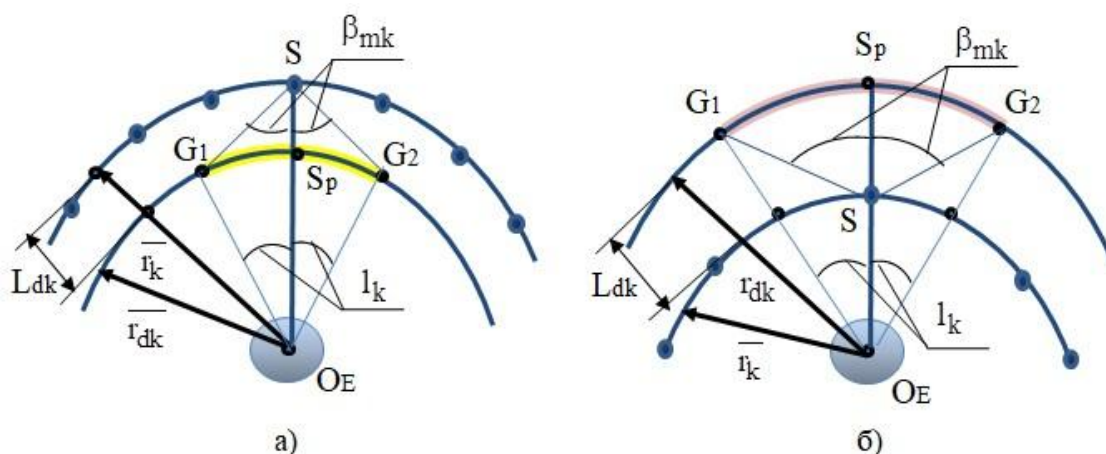


Рисунок 4 – До розрахунку зони покриття пристрою спостереження на визначальній сфері: при «спостереженні вниз» (а); при «спостереженні вверх» (б)

Повне покриття області між сферами з радіусами r_{mk} і r_{dk} миттєвими зонами можливого застосування пристроїв спостереження буде досягнуте, якщо для кожного моменту часу буде забезпечено повне покриття поверхні сфери радіусу r_{dk} . Задача неперервного покриття поверхні сфери з радіусом r_{dk} повністю ідентична задачі повного покриття поверхні Землі (яка розглядається як сфера), зонами можливого застосування відповідних пристроїв супутникової системи зв'язку або дистанційного зондування. На рис. 4а представлена відповідна картина у площині одної з орбіт k -го угруповання системи. Аналогічно задачам покриття поверхні Землі для кожного космічного апарату (на рис. 4а позначення точки його розташування – S) можна побачити під супутникову точку S_p , дугу на поверхні сфери r_{kd} з кутовою величиною $2l_k$, що визначає розмір миттєвої зони покриття на сфері пристроєм, встановленим на відповідному космічному апараті).

Визначимо величину дуги l_k аналогічно тому, як це значення знаходиться при визначенні умов прямої видимості космічного апарату з наземної станції (умова появи космічного апарату над площиною місцевого горизонту наземної станції), тільки при цьому застосовуються не радіус Землі і радіус орбіти космічного апарату, а, відповідно, радіус r_{dk} сфери, що покривається зонами спостереження, та радіус r_k сфери розташування космічних апаратів. Виходячи з цього

$$l_k = \arccos\left(\frac{r_{dk}}{r_k}\right). \quad (1)$$

Вважається також, що значення β_k може бути реалізованим при визначенні його з тієї ж умови із застосуванням відповідного до цього рівняння.

Далі задача зводиться до покриття поверхні сфери радіусу r_{dk} зонами, що визначені дугами $2l_k$. Для цього обраний метод, запропонований у джерелі [15] для випадку полярних орбіт. Згідно з базовими положеннями методу, представленого в роботі [15], обирається кількість орбітальних площин m_k і однакова кількість n_{mk} космічних апаратів в кожній з цих орбітальних площин (пояснення дані на рис. 5а і 5б). При цьому кількість космічних апаратів n_{mk} має забезпечити полосу неперервного огляду, ширину якої на сфері радіусу r_{dk} визначає дуга $2b_k$ ($b_k < l_k$). Дуга b_k показана на рис. 5а знизу окружності для випадку, коли підсупутникова точка знаходиться на окружності, а дуга l_k на тій же окружності, – зверху, також для випадку, коли підсупутникова точка знаходиться на окружності. Величина дуги b_k прямо пропорційна «ступеню перетинання» суміжних підсупутникових зон на одній орбіті (рис. 5а, рис. 5б). Інакше, b_k тим більша, чим менша величина дуги a_k на сфері радіусу r_{dk} , яка є половиною кутової відстані між підсупутниковими точками суміжних супутників на одній орбіті (на рис. 5б і рис. 6 це дуга між точками S_{pj} і K_1). Величина дуги a_k пов'язана з кількістю N_{mk} космічних апаратів на орбіті співвідношенням [15]

$$N_{km} = \left\lceil \frac{\pi}{a} \right\rceil + 1, \quad (2)$$

тут $\lceil \cdot \rceil$ в формулах означає цілу частину числа. Повне покриття поверхні сфери полосами неперервного огляду з шириною b_k забезпечить кількість орбітальних площин (відповідно – кількість полос), що визначена рівнянням [15]

$$M_k = \left\lceil \frac{\pi}{2b} \right\rceil + 1 \quad (3)$$

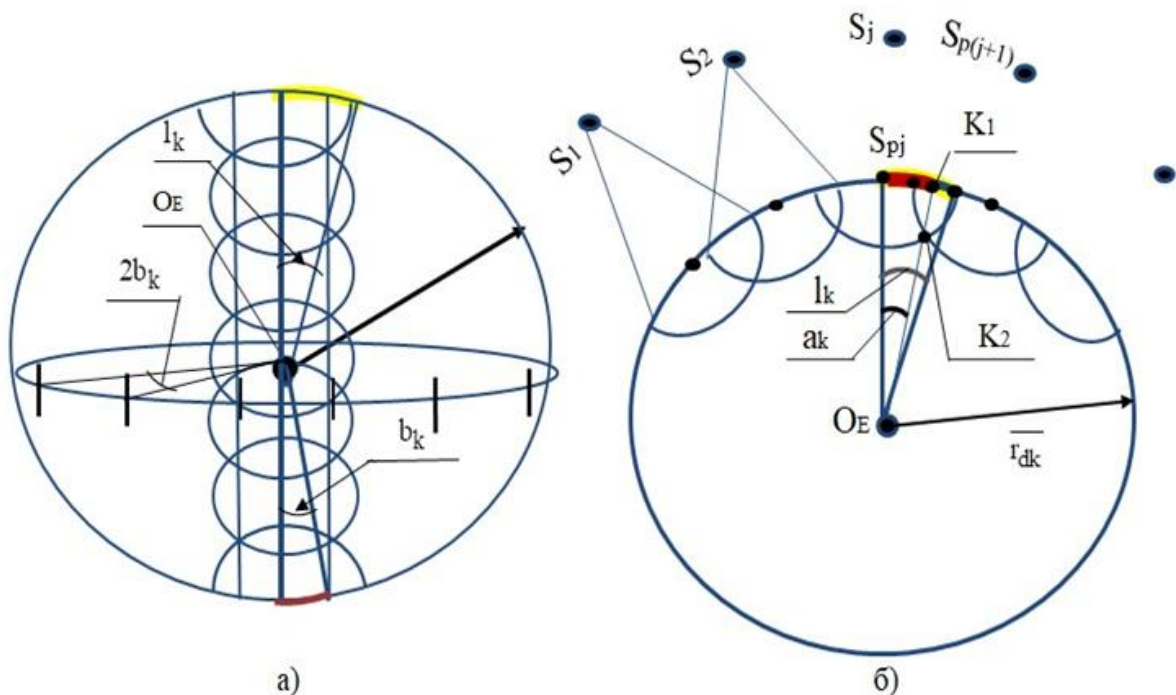


Рисунок 5 – До визначення кількості орбітальних площин і кількості супутників у площинах за джерелом (зроблено на основі рисунків джерела [15]): глобальне покриття полосами неперервного обзору (а); розташування космічних апаратів у орбітальній площині (б)

Величина дуг a_k і b_k також зв'язані між собою. Цей зв'язок можна виразити формулою, вирішуючи задачу сферичної тригонометрії при розгляді сферичного прямокутного трикутника $S_{p(j+1)}K_3K_2$ на рис. 6 (на якому випадок розташування на екваторі підсупутникової точки $S_{p(j+1)}$). Величина дуги K_3K_2 дорівнює величині дуги $S_{p(j+1)}K_1$ (дуга $S_{p(j+1)}K_1$ також показана на рис. 5в) і має кутову величину a_k . Дуга $S_{p(j+1)}K_3$ дорівнює значенню b_k половини ширини полоси неперервного огляду (є частиною дуги $S_{p(j+1)}K_4$, що має величину l_k). З трикутника $S_{p(j+1)}K_3K_2$ маємо

$$a_k = \arccos \left(\frac{\cos l_k}{\cos b_k} \right). \quad (4)$$

Запишемо величину дуги b_k , як частину дуги l_k при заданому коефіцієнті пропорційності k_{bk} , що зв'яже ці величини

$$b_k = k_{bk} l_k. \quad (5)$$

З врахуванням виразів (4) і (5) рівняння (2) і (3) будуть записані так:

$$M_k = \left\lceil \frac{\pi}{2 k_{bk} l_k} \right\rceil + 1, \quad (6)$$

$$N_{km} = \left\lceil \frac{\pi}{\arccos \left(\frac{\cos l_k}{\cos(k_{bk} l_k)} \right)} \right\rceil + 1. \quad (7)$$

Варіюючи значення коефіцієнту k_{bk} з прийнятим кроком в межах інтервалу $]0, 1[$, можна обрати значення кількості орбітальних площин M_k і кількості N_{km} космічних апаратів у кожній площині, які забезпечать при даній постановці задачі мінімальне значення N_k ($N_k = M_k N_{km}$) космічних апаратів в угрупованні k -го сегменту системи (показано у прикладі далі).

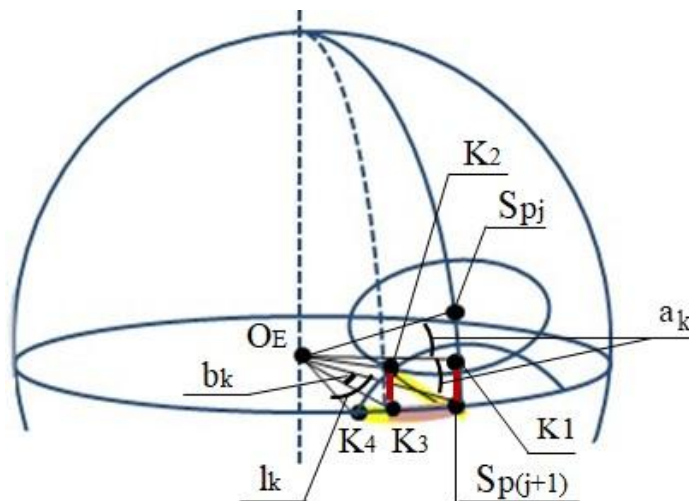


Рисунок 6 – До вирішення задачі сферичної тригонометрії задля знаходження зв'язку між дугами a_k і b_k (на основі рисунку джерела [15])

Далі розглянемо елементарну задачу для випадку коли пристрої орбітального угруповання ведуть «спостереження вверху» (рис. 2в, рис. 3б). За аналогією з попереднім випадком, коли ведеться спостереження вниз, позначимо радіус сфери, на який розташовані орбіти космічних апаратів-спостерігачів k -го орбітального угруповання, – r_k . Радіус сфери, що визначає границю дальності, на яку можуть здійснюватися спостереження, – $r_{mk} = r_k + L_m$. Радіус сфери, що визначає ту границю, від якої буде враховано покриття простору пристроями спостереження k -го сегменту системи, – $r_{dk} = r_k + L_d$. Тобто спостерігатися

буде полоса висот шириною L_z між концентричними сферами радіусів r_{dk} і r_m . Знов реалізується покриття сфери, ближчої до висоти орбіт розташування космічних апаратів, що несуть пристрої спостереження. Тут проекція точки розташування супутника S на поверхню сфери з радіусом r_{dk} – також S_p . Величина дуги l_k визначається при заданому значенні β_{mk} ($\beta_{mk} > l_k$) на основі рівняння

$$l_k = \beta_{mk} - \arcsin\left(\frac{r_k}{r_{dk}} \sin \beta_{mk}\right). \quad (8)$$

Після визначення величини дуги l_k алгоритм подальшого рішення задачі вибору кількості орбітальних площин N_{km} і кількості космічних апаратів в кожній орбітальній площині M_k співпадає з описаною вище задачею реалізації спостережень вверх.

4.°Приклади розрахунків. За даним підходом наведено два приклади розрахунків. В табл. 1 представлений вибір структури системи, побудованої на трьох орбітальних угрупованнях, яка призначена для спостереження вниз і повністю покриває зонами застосування пристроїв спостереження область висот над поверхнею Землі в діапазоні від 400 км до 1000 км при заданому значенні границі максимальної дальності спостереження $L_m = 300$ км. Для кожного k -го сегменту системи прийнято значення $L_{dk} = 100$.

Таблиця 1

Приклад орбітального угруповання космічних апаратів, що реалізують спостереження вниз при значенні $L_m = 300$ км

Номер сегменту k	Кількість КА	Кількість орбітальних площин	Кількість КА у площині	Висота угруповання	Полоса покриття	
					$h_{\min k}$	$h_{\max k}$
3	378	14	27	1100	800	1000
2	364	14	26	900	600	800
1	357	13	27	700	400	600

Приклад вибору значення коефіцієнту k_{bk} , представлений для третього сегменту обраного в прикладі орбітального угруповання системи, що реалізує «спостереження вниз», представлено на рис. 7 (при варіюванні значення k_{b3} обрано значення $k_{b3} = 0,69$, яке забезпечує найменше значення $N_3(k_{b3}) = 378$ кількості космічних апаратів в угрупованні).

В табл. 2 представлений вибір структури системи, побудованої на трьох орбітальних угрупованнях, яка веде спостереження вверх і повністю покриває зонами застосування пристроїв спостереження область висот над поверхнею Землі в діапазоні від 500 км до 1100 км при заданому значенні границі максимальної дальності спостереження $L_m = 400$ км. Для кожного k -го сегменту системи прийнято значення $L_{dk} = 200$.

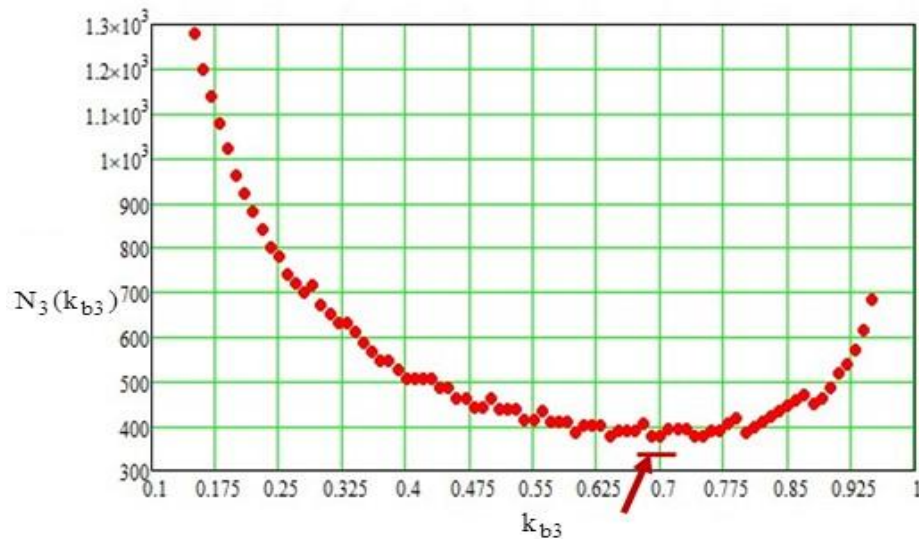


Рисунок 7 – До вибору значення коефіцієнту k_{bk} – вибір значення коефіцієнта k_{b3} сегменту 3 ($k = 3$) у представленому прикладі побудови угруповання спостереження вниз

Таблиця 2

Приклад орбітального угруповання космічних апаратів, що реалізують спостереження вверх при значенні $L_m = 400$ км

Номер сегменту k	Кількість КА	Кількість орбітальних площин	Кількість КА у площині	Висота угруповання	Полоса покриття	
					$h_{\min k}$	$h_{\max k}$
3	242	11	21	700	900	1100
2	240	12	20	500	700	900
1	231	11	22	300	500	700

Висновки. Основним результатом даної роботи є представлення простого алгоритму, який може бути покладений в основу ітераційного алгоритму пошуку квазіоптимального рішення задачі при варіюванні даних, що обираються в процесі рішення (а також при варіюванні вхідних даних). В описаній постановці головним чином обираються значення відстані (L_{kd}) для кожного k -го

угруповання системи спостереження, на яку «полоса висот» гарантованого покриття простору миттєвими зонами застосування пристроїв спостереження цього угруповання відсунута від висоти, на якій знаходиться угруповання. Значення відстаней L_{kd} і L_m визначають ширину полоси висот зони гарантованого використання пристроїв спостереження, кількість цих полос та їх розташування. Вибір значення L_{kd} може бути заснованим тільки на врахуванні заданої мінімальної дальності, яка має бути між об'єктом спостереження та космічним апаратом, що є носієм пристрою спостереження. А в основному її вибір продиктований прагненням «відсунути» ближчу границю зони спостереження задля зменшення кількості космічних апаратів угруповання, які реалізують спостереження (перевага від цього більше явна при реалізації «спостережень вверх»). Але на альтернативу цьому є і прагнення наблизити до пристроїв спостереження границі області, яка гарантовано спостерігається, задля підвищення якості спостережень. При цьому слід відзначити, що пристрої спостереження можуть бути застосовані і до границі, з якої при проектуванні угруповання починається виконання умови можливості їх застосування у будь якій точці.

В представленому підході «полоси висот» гарантованого спостереження простору не перекриваються, але побудова систем спостереження множини орбітальних об'єктів у навіколоземному просторі має передбачати їх перекриття в різний спосіб, задля чого базові підходи і алгоритми вирішення елементарних задач, представлені в даній роботі, також стануть у пригоді.

Ще одним результатом є отримання (як прикладу) деяких показників щодо необхідної кількості космічних апаратів-спостерігачів, які будуть здатні гарантовано покрити миттєвими зонами застосування пристроїв спостереження орбітального базування велику частину області висот низькоорбітальних супутникових систем. Приведені приклади показують, що при дальності застосування пристроїв спостереження в декілька сотень кілометрів мінімально необхідний склад орбітального угруповання супутникової системи спостереження – біля тисячі одиниць космічних апаратів в декількох різновисоких орбітальних угрупованнях, що для сучасних глобальних систем не є критично великим числом.

Робота присвячена космічній тематиці, але, узагальнивши постановку задачі, варіюючи низку умов цієї постановки та змінюючи «масштаб» вхідних даних, можна прийти до різноманіття технічних задач, де будуть доцільні і прийнятні (частково або повністю) запропонована концепція та алгоритми, за-

стосовані при її реалізації. Зокрема – при створенні систем спостереження або систем комплексного застосування пристроїв виконання технічних операцій.

ЛІТЕРАТУРА

1. ESA'S annual space environment report [Електронний ресурс]: Prepared by ESA Space Debris Office. – Date of Issue 12 June 2023. Режим доступу: https://www.sdo.esoc.esa.int/environment_report/Space_Environment_Report_latest.pdf
2. Our world in Data. Number of objects in the orbits of the Earth. [Електронний ресурс]: Режим доступу: <https://ourworldindata.org/grapher/low-earth-orbits-objects?country=Debris~Payloads~Rocket+bodies>
3. Space explained: How much space junk is there?: [Електронний ресурс]: <https://www.inmarsat.com/en/insights/corporate/2022/how-much-space-junk-is-there.html>
4. Васильев В.В. Орбитальний сервіс — крок до подальшого освоєння навколоземного космосу / В.В. Васильев, Л.Я. Годунок, С.А. Матвієнко // Космічна наука і технологія. 2021. 27, № 3 (130). С. 39-50. <https://doi.org/10.15407/knit2021.03>
5. Випорханюк Д.М. Основи космічної ситуаційної обізнаності (Space Situation Awareness, SSA). Іноземний і вітчизняний досвід космічної діяльності у сфері оборони: монографія./ Д.М. Випорханюк, С.В. Ковбасюк // Житормир: Видавництво О.О. Євєнук, 2018. – 532 с.
6. Беспалко І.А. Концепція інформаційної системи для забезпечення моніторингу космічного простору з метою підвищення воєнної безпеки. / І.А. Беспалко, Л.Д. Греков, Д.В. Пекарьєв, Д.Л. Федорчук // Космічна наука і технологія. 2022. 28, № 4 (137). С. 3—17. <https://doi.org/10.15407/knit2022.04.003>
7. Лабуткіна Т.В. Всеобщая глобальная космическая система наблюдения Земли и космоса в аспекте мира и безопасности землян, акцент на орбитальной составляющей / Т.В. Лабуткіна, А.В. Хлапоніна // Наукові читання «Дніпровська орбіта-2020»: Збірник доповідей. - Дніпро, НЦАОМ, 2020. - С. 120-130. URL: https://dneprorbita.org.ua/_files/doc/sbornik2020.pdf
8. Ананко Р.В., Лабуткіна Т.В. Навколоземний космос, контрольований людством: системність підходів, глобальність рішень, система-спостерігач на навколоземних орбітах. // Друга науково-практична Інтернет-конференція «Космічні горизонти», третій етап конференції - «Космос для людства». Збірник тез, НЦАОМ, Дніпро, 1-3 грудня, 2022. - С. 33-43. URL: https://space-horizons.org.ua/uploads/source/archiv_2022_3/tezu_3_2022.pdf
9. Mark R. Ackermann, Colonel Rex R. Kiziah, Peter C. Zimmer, John T. McGraw, David D A systematic examination of ground-based and space-based approaches to optical detection and tracking of satellites. // 31st Space Symposium, Technical Track,

Colorado Springs, Colorado, United States of America Presented on April 14, 2015. – SAND2015-3726C.

<https://www.osti.gov/servlets/purl/1253293>Emiliano Cordelli

10. Zhao Li, Yidi Wang, Wei Zheg Space-Based Optical Observation of Space Debris via Multipoint of View // Hindawi International Journal of Aerospace Engineering. // Volume 2020, Article ID 8328405, 12 pages. <https://doi.org/10.155/2020-83284>

11. Лабуткіна Т.В., Скородень Я.А., Борщева А.В., Тихонова А.А. Неітераційний метод планування спостереження орбітальних об'єктів з космічного апарата // Системне проектування та аналіз характеристик аерокосмічної техніки. – 2016. – Том XXI – С. 53-69.

12. HU Yunpeng¹, LI Kebo¹, LIANG Yan'gang¹, and CHEN Lei^{2,*} Review on strategies of space-based optical space situational awareness

<https://ieeexplore.ieee.org/document/9612138>

13. Лабуткіна Т.В Концепція системи з наземними і орбітальними засобами спостереження орбітальних об'єктів: стратегії використання засобів /Т.В. Лабуткіна, А.В. Хлапоніна, О.Р. Акіншев // Multidisciplinary academic notes. Theory, methodology and practice. Proceedings of the XVII International Scientific and Practical Conference. Tokyo, Japan. 2022. Pp. 1060-1069 URL: <https://isg-konf.com/multidisciplinary-academic-notes-theory-methodology-and-practice/>

14. Лабуткіна Т.В., Ананко Р.В. До концепції складової супутникової системи спостереження орбітальних об'єктів на основі стабільних регулярних угруповань космічних апаратів // Proceedings of the V International Scientific and Practical Conference. Stockholm, Sweden. 2023. Pp. 616-625. URL: <https://isg-konf.com/prospects-of-modern-science-and-education/>.

15. Баринов К.Н., Бурдаев М.Н., Мамон П.А. Динамика и принципы построения орбитальных систем космических аппаратов. М.: Машиностроение, 1975. 270 с.13.

16. Walker J.G. Satellite Constellations. Journal of the British Interplanetary Society, 1984, vol. 24, pp. 369–384.

17. Lang T.J. A Parametric Examination of Satellite Constellations to Minimize Revisit Time for Low Earth Orbits Using a Genetic Algorithm. AAS/AIAA Astrodynamics Specialist Conference, Quebec, Canada, 30 July–2 August 2001, vol. 109, pp. 625–640

REFERENCES

1. ESA'S annual space environment report [Електронний ресурс]: Prepared by ESA Space Debris Office. – Date of Issue 12 June 2023. Режим доступу: https://www.sdo.esoc.esa.int/environment_report/Space_Environment_Report_latest.pdf

2. Our world in Data. Number of objects in the orbits of the Earth. [Електронний ресурс]: Режим доступу: <https://ourworldindata.org/grapher/low-earth-orbits-objects?country=Debris~Payloads~Rocket+bodies>
3. Space explained: How much space junk is there?: [Електронний ресурс]: <https://www.inmarsat.com/en/insights/corporate/2022/how-much-space-junk-is-there.html>
4. Vasyliev V.V. Orbitalnyi servis — krok do podalshoho osvoiennia navkolozemnoho kosmosu / V.V. Vasyliev, L.Ia. Hodunok, S.A. Matviienko // Kosmichna nauka i tekhnolohiia. 2021. 27, № 3 (130). С. 39-50.
<https://doi.org/10.15407/knit2021.03>
5. Vyporkhaniuk D.M. Osnovy kosmichnoi sytuatsiinoi obiznanosti (Space Situation Awareness, SSA). Inozemnyi i vitchyzniani dosvid kosmichnoi diialnosti u sferi oborony: monohrafiia./ D.M. Vyporkhaniuk, S.V. Kovbasiuk // Zhytormyr: Vydavnytstvo O.O. Yeveniuk, 2018. – 532 pp.
6. Bepalko I.A. Kontseptsiiia informatsiinoi systemy dlia zabezpechennia monitorynhu kosmichnoho prostoru z metoiu pidvyshchennia voiennoi bezpeky. / I.A. Bepalko, L.D. Hrekov, D.V. Pekariev, D.L. Fedorchuk // Kosmichna nauka i tekhnolohiia. 2022. 28, № 4 (137). С. 3—17 <https://doi.org/10.15407/knit2022.04.003>
7. Labutkina T.V. Vseobshaya globalnaya kosmicheskaya sistema nablyudeniya Zemli i kosmosa v aspekte mira i bezopasnosti zemlyan, akcent na orbitalnoj sostavlyayushej / T.V. Labutkina, A.V. Khlaponina // Naukovi chytannia «Dniprovska orbita-2020»: Zbirnyk dopovidei. - Dnipro, NTsAOM, 2020. - С. 120-130. URL: https://dneprorbita.org.ua/_files/doc/sbornik2020.pdf
8. Ananko R.V., Labutkina T.V. Navkolozemnyi kosmos, kontrolovanyi liudstvom: systemnist pidkhodiv, hlobalnist rishen, systema-sposterihach na navkolozemnykh orbitakh. // Druha naukovo-praktychna Internet-konferentsiia «Kosmichni horyzonty», tretii etap konferentsii - «Kosmos dlia liudstva». Zbirnyk tez, NTsAOM, Dnipro, 1-3 hrudnia, 2022. - С. 33-43.
URL: https://space-horizons.org.ua/uploads/source/archiv_2022_3/tezu_3_2022.pdf
9. Mark R. Ackermann, Colonel Rex R. Kiziah, Peter C. Zimmer, John T. McGraw, David D A systematic examination of ground-based and space-based approaches to optical detection and tracking of satellites. // 31st Space Symposium, Technical Track, Colorado Springs, Colorado, United States of America Presented on April 14, 2015. - SAND2015-3726C.
<https://www.osti.gov/servlets/purl/1253293EmilianoCordelli>
10. Zhao Li. Yidi Wang, Wei Zheg Space-Based Optical Observation of Space Debris via Multipoint of View // Hindawi International Journal of Aerospace Engineering. // Volume 2020, Article ID 8328405, 12 pages. <https://doi.org/10.155/2020-83284>

11. Labutkina T.V., Skoroden Ya.A., Borshcheva A.V., Tykhonova A.A. Neyteratsyonnyi metod planirovaniya nabliudeniya orbitalnykh ob'ektov s kosmicheskogo apparata // Systemne proektuvannia ta analiz kharakterystyk aerokosmichnoi tekhniki. – 2016. – Tom XXI – С. 53-69.
12. HU Yunpeng¹ , LI Kebo¹ , LIANG Yan'gang¹ , and CHEN Lei^{2,*} Review on strategies of space-based optical space situational awareness <https://ieeexplore.ieee.org/document/9612138>
13. Labutkina T.V. Kontseptsiiia systemy z nazemnymi i orbitalnymi zasobamy sposterezhennia orbitalnykh ob'ektiv: stratehii vykorystannia zasobiv /T.V. Labutkina, A.V. Khlaponina, O.R. Akinshev // Multidisciplinary academic notes. Theory, methodology and practice. Proceedings of the XVII International Scientific and Practical Conference. Tokyo, Japan. 2022. Pp. 1060-1069 URL: <https://isg-konf.com/multidisciplinary-academic-notes-theory-methodology-and-practice/>
14. Labutkina T.V., Ananko R.V. Do kontseptsii skladovoi suputnykovoi systemy sposterezhennia orbitalnykh ob'ektiv na osnovi stabilnykh rehuliarnykh uhrupovan kosmichnykh aparativ // Proceedings of the V International Scientific and Practical Conference. Stockholm, Sweden. 2023. Pp. 616-625. URL: <https://isg-konf.com/prospects-of-modern-science-and-education/>.
15. Barinov K.N., Burdaev M.N., Mamon P.A. Dinamika i principy postroeniya orbitalnykh sistem kosmicheskikh apparatov. M.: Mashinostroenie, 1975. 270 s.13.
16. Walker J.G. Satellite Constellations. Journal of the British Interplanetary Society, 1984, vol. 24, pp. 369–384.
17. Lang T.J. A Parametric Examination of Satellite Constellations to Minimize Re-visit Time for Low Earth Orbits Using a Genetic Algorithm. AAS/AIAA Astrodynamics Specialist Conference, Quebec, Canada, 30 July–2 August 2001, vol. 109, pp. 625–640.

Received 04.05.2023.

Accepted 08.05.2023.

***Global near-earth space coverage by zones of the use of its observation devices:
concept and algorithms***

The results of the study are presented within the framework of the task of ensuring full coverage of a given area of heights above the Earth's surface (the area of space between two spheres with a common center at the center of the Earth) by instantaneous zones of possible application of orbital-based surveillance devices located on spacecraft in orbital groups of different heights in circular orbits. In the general case, the solution of the problem involves the use of several orbital groupings of different heights on circular quasi-polar orbits, which in the simplified statement of the problem are assumed to be polar. The instantaneous zone of possible application of the surveillance device is simplified in the form of a cone. The cases of using observation devices "up" (above the plane of

the instantaneous local horizon of the spacecraft, which is the carrier of the observation device) and observations "down" (below this plane) are considered. The concept of solving the problem is proposed, which is based on the selection (based on the development of methods of applying known algorithms) of such a structure of each orbital grouping, which will ensure continuous coverage of a part of the given observation space (area of guaranteed observation), the boundaries of which are moved away from the location of observation devices, and then - filling the space with these areas. The work is devoted to the space theme, but by generalizing the statement of the problem, varying a number of conditions of this statement and changing the "scale" of the input data, it is possible to arrive at a variety of technical problems where the proposed concept and algorithms used in its implementation will be appropriate and acceptable (in part or in full). In particular, when some surveillance systems or systems of complex application of technical operations devices are created.

Лабуткіна Тетяна Вікторівна – к.т.н., доцент, доцент кафедри кібербезпеки та комп'ютерно-інтегрованих технологій фізико-технічного факультету Дніпровського національного університету імені Олеся Гончара.

Ананко Руслан Віталійович – аспірант кафедри кібербезпеки та комп'ютерно-інтегрованих технологій фізико-технічного факультету Дніпровського національного університету імені Олеся Гончара.

Labutkina Tetyana Viktorivna - Ph.D., Associate Professor, Associate Professor of the Department of Cyber Security and Computer-Integrated Technologies of the Faculty of Physics and Technology of Oles Honchar Dnipro National University.

Ananko Ruslan Vitaliyovych - graduate student of the Department of Cyber Security and Computer-Integrated Technologies of the Faculty of Physics and Technology of Oles Honchar Dnipro National University.

МЕТОД ПОБУДОВИ ЦИФРОВОГО ДВІЙНИКА ВІБРОЗАХИСНОГО ПРОЦЕСУ

Анотація. У статті запропоновано метод побудови цифрових двійників віброзахисного процесу для висотних гнучких будівель. В основі методу лежить використання кульового гасника, що дозволяє підбирати оптимальні значення для зменшення руйнівного впливу коливань.

Ключові слова: цифровий двійник, віброзахисна система, кульовий гасник, швидкодія методу, вимушені коливання.

Постановка проблеми. Технологія цифрових двійників набуває все більшого значення, як для виробничих підприємств, так і для побутового життя. Попри всю потужність технології та велику кількість переваг вона не завжди може бути застосована через ряд різноманітних причин. Одним із таких випадків є створення цифрового двійника віброзахисного процесу для висотних гнучких споруд. При експлуатації висотних гнучких споруд (телевеж, радіощогл, металевих вентиляційних труб, димарів тощо) часто виникають вимушені коливання, боротьба з якими перетворюється на велику технічну проблему. Оскільки даний вид споруд є надзвичайно важливим для людства, постає науково-технічна задача створення методу побудови цифрового двійника віброзахисного процесу для висотних гнучких будівель. При цьому важливою задачею є знаходження оптимальних параметрів цифрового двійника для зменшення руйнівного впливу від коливань.

Аналіз останніх досліджень і публікацій. Аналіз основних підходів до створення цифрових двійників [1] показує, що жоден з них не має явних переваг над іншим, а найкращий результат досягається при їх комбінації. Проте варто зауважити, що підхід на основі даних вимагає збору великої кількості інформації за допомогою сенсорів на існуючих об'єктах. У випадку з віброзахисними системами для висотних гнучких будівель даний варіант не можливий, оскільки у більшості випадків вони не були впроваджені. Ось чому потрібно розглядати підхід на основі моделювання фізичних об'єктів/систем за допомогою високоточних математичних моделей [2]. Для аналізу існуючих математичних

моделей віброзахисного процесу, розглянемо детальніше сам процес появи та гасіння вимушених коливань.

Вимушені коливання висотних об'єктів виникають при їх взаємодії із повітряним потоком і здійснюються як у площині вектора повітряного потоку (коливання під дією вітрової пульсації), так і в ортогональній напрямку вектора повітряного потоку площині (автоколивання типу “повітряний резонанс”). При цьому частота основного тону власних коливань низькочастотних висотних споруд, що розглядаються, знаходиться в діапазоні 0,2–2,0 рад/с [3].

До останнього часу для розв'язання цієї проблеми використовувались динамічні гасники маятникового типу на підвісі [4]. Функціонування маятників такої конструкції однозначно визначається довжиною підвісу їх робочого тіла. Для телевеж та баштових споруд розрахункова довжина підвісу таких гасників має знаходитись у межах від 4 до 15 м. Тому спроби використати маятникові гасники на малих частотах (менше 2,0 рад/с) були невдалими. Ситуація, що склалася, і обумовила необхідність розробки гасителів такої конструкції, яка б дозволила зберегти маятниковий характер функціонування робочого тіла гасителя.

В цих умовах найбільш ефективним є новий метод віброзахисту низькочастотних висотних споруд із використанням гасителів коткового типу [5]. Ефект віброзахисту таких гасителів ґрунтується на перекочуванні важких куль, циліндрів, роликів без ковзання по сферичних або циклоїдальних поверхнях несучих тіл. Експериментальні дослідження функціонування коткових гасителів довели їх ефективність при демпфіруванні вимушених коливань гнучких висотних споруд у низькочастотному діапазоні (до 2,0 рад/с) та при великих переміщеннях верхніх перерізів споруд (до 2,5 – 3 м).

У дослідженні [6] було запропоновано та математично обґрунтовано новий метод оптимального налаштування (тюнінгу) основних параметрів кульових гасників цього класу на власну частоту несучого тіла, що дозволяє використовувати його у сучасній практиці віброзахисту висотних гнучких об'єктів.

Мета дослідження. Основною метою даної роботи є побудова методу створення цифрових двійників віброзахисного процесу, який забезпечить можливість визначення регулюючих параметрів кульового гасника для оптимального гасіння вимушених коливань.

Викладення основного матеріалу дослідження. У даній роботі розглядається процес придушення вимушених коливань віброзахисної системи з ви-

користанням кульових гасників (BVA - Ball Vibration Absorber). Принципова схема одної із таких систем показана на рис. 1.

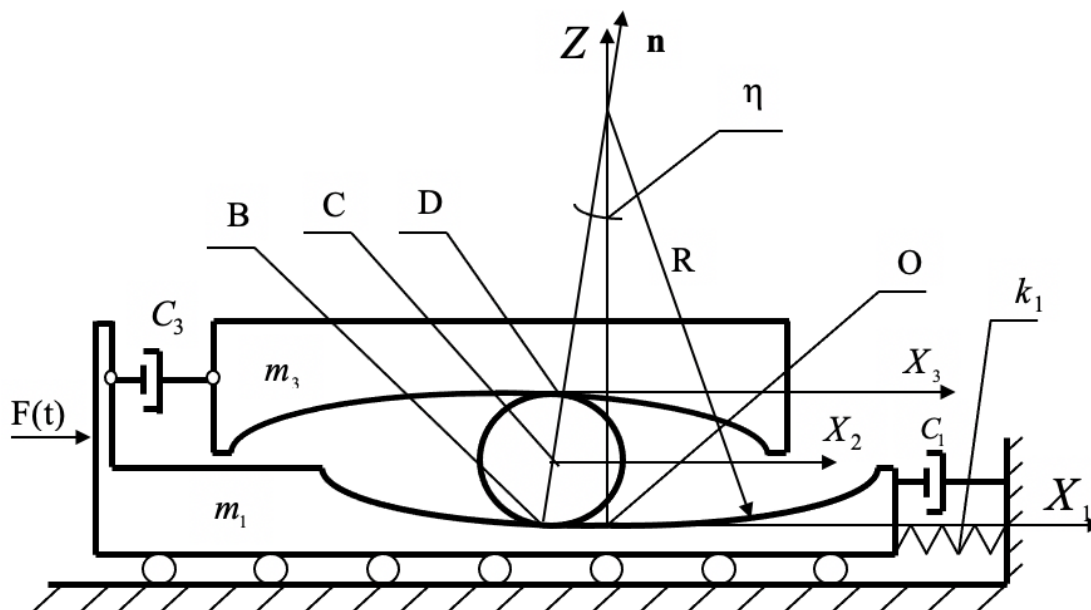


Рисунок 1 - Принципова схема кульового гасника

У ході попереднього дослідження [6] для такої системи було описано наступне рівняння амплітудно-частотної характеристики $A(\omega)$ системи “несуче тіло – котковий гасник”:

$$A(\omega) = \frac{\bar{F}_0 \sqrt{F_6(\omega)}}{\sqrt{F_6(\omega)F_7(\omega) + F_9(\omega) + F_{10}(\omega)}}, \quad (1)$$

де $F_1(\omega) = \omega_0^2 - \omega^2(1 + \nu)$; $F_2(\omega) = 2n_X\omega$; $F_3(\omega) = g - 2\bar{R}\omega^2$; $F_4(\omega) = 4n_\eta\bar{R}\omega$; $F_5(\omega) = 2\nu\bar{R}\omega^4$; $F_6(\omega) = (F_3(\omega))^2 + (F_4(\omega))^2$; $F_7(\omega) = (F_1(\omega))^2 + (F_2(\omega))^2$; $F_8(\omega) = F_2(\omega)F_4(\omega) - F_1(\omega)F_3(\omega)$; $F_9(\omega) = 2F_5(\omega)F_8(\omega)$; $F_{10}(\omega) = (F_5(\omega))^2$.

Визначимо параметри віброзахисної системи, базуюсь на принциповій схемі системи (рис. 1) та отриманому рівнянні амплітудно-частотної характеристики (1).

У результаті отримуємо наступний набір параметрів:

- m_1 – маса несучого тіла;
- m_3 – маса робочого тіла;
- g – прискорення вільного падіння;
- ω_0 – кругова або циклічна частота;
- $\bar{F}_0$ – максимальна амплітуда гармонічної збуджуючої сили;
- k_1 – коефіцієнт пружності несучого тіла;
- R – радіус виїмок;

- r – радіус кулі, що знаходиться всередині виїмок;
- v – відношення мас робочого та несучого тіл;
- C_1 – коефіцієнт в'язкого опору демпфера, пов'язаного з несучим тілом;
- C_3 – коефіцієнт в'язкого опору демпфера, пов'язаного з робочим тілом;
- n_x – коефіцієнт демпфірування демпфера, пов'язаного з несучим тілом;
- n_η – коефіцієнт демпфірування демпфера, пов'язаного з робочим тілом;
- $\bar{R}$ – різниця радіусів сферичних виїмок і кулі, що перекочується між ними.

Беручи до уваги те, що з будівельної точки зору віброзахисна система найлегше буде конфігуруватися через параметри саме робочого тіла гасника [7], тобто коефіцієнт демпфірування n_η та різницю радіусів виїмок та кулі $\bar{R}$, оберемо їх за вихідні параметри методу. У такому випадку ми будемо мати п'ять вхідних параметрів: $\bar{F}_0$ – максимальна амплітуда зовнішнього силового збудження; v – відношення мас робочого та несучого тіл; n_x – коефіцієнт демпфірування демпфера, пов'язаного з несучим тілом; ω_0 – кругова частота власних коливань несучого тіла, яка є найменшою (або головною) з його спектру частот та точність пошуку вихідних параметрів.

Оскільки було визначено вхідні та вихідні параметри, далі потрібно визначити, яким чином буде здійснюватись пошук оптимальних значень вихідних параметрів. Найпростішим способом знаходження оптимальних значень є звичайні ітеративні операції, де обидва вихідні параметри поступово інкрементуються, а отримані значення амплітуди при них зберігаються у певному словнику, а потім порівнюються. Такий спосіб є доволі очевидним та незалежним від умов, проте дуже повільним.

З точки зору оптимізації швидкодії необхідно виконати певну модифікацію, щоб пришвидшити метод побудови цифрового двійника. Для цього можемо використати аналітично-графічний метод знаходження оптимальних параметрів налаштування гасників коткового типу [8-9]. В основі цього аналітично-графічного методу лежить принцип «рівності двох максимумів», які досягаються на двох частотах ω_1 і ω_2 в околі резонансної частоти ω_0 . Як параметри налаштування розглядаються геометрична характеристика $\bar{R} = R - r$ гасника та коефіцієнт демпфірування n_η його робочого тіла, що і є вихідними параметрами запропонованого методу.

Аналітично-графічний метод оптимального налаштування гасника по частоті полягає у побудові фрагментів функціональних залежностей амплітуди A несучого тіла від характеристики $\bar{R}$ гасника в околі двох частот ω_1 і ω_2 за умови фіксації інших параметрів системи. Перетин двох, знайдених таким чином фрагментів кривих АЧХ, і визначає оптимальну величину характеристик $\bar{R}$ чи n_η гасника. Розглянемо детальніше аналітично-графічний метод.

Спочатку параметри $\bar{R}$ і n_η аналітично визначаються із умови рівності двох максимумів графіка функції $A = A(\omega)$, один з яких відповідає зведений масі m_1 несучого тіла (споруди), а другий – масі m_2 робочого тіла гасника. Для початку роботи з методом знаходимо початкові частоти ω_1 і ω_2 , на яких досягаються два максимуми амплітудно-частотної характеристики для віброзахисної системи, за формулою (2) [9]:

$$\gamma_{1,2}^2 = \frac{1}{1+v} \left(1 \pm \sqrt{\frac{v}{2+v}} \right), \text{ де } \gamma_{1,2} = \frac{\omega_{1,2}}{\omega_0} \quad (2)$$

Далі графічно визначаємо точку перетину двох графіків A_1 та A_2 побудованих для фіксованих частот $\omega_1 = \gamma_1 \omega_0$ та $\omega_2 = \gamma_2 \omega_0$ при змінюванні параметра $\bar{R}$ гасника, на яких функція $A = A(\omega)$ досягає двох рівних максимумів. Після визначення параметра $\bar{R}$ гасника застосовуємо такий же підхід з побудовою двох графіків функції для коефіцієнта демпфірування n_η . У результаті отримуємо наступні блок-схеми запропонованих методів (див. рис.2-3).

Аналіз отриманих результатів застосування створеного методу

Запропонований метод створення цифрових двійників був реалізований за допомогою мови програмування Python та протестований на наборі різних даних, зібраних на основі реальних об'єктів. Для виконання операцій з математичними функціями була використана бібліотека `sympy`. З огляду на пришвидшення роботи методу для знаходження нулів математичних функцій були застосовані чисельні методи бібліотеки `sympy`.

Input: F_0, w_0, n_x, v
Output: R, n_n



Рисунок 2 - Блок-схема методу побудови цифрового двійника віброзахисної системи з оптимальними параметрами

Input: F_0, w_0, n_x, v
Output: R, n_n

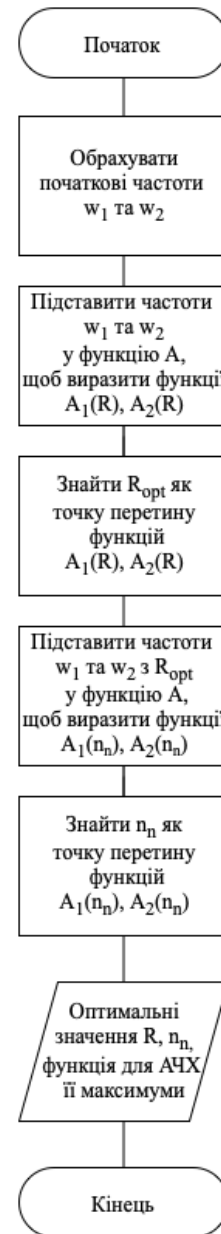


Рисунок 3 - Блок-схема аналітично-графічного методу визначення оптимальних параметрів віброзахисної системи

Далі було проведено дослідження залежності ефективності застосування цифрового двійника віброзахисного процесу з оптимально налаштованими параметрами. Дослідження показало, що метод дозволяє зменшити максимальну амплітуду коливань, а, отже, і руйнівну дію, в середньому у 3–4 рази, а у найкращому майже у 20 разів. Окрім того, було виявлено, що точність пошуку вихідних параметрів методів суттєво не впливає на ефективність його застосування (див. рис. 4), що дозволяє зменшити час пошуку.

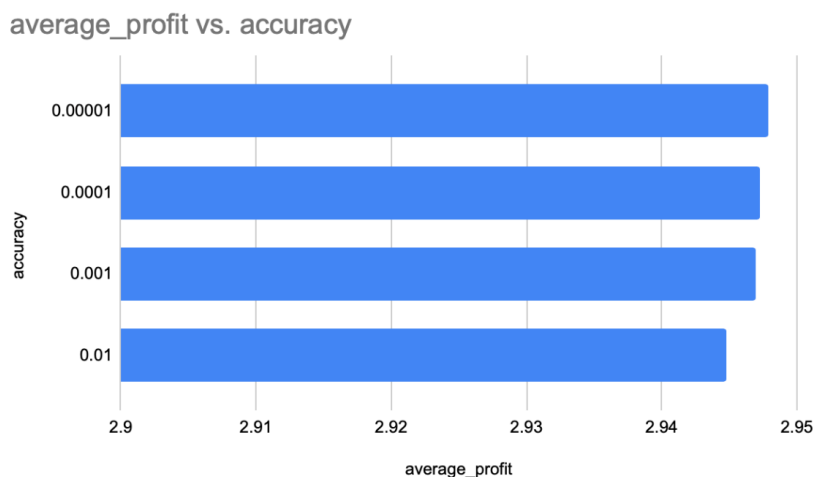


Рисунок 4 - Діаграма залежності ефективності методу від заданої точності

Далі були виконані емпіричні заміри швидкодії оригінального методу та його модифікацій з використанням аналітично-графічного методу пошуку оптимальних параметрів та деяких чисельних методів (див. рис. 5).

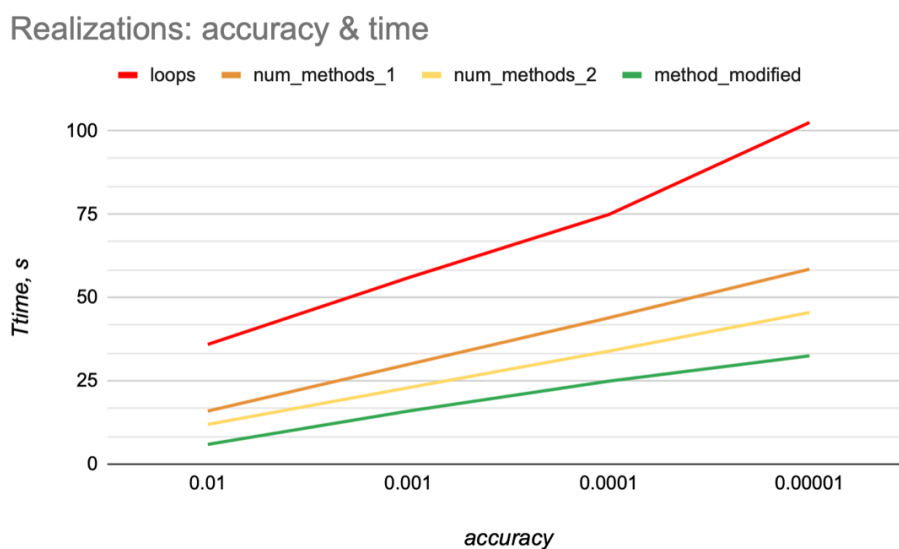


Рисунок 5 - Графік залежності часу побудови цифрового двійника реалізацій в залежності від заданої точності

Аналітично-графічний метод був реалізований за допомогою бібліотеки matplotlib, а чисельні методи – sympy.

Як бачимо, запропонований аналітично-графічний метод є найбільш швидким, адже дозволяє зменшити час виклику в середньому втричі, а також зменшити навантаження на процесор та обробляти випадки, коли знаходження оптимальних параметрів цифрового двійника є доволі складною аналітичною задачею.

Висновки. Під час даного дослідження було розроблено алгоритмічний метод зі створення оптимального цифрового двійника, що відображає процес придушення небезпечних вимушених коливань висотних гнучких споруд на основі використання кульових гасників. При умові малих вимушених коливань отримано рівняння АЧХ лінійної віброзахисної системи. Було описано та проаналізовано вхідні та вихідні параметри методу, способи їх отримання.

На основі числового аналізу встановлено, що форма кривої АЧХ суттєво залежить від спеціального налаштування регулюючих параметрів кульового гасника. При його оптимальних параметрах крива АЧХ віброзахисної системи набуває симетричного вигляду з двома рівними за амплітудою максимумами на відповідних частотах ω_1 та ω_2 , що вказує на рівний розподіл вхідної енергії зовнішнього збурення між робочим тілом гасника і несучим тілом. Порівняльний числовий аналіз показав високу ефективність функціонування такої віброзахисної системи з оптимально налаштованими параметрами кульового гасника.

Окрім цього, були виконані дослідження швидкодії методу за допомогою програмної реалізації цифрового двійника на мові Python. В результаті оцінки швидкості методу була запропонована модифікація методу, що дозволила скоротити час пошуку оптимальних параметрів цифрового двійника. Порівняльний аналіз показав високу ефективність функціонування запропонованої віброзахисної системи з оптимально налаштованими параметрами цифрового двійника.

Подальші дослідження будуть спрямовані на оцінку точності знаходження оптимальних параметрів цифрового двійника та оцінку швидкості його роботи. Окрім того, будуть проаналізовані інші варіанти побудови амплітудно-частотної характеристики, зокрема у нелінійній постановці, та відповідно умови й можливість їх застосування для створення цифрового двійника віброзахисного процесу.

ЛІТЕРАТУРА

1. Review and comparison of the methods of designing the Digital Twin / [Dmytro Adamenko, Steffen Kunnen, Robin Pluhnu, André Loibl, Arun Nagarajah] // Procedia CIRP.–2020.–№91.–P.27–32
2. A probabilistic graphical model foundation for enabling predictive digital twins at scale / [Kapteyn Michael, Pretorius Jacob, Willcox Karen] // Nature Computational Science .–2021.–№1.–P.337–347
3. Легеза В.П. Теория виброзащиты систем с применением изохронных катковых гасителей: модели, методы, динамический анализ, технические решения. Lambert Academic Publishing, Saarbrücken, Deutschland. 2013. – 108 с.
4. Dynamics of vibroprotective systems with roller dampers of low-frequency vibrations / [V. P. Legeza] // Strength of Materials 36. – 2006. – P.185–194
5. Non-linear Model of a Ball Vibration Absorber / [Náprstek Jiří, Fischer Cyril, Pirner, Mark, Fischer Ondřej] // Computational Methods in Earthquake Engineering. –2013.–P.381–396
6. В.П.Легеза, О.В.Атаманюк. Ампліудно-частотна характеристика віброзахисної системи та метод визначення оптимальних параметрів її кульового гасника //Доповідь на 14 науковій конференції ПМК-2021,КПІ, «Політехніка»,2021,С.4-8.
7. Den Hartog, J. P. Mechanical Vibrations. Courier Corporation, 2013. – 464 p.
8. W.Weaver Jr., S.P.Timoshenko, D.H.Young. Vibration Problems in Engineering, 5th Edition, Wiley, 1991. - 624 p
9. Theoretical study and experimental verification of vibration control of offshore wind turbines by a ball vibration absorber (BVA) / [Z.-L.Zhang, J.-Bing Chen, Jie Li.] // Structure and Infrastructure Engineering, V.10, #8, 2014, P.P. 1087-1100

REFERENCES

1. Review and comparison of the methods of designing the Digital Twin / [Dmytro Adamenko, Steffen Kunnen, Robin Pluhnu, André Loibl, Arun Nagarajah] // Procedia CIRP.–2020.–№91.–P.27–32
2. A probabilistic graphical model foundation for enabling predictive digital twins at scale / [Kapteyn Michael, Pretorius Jacob, Willcox Karen] // Nature Computational Science .–2021.–№1.–P.337–347
3. V. P. Legeza, Theory of vibration protection of systems using isochronous roller dampers: models, methods, dynamic analysis, technical solutions. Lambert Academic Publishing, Saarbrücken, Deutschland. 2013. – 108 с.
4. Dynamics of vibroprotective systems with roller dampers of low-frequency vibrations / [V. P. Legeza] // Strength of Materials 36. – 2006. – P.185–194

5. Non-linear Model of a Ball Vibration Absorber / [Náprstek Jiří, Fischer Cyril, Pirner, Mark, Fischer Ondřej] // Computational Methods in Earthquake Engineering. –2013.–P.381–396
6. V. P. Legeza, O. V. Atamaniuk, Amplitude-frequency characteristics of the anti-vibration system and the method of determining the optimal parameters of its ball extinguisher // Report at the 14th scientific conference AMC-2021,KPI, 2021, P. 4 - 8.
7. Den Hartog, J. P. Mechanical Vibrations. Courier Corporation, 2013.– 464 p.
8. W.Weaver Jr., S.P.Timoshenko, D.H.Young. Vibration Problems in Engineering, 5th Edition, Wiley, 1991. - 624 p
9. Theoretical study and experimental verification of vibration control of offshore wind turbines by a ball vibration absorber (BVA) /[Z.-L.Zhang, J.-Bing Chen, Jie Li.] // Structure and Infrastructure Engineering, V.10, #8, 2014, P.P. 1087-1100

Received 12.05.2023.

Accepted 15.05.2023.

Method of creation a digital twin of a vibration protection process

Various approaches to building digital twins are considered. The data-based approach has a big disadvantage due to need of the huge amount of information. The system-based approach can not be used in some cases due to the lack of a mathematically justified method. One of such cases is a ball vibration absorber but they can be really useful for the vibration protection of high-rise flexible objects.

The purpose of the research is to develop an algorithmic method of creating digital twins of the vibration protection process, which will provide the possibility of determining the optimal control parameters of the ball vibration absorber.

The paper examines small steady oscillations of the dynamic system "supporting body - ball vibration absorber". Under the condition of small forced oscillations, the equation of the amplitude-frequency characteristic of the linear anti-vibration system was obtained. In view of the use in construction, the input and output parameters of the method of building a digital twin of a flexible structure were described and analyzed, as well as the methods of obtaining them. As a result of the evaluation of the speed of the method, a modification of the search way for the optimal parameters of the digital twin was proposed.

The comparative analysis showed the high efficiency of the proposed anti-vibration system with optimally adjusted parameters of the digital twin. The proposed method allows to reduce the maximum value of the amplitude by approximately four times. Modifications of the method made it possible to speed it up by an average of three times, reduce the load on the processor and handle cases when finding the optimal parameters of a dig-

ital twin is a rather difficult analytical problem.

The input and output parameters of the method and ways of obtaining them were described and analyzed. A comparative numerical analysis showed the high efficiency of the functioning of such a vibration protection system with optimally adjusted parameters of the ball vibration absorber.

Атаманюк Олексій Віталійович – магістрант 2 курсу кафедри програмного забезпечення комп'ютерних систем, Національний технічний університет України «Київський політехнічний інститут імені Ігоря Сікорського», Київ.

Oleksii Atamaniuk – 2nd year master's student of Computer Systems Software Department, National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute", Kyiv.

Легеза Віктор Петрович – докт. техн. наук, професор, професор кафедри програмного забезпечення комп'ютерних систем. Національний технічний університет України «Київський політехнічний інститут імені Ігоря Сікорського», Київ.

Viktor Legeza – Doc. Sc., Professor, Professor of Computer Systems Software Department. National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute", Kyiv.

М.О. Бубенко, В.В. Герасимов, Н.В. Карпенко, О.С. Морозов
АНАЛІЗ МЕТОДІВ ТЕСТУВАННЯ ВЕБ-ДОДАТКІВ

Анотація. У статті розглядається тестування веб додатків через призму власного досвіду роботи в ІТ компанії. Показано, що під час тестування додатків різних видів слід звертати увагу на їх особливості та на критичні зони, які повинні бути перевірені. Авторами зроблено порівняння тестування web додатків та desktop програм за різними параметрами.

Ключові слова: веб додаток, тестування, критичні зони, модель тестування за сценарієм.

Постановка проблеми. Важливим практичним завданням, що стоїть перед розробниками, є швидке створення та супровід якісного багаторівневого програмного забезпечення (ПЗ). При цьому передбачається, що розроблений продукт відповідатиме характеристикам якості, таким як стійкість, тестованість, корисність, доступність, масштабованість, відкритість, гнучкість. Це накладає на процес розробки додаткові обмеження та правила, а саме: документування розробки на різних рівнях; покриття тестами компонентів, модулів, підсистем; управління проектами, процесами тощо [1]. І, якщо казати про тестування додатків різних видів, то слід звертати увагу на їх особливості, наприклад, веб-додатки мають критичні зони, які повинні бути перевірені. Таким чином, метою даної статті є аналіз різноманітних методів тестування веб-додатків.

Аналіз останніх досліджень і публікацій. Мануальне та автоматизоване тестування ПЗ мають спільну теорію тестування. Досить корисною в плані ознайомлення з термінами, які використовуються під час тестування, є стаття [2], в якій також визначені цілі тестування. Але ця стаття не відповідає на такі питання: «Що потрібно тестувати?», «Яким чином потрібно виконувати дії з тестування?», «Що знадобиться для того, щоб виконати тестування ПЗ?» та інші.

Не дивлячись на те, що автоматизоване тестування широко застосовується в ІТ-компаніях під час розробки різноманітного ПЗ, існує досить мало статей, які б показували практичні аспекти такого тестування. Однією з найбільш цікавих статей в цьому напрямку є [3], де автори порівняли

декілька поколінь популярного інструменту автоматизації тестування веб-додатків Selenium, проаналізували його функціональні особливості, можливості та зручність роботи. Автори статті [4] описали деякі популярні засоби для тестування API, а також надали рекомендації щодо областей їх застосування. Стаття [5] повністю присвячена огляду інструментів автоматизованого тестування ПЗ, таких як Selenium, IBM Rational, SilkPerformer, TestComplete, HP QuickTest Professional, JUnit та їх порівнянню за ефективністю використання у різних додатках.

Практичні поради щодо тестування веб-додатків можна знайти в [6], де автор надає чек-лист з пунктами, які потрібно перевірити під час функціонального, інтеграційного, крос-платформного тестування, а також тестування безпеки, зручності використання, локалізації та глобалізації. Автор [7] пропонує універсальну схему для тестування веб-додатків, яка доповнює чек-лист попереднього автора та надає поради з тестування Performance та Configuration. Автор статті [8] звертає увагу на те, що інформативними є Ad-hock та негативне тестування, а також написання еквівалентних тестів і проведення інтуїтивного тестування (Exploratory testing). Окрім цього, важливим кроком є виконання регресійного тестування, що не було акцентовано в попередніх статтях.

Всі вищезазначені джерела, окрім [2], стосуються практики тестування веб-додатків, однак слід згадати й ті статті, з яких почався розвиток трьох теорій тестування: Goodenough-Gerhart [9], Weyuker-Ostrand [10] та Gourlay [11]. Ці роботи дозволяють зрозуміти основні проблеми тестування та шляхи їх вирішення.

Основна частина. Для того, щоб виявити ті фактори, які відрізняють тестування web-додатків від тестування desktop-додатків, складено таблицю з їх порівнянням за різними критеріями (табл.1). Це порівняння показує, що веб-додатки мають кілька особливостей, які потрібно враховувати під час їх тестування. Так, навколишнє середовище висуває вимогу щодо перевірки на правильність відображення інтерфейсів веб-додатку в різних браузерах. Під правильністю слід розуміти відповідність макетам інтерфейсів, які були попередньо розроблені дизайнерами та погоджені замовником. Таким чином, при кросбраузерному тестуванні перевіряється правильність відображення елементів графіки, шрифтів, розмір текстових символів; доступність елементів керування, чи є вони інтерактивними тощо.

Порівняння web-додатків та desktop-програм

	Web-додаток	Desktop-додаток
Навколишнє середовище	Додаток розгорнутий на сервері, до якого інші комп'ютери можуть отримати доступ лише за допомогою браузера.	Додаток обмежений комп'ютером, на якому він розгорнутий.
Платформа	Не залежать від платформи, можуть працювати на кількох платформах без окремої розробки.	Для різних платформ потрібно розробляти окремі настільні програми.
Розгортання / Оновлення	Розгортання або будь-яке оновлення / виправлення потрібно виконувати лише на одному наборі серверних машин.	Розгортання або будь-яке оновлення коду потрібно виконувати на всіх клієнтських машинах.
Архітектура	Клієнтська частина зазвичай базується на веб-технологіях та працює в браузері, використовується протокол HTTP для передачі даних. Серверна частина може бути розподілена на багато серверів.	Програма встановлюється на комп'ютер користувача і використовує механізми операційної системи для комунікації з іншими процесами або мережевими сервісами.
Підключення	Веб-програми залежать від підключення до Інтернету для їх функціонування.	Настільні програми не завжди вимагають підключення до Інтернету.
Доступність	Веб-програми не мають жодних обмежень щодо розташування, оскільки вони розгорнуті на серверах, до яких можна отримати доступ із будь-якого місця.	Настільні програми обмежені фізичними розташуваннями, тому до них можна отримати доступ лише з машин, на яких вони за своєю суттю розгорнуті.

Веб-додаток за допомогою різних елементів керування надає користувачу можливість виконувати дії, які визначені в функціональних вимогах. Таким чином, слід перевірити, чи всі функції реалізовані в програмному продукті, а

також коректність їх роботи. Під час перевірки функціоналу веб-додатку, слід звернути особливу увагу на форми для заповнення, через які здійснюється взаємодія клієнта з сервером, а саме на таке:

— Обов'язкові поля для заповнення: чи зазначено, що вони є обов'язковими? Чи обов'язкові поля насправді є обов'язковими? Що буде, якщо поле залишити пустим?

— Введення в поле спецсимволів #, %, @, №, &, ^, \$, *, <, >, { }, [] залежить від вимог – є поля, де вони потрібні, є поля, де їх потрібно заборонити. Наприклад, в полі з ім'ям, прізвищем та по-батькові спецсимволів не повинно бути, в полі email повинен бути один обов'язковий символ "@". Поля пароля можуть включати введення, як мінімум, одного будь-якого спеціального символу тощо.

— Наявність валідації: що відбувається, коли користувач вводить невалідні значення? Чи отримує користувач повідомлення про помилку? Що відбувається з даними після введення?

Окрім форм для заповнення є й інші елементи керування, такі як радіокнопки, випадаючі списки, прапорці, які також слід перевірити. Оскільки таких елементів може бути досить багато і, відповідно, випадків для перевірки, то рекомендується складати чек-лист UI-компонентів, які потрібно перевірити.

При тестуванні веб-додатків не можна забувати про базу даних (БД). Отже, потрібно перевіряти:

- чи збереглися в БД дані, введені з боку клієнта;
- чи коректно відображаються на стороні клієнта дані, записані в БД;
- чи можна видалити / змінити дані;
- кількість відкритих з'єднань з БД;
- швидкість обробки запитів.

Одночасно з веб-програмою може працювати велика кількість користувачів. Тому потрібно:

— проводити навантажувальне тестування для з'ясування того, скільки користувачів можуть без проблем працювати одночасно;

— перевіряти аутентифікацію та авторизацію — на сервері перевіряються ідентифікаційні дані користувача та його доступ до ресурсів веб-додатка. Після успішної аутентифікації сервер зберігає дані про користувача та його доступ, що використовуються для подальшої авторизації;

— тестувати конкурентний (паралельний) доступ до ресурсів — швидкість роботи додатків не повинна змінюватися зі збільшенням кількості користувачів;

— рівні доступу користувачів.

Виходячи з вищесказаного, можна виділити такі основні **критичні зони** веб-додатків, які потребують додаткового опису та інструкцій:

- єдність дизайну;
- навігація та «дружність» до користувача;
- функціональність;
- безпека;
- сумісність з браузерами та операційними системами;
- продуктивність.

Єдність дизайну веб-ресурсу є запорукою того, що ресурс справить враження цілісної структури, яка є гармонійною та логічною, тому перевірка єдності потребує чіткої документації, унікальної для кожного окремого ресурсу, проте не вимагає значних затрат ресурсів та часу.

Тестування продукту на навігацію та «дружність» допомагає спрямувати зусилля на покращення першого враження користувача про продукт та перетворити випадкову аудиторію на прихильників. Тому таке тестування є доволі важливим, при цьому не потребує багато ресурсів, але бажана наявність чітких рекомендацій чи втручання спеціаліста з предметної області.

У переважній кількості розроблюваних додатків функціональність є ключовою зоною уваги як при тестуванні, так і під час користування кінцевим користувачем. Отже, функціональне тестування вимагає ресурсів, досвіду від працівників, чітких вимог та рекомендацій з точки зору спеціаліста у відповідній сфері. Позитивним моментом є те, що функціональні тестові сценарії добре підлягають автоматизації, а це скорочує час перевірки та ресурси.

Перевірка на безпеку є особливо критичним аспектом саме у тестуванні веб-додатків, оскільки інтерфейс кожного користувача безпосередньо взаємодіє з єдиним сервером, тому за певних умов та зусиль секретна інформація може опинитися в руках людини, що не повинна мати доступ до неї. Це може призвести до негативних наслідків та збитків як для представника послуг, так і для інших користувачів. Часто таке тестування може ефективно виконуватися командою з 1-3 людей з досвідом та кваліфікацією у відповідній сфері і не залежить від специфіки певного веб-ресурсу.

Сумісність з браузерами та операційними системами є принципово важливою для продуктів, спрямованих на велику аудиторію (наприклад Інтернет магазини, соціальні мережі). У час розвинених технологій та широкої різноманітності пристроїв, що мають доступ до мережі Інтернет, кожен наявний пристрій додає перевірок. Зазвичай фізично неможливо, а за можливості вкрай не рентабельно проводити тестування на всіх існуючих пристроях, а тому, враховуючи аудиторію користувача та популярність пристроїв, для перевірок обираються найпоширеніші з них. Тестування сумісності є ресурсозатратним незалежно від специфіки продукту.

Перевірка на продуктивність є ключовою для ресурсів, що розробляються для великої кількості користувачів, які користуються веб-сервісом одночасно. Оскільки загалом такі тести виконуються за допомогою емуляторів та спеціальних програм імітаторів навантаження, то така перевірка є автоматизованою, але вимагає високого рівня практики та знань від виконавця.

Здатність веб-додатків обробляти і поширювати інформацію через Інтернет робить їх уразливішими, порівняно з desktop-додатками. Веб-додатки повинні забезпечувати публічний доступ до даних, а це значно підвищує вимоги до стійкості цього ПЗ. По-друге, що важливіше з точки зору тестування на проникнення, вони обробляють дані, отримані із запитів *HTTP*-протоколу — протоколу, який може використовувати безліч способів кодування і інкапсуляції інформації.

Таким чином, можна виділити тести, які зазвичай використовують під час перевірки web та desktop додатків (табл. 2).

Таблиця 2

Тести, які використовують для перевірки web та desktop-додатків

Web-додаток	Desktop-додаток
User Interface Testing, Load, Stress, Volume Browser/OS Compatibility, Security Data, Volume Testing, Cross Browser Testing, Static Page Testing, Broken Link Testing	GUI, Functionality, Load, Backend

Найпопулярніший інструмент для автоматизації тестування — Selenium Web Driver. Найбільш популярною сферою застосування Selenium є автоматизація тестування веб-додатків. Однак, за допомогою Selenium можна автоматизувати будь-які інші рутинні дії, що виконуються через браузер. Розробка Selenium підтримується виробниками майже всіх браузерів. Вони адап-

ISSN 1562-9945 (Print)
ISSN 2707-7977 (Online)

тують браузери для більш тісної інтеграції із Selenium, а іноді навіть реалізують вбудовану підтримку Selenium у браузері.

Висновки. Зроблено порівняння web-додатків та desktop-програм за такими критеріями, як навколишнє середовище, платформа, архітектура, розгортання тощо. Відповідно до цього було наведено рекомендації щодо особливостей тестування web-додатків та виділено їх критичні зони.

Описано специфіку тестування критичних зон та наведено оцінку ресурсозатратності їхнього тестування.

Рекомендовано тести, які доцільно використовувати для тестування web та desktop-додатків.

ЛІТЕРАТУРА

1. Литвинов А.А., Карпенко Н.В. Тестирование информационных систем: модульное, интеграционное, системное: учебное пособие. / А. А. Литвинов, Н. В. Карпенко – Д.: Лира, 2016. – 284 с.
2. Mishchevskii Gennadii. Тестирование. Фундаментальная теория. URL: <https://dou.ua/forums/topic/13389/> (дата звернення 12.02.2023).
3. Герасимов В.В., Беличенко А.Я. Развитие инструмента автоматизированного тестирования – SELENIUM // Системні технології. Регіональний міжвузівський збірник наукових праць. — Вип. 1 (114). — Дніпро, 2018. С. 38–44.
4. Герасимов В. В., Никитин Н. Э., Щербак А. Е., Карпенко Н. В. Тестирование API и актуальные средства его реализации // Системні технології. Регіональний міжвузівський збірник наукових праць. - Вип. 1 (120). — Дніпро, 2019. С. 71–80.
5. Герасимов В. В., Кронфельд М. Ф., Озерова Д. М. Исследование технологий автоматизированного тестирования // Системні технології. Регіональний міжвузівський збірник наукових праць. — Вип. 1(96). 2015. – С. 130-136.
6. Чек-лист тестирования WEB приложений.
URL: <https://habr.com/ru/post/542422/> (дата звернення 12.02.2023).
7. Nani Davitadze. Ничего не забыть: универсальная схема для тестирования веб-приложений. URL: <https://dou.ua/lenta/articles/scheme-for-qa/> (дата звернення 12.02.2023).
8. Тестирование веб-проектов: основные этапы и советы.
URL: <https://qalight.ua/ru/baza-znaniy/testirovanie-veb-proektov-osnovnye-etapy-i-sovety/> (дата звернення 12.02.2023).
9. Goodenough J. B. Toward a theory of test data selection. / J. B. Goodenough, S. L. Gerhart // SIGPLAN Notices. – 1975. – 10 (6). – P. 495-510.
10. E. J. Weyuker , T. J. Ostrand. Theories of Program Testing and the Application of Revealing Subdomains. // IEEE Trans. Software Engrg. SE-6 (3) (1980). – P. 236-246.

11. Gourlay J. S. A mathematical framework for the investigation of testing. / J. S. Gourlay // IEEE Trans. Software Engrg. SE. – 1983. – 9 (21). – P. 686-709.

REFERENCES

1. Lytvynov A.A., Karpenko N.V. Testyrovanye informatsionnykh system: modulnoe, intehtatsionnoe, systemnoe: uchebnoe posobie. / A. A. Lytvynov, N. V. Karpenko – D.: Lyra, 2016. – 284 p.
2. Mishchevskii Gennadii. Testirovanie. Fundamentalnaia teoriya. URL: <https://dou.ua/forums/topic/13389/>
3. Herasymov V. V., Belychenko A. Ya. Razvitie instrumenta avtomatyzyro-vannoho testyrovanyia – SELENIUM // Systemni tekhnolohii. Rehionalnyi mizh-vuzivskiy zbirnyk naukovykh prats. – N1 (114). – Dnipro, 2018. P. 38–44.
4. Herasymov V. V., Nyktytn N. E., Shcherbak A. E., Karpenko N. V. Testyrovanye API i aktualnye sredstva eho realizatsyy // Systemni tekhnolohii. Rehionalnyi mizhvuzivskiy zbirnyk naukovykh prats. – N1 (120). – Dnipro, 2019. P. 71–80.
5. Herasymov V. V., Kronfeld M. F., Ozerova D. M. Issledovanye tekhnolohiy avtomatyzyrovannoho testyrovanyia // Systemni tekhnolohii. Rehionalnyi mizhvuzivskiy zbirnyk naukovykh prats. – N1(96). 2015. P. 130-136.
6. Chek-lyst testyrovanyia WEB prylozheniy. URL: <https://habr.com/ru/post/542422/>
7. Nani Davitadze. Nycheho ne zabyt: unyversalnaia skhema dlia testyrovanyia web-prylozheniy. URL: <https://dou.ua/lenta/articles/scheme-for-qa/>
8. Testyrovanye web-proektov: osnovnye etapy i sovety. URL: <https://qalight.ua/ru/baza-znaniy/testirovanie-veb-proektov-osnovnye-etapy-i-sovety/>
9. Goodenough J. B. Toward a theory of test data selection. / J. B. Goodenough, S. L. Gerhart // SIGPLAN Notices. – 1975. – 10 (6). – P. 495-510.
10. E. J. Weyuker, T. J. Ostrand. Theories of Program Testing and the Application of Revealing Subdomains. // IEEE Trans. Software Engrg. SE-6 (3) (1980). – P. 236-246.
11. Gourlay J. S. A mathematical framework for the investigation of testing. / J. S. Gourlay // IEEE Trans. Software Engrg. SE. – 1983. – 9 (21). – P. 686-709.

Received 16.05.2023.

Accepted 18.05.2023.

Analysis of web application testing methods

An important practical task for developers is the rapid creation and maintenance of high-quality multi-level software. It is assumed that the developed product will meet the quality characteristics. And, if we talk about testing applications of different types, then you should pay attention to their features. For example, web applications have critical areas that must be checked. Thus, the purpose of this article is to analyse various methods and technics for testing web applications.

The article provides a detailed analysis of the latest publications related to testing web applications. It turned out that most of the articles are aimed at describing terms or general information about testing. Several articles describe automated testing with Selenium, IBM Rational, SilkPerformer, TestComplete, HP QuickTest Professional, JUnit and compare them in terms of efficiency in various applications. However, most of the articles are devoted to various aspects of manual testing.

In order to identify the factors that distinguish web application testing from desktop application testing, a table has been compiled comparing them according to the following criteria: environment, platform, deployment and updating, architecture, connectivity, availability. This comparison shows that web applications have several features that need to be considered when testing them.

In our opinion, the main critical areas of web applications that require additional description and instructions are unity of design, navigation and "friendliness" to the user, functionality, security, compatibility with browsers and operating systems, and productivity. The article describes the specifics of testing critical zones and gives an estimate of the resource consumption of their testing. Tests are also recommended, which are useful for testing web and desktop applications.

Бубенко Максим Олександрович - бакалавр 4 курсу кафедри ЕОМ Дніпровського національного університету імені Олеся Гончара.

Карпенко Надія Валеріївна - к.ф.-м.н., доцент кафедри ЕОМ Дніпровського національного університету імені Олеся Гончара.

Герасимов Володимир Володимирович - к.т.н, доцент кафедри ЕОМ Дніпровського національного університету імені Олеся Гончара.

Морозов Олександр Сергійович - асистент кафедри ЕОМ Дніпровського національного університету ім. Олеся Гончара.

Bubenko Maksym - 4th year bachelor's student, Computer Engineering Department, Oles Honchar Dnipro National University.

Karpenko Nadiia - Ph.D in Physics and Mathematical Sciences, Associate Professor, Computer Engineering Department, Oles Honchar Dnipro National University.

Gerasimov Volodymyr - Ph.D. in Technical Sciences, Associate Professor, Computer Engineering Department, Oles Honchar Dnipro National University.

Morozov Alexander Sergeyevich – Assistant, Computer Engineering Department, Oles Honchar Dnipro National University.

МЕТОДИ ПІДВИЩЕННЯ РІВНЯ ЕФЕКТИВНОСТІ РОБОТИ АВТОМАТИЗОВАНИХ СИСТЕМ

Анотація. Автоматизовані системи відіграють важливу роль у сучасних промислових, технологічних та інформаційних сферах, забезпечуючи ефективність, надійність та оптимізацію процесів. Однак, в умовах швидкого розвитку технологій та зміни вимог, підвищення рівня ефективності автоматизованих систем залишається актуальною проблемою. Розглянуто існуючі методології та підходи до підвищення ефективності автоматизованих систем. Представлена порівняльна характеристика методів підвищення ефективності, визначені переваги та недоліки. Враховано такі параметри, як гнучкість, адаптивність, надійність та ефективність. Окремо досліджено простоту впровадження, масштабування та підтримка стандартів. Для сфери застосування наведено приклади конкретних методів та їх вплив на підвищення ефективності автоматизованих систем. Запропоновано новий алгоритм підвищення рівня ефективності, що ґрунтується на комбінації оптимізації процесів, автоматизації та інтеграції систем для досягнення найкращого результату. Пропонуються можливі напрями підвищення ефективності автоматизованих систем.

Ключові слова: автоматизовані системи, ефективність, оптимізація процесів, моніторинг, масштабованість, ітеративність.

Постановка проблеми. Сучасний промисловий та технологічний прогрес нерозривно пов'язаний з використанням автоматизованих систем у різних галузях. Ці системи відіграють ключову роль у забезпеченні ефективності та надійності процесів, збільшенні продуктивності та зниженні операційних витрат. Однак, в умовах швидкого розвитку технологій та збільшення вимог до продуктивності – ефективність автоматизованих систем стає ще більш важливим завданням. [6]

Незважаючи на широке застосування автоматизованих систем, залишається низка суттєвих проблем, пов'язаних з їхньою ефективністю. Перш за все, в умовах складних і виробничих середовищ, що динамічно змінюються, пошук оптимальних методів підвищення ефективності стає складним завданням. Існуючі методи та підходи можуть бути недостатньо адаптивними та гнучкими для ефективного управління автоматизованими системами у різних

умовах. Крім того, важливим аспектом є баланс між підвищенням ефективності автоматизованих систем і збереженням безпеки та надійності їхньої роботи. Помилки та збої в автоматизованих системах можуть мати серйозні наслідки, тому необхідно створити способи, які дозволять покращити ефективність, без втрати надійності системи.

Аналіз останніх досліджень і публікацій. Дослідження, присвячені підвищенню ефективності автоматизованих систем, охоплюють різні галузі, включаючи виробничі підприємства, транспортні мережі, енергетику, охорону здоров'я та інші. У ряді робіт акцент робиться на розробці та оптимізації алгоритмів і моделей, які використовуються в автоматизованих системах.

Підхід, представлений у [3], пов'язаний із оптимізацією процесів обробки даних. В роботі активно досліджуються методи стиснення даних, фільтрації шуму, прискорення обчислень та інші техніки, які створені задля поліпшення продуктивності системи. Такі підходи особливо актуальні в контексті великих обсягів даних, які потребують швидкої та ефективної обробки. Також важливим аспектом є покращення взаємодії між автоматизованими системами та людьми. Робота [4], присвячена інтерфейсам та ергономіці систем, наголошує на важливості розробки інтуїтивно зрозумілих та зручних у використанні інтерфейсів. Крім того, в останні роки спостерігається зростання інтересу до нових технологій, таких як інтернет речей (IoT), розподілені системи та хмарні обчислення. Дослідження у цій галузі [7] показують, що впровадження таких технологій може значно покращити роботу автоматизованих систем. Разом з тим слід зазначити, що питання надійності автоматизованих систем також привертають велику увагу. У роботі [6] розглядаються методи та підходи до забезпечення захисту систем від кібератак, а також методи резервування та резервного копіювання даних для запобігання втраті інформації. Аналізуючи літературу, можна виділити загальні тенденції, такі як використання оптимізацію процесів обробки даних, покращення інтерфейсів та впровадження нових технологій. Однак, незважаючи на значні досягнення, все ще залишається багато напрямків для подальших досліджень та розробок.

Мета статті: дослідження та систематизація різних існуючих методів та підходів, а також пропозиція нових алгоритмів та рішень для оптимізації роботи автоматизованих систем.

Викладення основного матеріалу дослідження. Підвищення рівня ефективності автоматизованих систем є важливим завданням у різних галузях, таких як промисловість, транспорт, охорона здоров'я, інформаційні технології

та ін. Це процес оптимізації роботи системи з метою досягнення максимального рівня продуктивності, зниження витрат, підвищення якості, скорочення часу виконання завдань та покращення загальної ефективності [2].

Існує безліч методів та підходів, які можуть бути застосовані для підвищення ефективності автоматизованих систем. Проаналізуємо основні із них.

Оптимізація процесів. Аналіз та покращення робочих процесів дозволяють ідентифікувати вузькі місця, надлишкові операції та неефективні кроки в системі. Шляхом оптимізації процесів можна покращити продуктивність, скоротити час виконання завдань та збільшити пропускну спроможність системи.

Автоматизація та роботизація. Впровадження автоматизованих систем та роботів може прискорити виконання завдань, знизити ймовірність помилок та підвищити загальну ефективність роботи. Автоматизація рутинних і повторюваних завдань звільняє людський ресурс від складних завдань.

Інтеграція та синхронізація систем. Об'єднання різних автоматизованих систем та їх взаємодія може підвищити загальну ефективність роботи. Наприклад, інтеграція систем управління виробництвом, складськими системами та системами доставки дозволяє більш ефективно планувати та координувати процеси.

Контроль та моніторинг. Впровадження систем контролю та моніторингу дозволяє отримувати в реальному часі дані про роботу автоматизованих систем, виявляти проблеми та несправності.

Кожна автоматизована система має свої особливості та потребує індивідуального підходу до підвищення її ефективності. Також слід враховувати конкретні цілі, щоб вибрати найбільш підходящі методи.

Методології підвищення ефективності роботи автоматизованих систем – набір підходів, стратегій і методів, які допомагають оптимізувати процеси і досягати оптимальної продуктивності системи:

Lean-підхід. Lean-підхід, спочатку розроблений для виробничих систем, може бути застосований і в галузі автоматизації. Він прагне елімінування втрат, оптимізації процесів та підвищення якості продукції чи послуг. Принципи Lean включають усунення надлишкової роботи, скорочення часу простою та підвищення гнучкості системи.

Six Sigma. Методологія Six Sigma фокусується на зменшенні варіації та дефектів у процесах. Вона використовує статистичний аналіз та методи покращення якості для досягнення високого ступеня точності та мінімізації поми-

лок. Застосування Six Sigma в автоматизованих системах дозволяє знизити кількість дефектів та підвищити ефективність.

Agile-методологія. Agile-методологія, яка широко використовується в розробці програмного забезпечення, може бути також застосована в автоматизованих системах. Вона заснована на ітеративному та інкрементальному підході до розробки, дозволяючи швидко адаптуватися до вимог, що змінюються, і скорочувати час циклу розробки. Agile-методологія сприяє більш гнучкому та ефективному управлінню проектами автоматизації.

Континуальне покращення процесів. Континуальне покращення процесів (також відоме як Kaizen) орієнтоване на постійне вдосконалення системи. Ця методологія передбачає регулярний аналіз процесів, виявлення проблемних областей та впровадження покращень. Континуальне покращення процесів сприяє поступовому підвищенню ефективності використання ресурсів.

Використання технологій та алгоритмів оптимізації. Для підвищення ефективності автоматизованих систем часто застосовуються технології та алгоритми оптимізації. Це може містити методи лінійного програмування, генетичні алгоритми, еволюційні стратегії та інші підходи, спрямовані на знаходження оптимальних рішень у складних системах [5].

Представлені тільки деякі з методологій, які можуть бути використані для підвищення ефективності роботи автоматизованих систем. Вибір методології залежить від конкретних цілей, контексту та характеристик системи. Важливо підбирати відповідні методи та налаштовувати їх відповідно до вимог та потреб проекту.

В таблиці 1 представлено порівняння методологій підвищення рівня ефективності автоматизованих систем. Кожна методологія оцінюється з погляду складності реалізації, інтеграції з існуючими системами, застосування до різних масштабів і типів систем та залучення персоналу

Створимо алгоритм підвищення рівня ефективності роботи автоматизованих систем – Алгоритм 3Е (Ефективність, Еволюція, Емпатія).

Крок 1. Аналіз та визначення ключових метрик ефективності

Ідентифікуються основні параметри, що впливають на ефективність системи. Визначаються ключові метрики, які можна використовувати для вимірювання ефективності роботи системи, таких як продуктивність, час виконання завдань, рівень автоматизації та інші.

Аналіз методологій ефективності автоматизованих систем

Методологія	Lean-підхід	Six Sigma	Agile-методологія	Континуальне вдосконалення	Технології оптимізації
Переваги	Усунення втрат та оптимізація процесів	Зниження варіації та дефектів	Гнучкість в адаптації до змін	Постійне вдосконалення процесів	Знаходження оптимальних рішень у складних системах
Недоліки	Потрібна культурна зміна в організації	Потрібен спеціаліст з навчання з Six Sigma	Чи можуть виникати труднощі в управлінні великими проектами	Потрібна постійна участь та підтримка	Потрібна висока експертиза застосування алгоритмів
Сфера використання	Виробництво, логістика та обслуговування	Виробництво, охорона здоров'я	Розробка програмного забезпечення	Будь-яка галузь, де потрібне постійне поліпшення	Виробництво, логістика та транспорт
Складність реалізації	Середня	Висока	Середня	Середня	Висока
Інтеграція з існуючими системами	Середня	Середня	Середня	Середня	Середня
Залучення персоналу	Висока	Висока	Середня	Середня	Середня

Крок 2. Оптимізація процесів та автоматизація.

Вивчення поточних процесів і виявлення вузьких місць та потенційних проблем. Застосування методів оптимізації процесів, такі як Lean-підхід або

Theory of Constraints, для усунення втрат та покращення ефективності. Автоматизація повторюваних завдань за допомогою технологій автоматизації, таких як роботизований процес автоматизації (RPA) або штучний інтелект (II).

Крок 3. Впровадження інновацій та еволюція.

Дослідження і впровадження нових технологій та інноваційних рішень, які можуть підвищити ефективність системи. Постійне відстеження змін та адаптація системи до нових вимог.

Крок 4. Облік емпатії та людського фактору.

Приділення уваги взаємодії з користувачами та співробітниками системи. Створення інтерфейсу користувача, який враховує потреби користувачів.

Крок 5. Вимірювання та аналіз результатів.

Регулярне вимірювання та аналіз метрики ефективності, визначені на першому кроці. Порівняння результати з початковими показниками та встановлення цілей для подальшого покращення. Використання отриманих даних для прийняття рішень та планування майбутніх кроків.

Рівень складності реалізації – Середній.

Інтеграція з існуючими системами – Висока.

Сфера використання алгоритму 3E – Алгоритм може бути застосований у різних галузях, включаючи виробництво, логістику, фінанси та інші, де автоматизовані системи відіграють важливу роль.

Переваги алгоритму 3E. Фокус на оптимізації процесів та автоматизації дозволяє покращити ефективність роботи системи та скоротити час виконання завдань. Впровадження інновацій та еволюція системи забезпечують адаптацію до змін та можливостей галузі. Регулярний аналіз результатів дозволяє оцінити ефективність алгоритму та приймати обґрунтовані рішення для подальшого покращення.

Недоліки алгоритму 3E. Впровадження нових технологій та інновацій може вимагати додаткових ресурсів та інвестицій. Залучення та навчання співробітників може бути тимчасово скрутним та вимагати додаткових зусиль.

Алгоритм 3E є комплексним підходом до підвищення ефективності роботи автоматизованих систем, враховуючи оптимізацію процесів, впровадження інновацій, облік емпатії та регулярне вимірювання результатів. Його застосування може призвести до покращення продуктивності, зниження витрат та підвищення задоволеності користувачів.

Висновки. Створений алгоритм 3E є комплексний підхід до підвищення ефективності роботи автоматизованих систем. Він поєднує кілька методологій

та підходів, таких як оптимізація процесів, впровадження інновацій, облік емпатії та регулярний аналіз результатів. Застосування цього алгоритму може призвести до значного покращення продуктивності, скорочення часу виконання завдань та зниження витрат.

Алгоритм починається з аналізу та визначення ключових метрик ефективності, що дозволяє фокусуватися на найважливіших параметрах. Потім слідує оптимізація процесів і автоматизація, які дозволяють усунути вузькі місця і підвищити ефективність системи. Впровадження інновацій та еволюція системи забезпечують її адаптацію до змін та можливостей галузі. Аналіз результатів допомагає оцінити ефективність алгоритму та приймати обґрунтовані рішення для подальшого покращення. Однак, слід зазначити, що реалізація алгоритму може вимагати додаткових ресурсів та інвестицій. Дотримуючись даного способу, організації зможуть досягти оптимального використання своїх автоматизованих систем і залишатися конкурентоспроможними в бізнес-середовищі.

ЛІТЕРАТУРА

1. Akhramovich V. Method of calculating the protection of personal data from the reputation of users / Akhramovich V. Hrebennikov A., Tsarenko B., Stefurak O. // Sciences of Europe, Praha, Czech Republic. – 2021. – VOL 1, No 80 (2021) P. 23-31.
2. Biyashev R.G. Algorithm for creation a digital signature with error detection and correction / Biyashev R.G., Nyssanbayeva S.E. // Cybernetics and Systems Analysis. – 2022. – P.489–495.
3. Dongyan NIE. Compression of surface texture acceleration signal based on spectrum characteristics./ Dongyan NIE, Xiaoying SUN. // Virtual Reality & Intelligent Hardware. – 2023. № 5(2). – P. 110–123
4. Hankiewicz Krzysztof. Ergonomic Context of the User Interface of Modern Enterprise Management Systems // EUROPEAN RESEARCH STUDIES JOURNAL. XXIV. – 2021. – P. 140-151. DOI:10.35808/ersj/2708.
5. Hernandez R. Kharitonov's theorem extension to interval polynomials which can drop in degree: a Nyquist approach / R. Hernandez, S. Dormido // IEEE Transactions on Automatic Control. – Vol. 41. - Issue 7. - 1996. - P. 1009 – 1012.
6. Matthew Peacock. An analysis of security issues in building automation systems / Matthew Peacock, Michael N. Johnstone //12th Australian Information Security Management Conference. Held on the 1-3 December, 2014 at Edith Cowan University, Joondalup Campus, Perth, Western Australia. This Conference Proceeding is posted at Research Online. – 2014. – P.100-104. DOI: 10.4225/75/57b691dfd9386

7. Sfar AR. A systematic and cognitive vision for IoT security: a case study of military live simulation and security challenges. / Sfar AR, Zied C, Challal Y. // International conference on smart, monitored and controlled cities (SM2C). – 2017. – P. 235-246. DOI: 10.1109/SM2C.2017.8071828 333
8. Siegwart R. Introduction to autonomous mobile robots / R. Siegwart, I. R. Nourbakhsh, D. Scaramuzza // Massachusetts Institute of Technology. 2004. – 321 p

REFERENCES

1. Akhramovich V. Method of calculating the protection of personal data from the reputation of users / Akhramovich V. Hrebennikov A., Tsarenko B., Stefurak O. // Sciences of Europe, Praha, Czech Republic. – 2021. – VOL 1, No 80 (2021) P. 23-31.
2. Biyashev R.G. Algorithm for creation a digital signature with error detection and correction / Biyashev R.G., Nyssanbayeva S.E. // Cybernetics and Systems Analysis. – 2022. – P.489–495.
3. Dongyan NIE. Compression of surface texture acceleration signal based on spectrum characteristics./ Dongyan NIE, Xiaoying SUN. // Virtual Reality & Intelligent Hardware. – 2023. № 5(2). – P. 110–123
4. Hankiewicz Krzysztof. Ergonomic Context of the User Interface of Modern Enterprise Management Systems // EUROPEAN RESEARCH STUDIES JOURNAL. XXIV. – 2021. – P. 140-151. DOI:10.35808/ersj/2708.
5. Hernandez R. Kharitonov's theorem extension to interval polynomials which can drop in degree: a Nyquist approach / R. Hernandez, S. Dormido // IEEE Transactions on Automatic Control. – Vol. 41. - Issue 7. - 1996. - P. 1009 – 1012.
6. Matthew Peacock. An analysis of security issues in building automation systems / Matthew Peacock, Michael N. Johnstone //12th Australian Information Security Management Conference. Held on the 1-3 December, 2014 at Edith Cowan University, Joondalup Campus, Perth, Western Australia. This Conference Proceeding is posted at Research Online. – 2014. – P.100-104. DOI: 10.4225/75/57b691dfd9386
7. Sfar AR. A systematic and cognitive vision for IoT security: a case study of military live simulation and security challenges. / Sfar AR, Zied C, Challal Y. // International conference on smart, monitored and controlled cities (SM2C). – 2017. – P. 235-246. DOI: 10.1109/SM2C.2017.8071828 333
8. Siegwart R. Introduction to autonomous mobile robots / R. Siegwart, I. R. Nourbakhsh, D. Scaramuzza // Massachusetts Institute of Technology. 2004. – 321 p.

Received 12.05.2023.
Accepted 19.05.2023.

Methods of increasing the level efficiency of automated systems

Automated systems play a key role in the modern world, ensuring efficiency and automation of various processes. However, with the constant development of technology and the increasing complexity of tasks, continuous improvement and efficiency of these systems is required. This article explores methods that can improve the efficiency of automated systems. Various aspects are analyzed, such as optimization of work, improvement of productivity, reduction of task execution time, reduction of errors, and increase of accuracy. The main goal of the article is to focus on the methodologies for increasing the level of efficiency. The table shows the methodologies with a description of their advantages, disadvantages, and areas of application. In addition, additional parameters such as the degree of automation, the degree of system flexibility, and the level of autonomy are proposed. The article also proposes a new algorithm for improving the efficiency of automated systems. The algorithm is based on the use of modern technologies and approaches, such as data analysis and process optimization. The proposed algorithm has the potential to improve the efficiency of automated systems and can be adapted many times over. The research represents a significant contribution to the field of improving the efficiency of automated systems. The algorithm can be useful for researchers, engineers, automation professionals, and managers interested in improving and optimizing their systems.

Keywords: automated systems, efficiency, process optimization, monitoring, scalability, iterability.

Тулуб Валентин Олександрович – аспірант кафедри робототехніки та спеціалізованих комп'ютерних систем Черкаського державного технологічного університету.

Tulub Valentyn – post-graduate student at the Department of Robotics and Specialized Computer Systems at Cherkasy State Technological University.

АВТОМАТИЗОВАНІ МОДЕЛІ ОБРОБКИ ВІЗУАЛЬНОЇ ІНФОРМАЦІЇ

Анотація. У контексті швидкого розвитку комп'ютерного зору та штучного інтелекту обробка візуальних даних відіграє ключову роль у різних галузях, включаючи медицину, робототехніку, безпеку та інші. В рамках огляду літератури проаналізовано існуючі методи та алгоритми обробки візуальної інформації. Результати аналізу дозволили виявити переваги та недоліки, а також визначити сфери їх застосування. Зокрема, розглянуто сегментацію зображень, розпізнавання образів, класифікацію і детекцію об'єктів та інші аспекти обробки візуальної інформації. Основна мета дослідження полягає у створенні комплексної моделі, здатної автоматично обробляти та аналізувати різні форми візуальних даних. Модель пока сала високу ефективність та точність в обробці візуальних даних, включаючи завдання сегментації, розпізнавання та класифікації об'єктів. Результати дослідження підтверджують ефективність та значущість автоматизованих моделей обробки візуальної інформації. За пропонується модель може бути корисною для розробки нових інформаційних систем, які базуються на обробці візуальних даних та сприяти розвитку комп'ютерного зору та штучного інтелекту.

Ключові слова: обробка зображень, розпізнавання образів, детектування об'єктів, машинне навчання, класифікація зображень, визначення ознак.

Постановка проблеми. В останні десятиліття зріс обсяг доступної візуальної інформації та її ефективна обробка стала одним із найважливіших напрямів досліджень у галузі комп'ютерного зору. Візуальна інформація, що міститься у зображеннях та відео, має величезний потенціал для різних галузей, включаючи медицину, автоматизацію виробництва, робототехніку, відеоспостереження та багато інших. Обробка цієї інформації є складним і трудомістким завданням, що потребує значних ресурсів та часу. [9]

Традиційні методи обробки візуальної інформації, такі як ручна обробка або використання простих алгоритмів, часто обмежені у своїй точності, ефективності та масштабованості. Ці методи часто вимагають великого обсягу ручної роботи, включаючи розмітку даних або завдання визначення правил обробки зображень. Крім того, такі методи не завжди здатні автоматично адаптуватися до різних типів та умов візуальних даних, що робить їх менш

гнучкими та ефективними. У зв'язку з цим виникає потреба у розробці автоматизованої моделі обробки візуальної інформації, що базується на передових методах машинного навчання. Така модель здатна обробляти та аналізувати візуальні дані з мінімальною участю людини, використовуючи великі набори даних для навчання та адаптації. Це дозволяє досягти більш точних і надійних результатів у завданнях обробки візуальної інформації.

Аналіз останніх досліджень і публікацій. Існує низка робіт, які представляють огляд існуючих методів та підходів до обробки візуальної інформації. У роботі [7] автори проводять огляд різних алгоритмів сегментації зображень, включаючи методи на основі порогового значення, графових алгоритмів, методи активного контуру та нейронних мереж. У роботі [1] досліджується застосування глибокого навчання до завдання детектування об'єктів на зображеннях. Автори представляють огляд різних архітектур нейронних мереж, таких як R-CNN, Fast R-CNN, Faster R-CNN, SSD та YOLO. Розглядаються методи обробки зображень, які використовуються для покращення точності детектування об'єктів. Автори статті [10] досліджують досягнення у сфері глибокого навчання класифікації зображень. Описані методи попередньої обробки зображень та оптимізації моделей для покращення точності класифікації. У статті [8] аналізуються методи обробки та розмітки даних, а також техніки постобробки для покращення якості сегментації. У роботі [4] розглядаються техніки попередньої обробки даних і адаптації моделей для досягнення високої точності та швидкості відстеження об'єктів.

Мета статті: дослідження комплексних моделей, які здатні автоматизовано обробляти та аналізувати різноманітні форми візуальної інформації.

Викладення основного матеріалу дослідження. На теперішній час обробка візуальної інформації має велику актуальність з кількох причин:

1. *Зростання обсягу візуальних даних.* З розвитком технологій зйомки, зберігання та передачі зображень обсяг візуальних даних зростає у великих масштабах – включаючи фотографії, відео, медичні зображення, супутникові знімки, зображення соціальних мереж. Обробка та аналіз такого обсягу даних потребує ефективних методів та моделей.

2. *Застосування у різних галузях.* Обробка візуальної інформації знаходить застосування у багатьох галузях, включаючи медицину, автомобільну промисловість, робототехніку, безпеку, рекламу та ігрову індустрію. Від виявлення захворювань на медичних зображеннях до автоматичного розпізнавання об'єктів на дорозі обробка візуальної інформації зараз відіграє ключову роль[6].

3. Поліпшення якості життя та зручності. Застосування обробки візуальної інформації призводить до покращення якості життя та створення зручностей для людей. Автоматична класифікація та організація фотографій, аналіз емоцій на обличчях людей, автоматичну рекомендацію товарів на основі візуальних уподобань та багато іншого.

Таблиця 1

Порівняльний аналіз моделей обробки візуальної інформації

Модель	Застосування	Переваги	Недоліки
Згорткові нейронні мережі (CNN)	Класифікація зображень, сегментація зображень	Виділення просторових ознак, висока точність	Не враховують контексту інформації
Глибокі нейронні мережі (DNN)	Класифікація зображень, генерація зображень	Визначення абстрактних ознак, висока гнучкість	Вимагають великої кількості даних та ресурсів
Рекурентні нейронні мережі (RNN)	Розпізнавання рукописного тексту, машинний переклад	Обробка послідовних даних	Обмежено визначають довгострокові залежності
Генеративні змагальні мережі (GAN)	Генерація зображень, підвищення роздільної здатності	Створення реалістичних даних	Нестійкість навчання, труднощі щодо оцінки якості генерації
Трансформери (Transformers)	Сегментація зображень, генерація описів	Ефективне використання механізму уваги	Висока обчислювальна складність

Обробка візуальної інформації стала одним із важливих аспектів розвитку штучного інтелекту. Глибоке навчання та нейронні мережі показали визначні результати в області обробки зображень та відео, перевершуючи людей у деяких завданнях, таких як класифікація зображень, розпізнавання об'єктів та сегментація.

Протягом останніх десятиліть запропоновано безліч моделей та алгоритмів для аналізу, інтерпретації та розуміння візуальних даних. В таблиці

1 представлено основні моделі обробки візуальної інформації, їх застосування, переваги та недоліки. Дані переваги та недоліки можуть змінюватись в залежності від конкретної реалізації моделі та контексту застосування. [2]

При оцінці моделі автоматизованої обробки візуальної інформації реальних даних використовують метрики. Основні метрики:

Точність. Метрика, яка вимірює частку правильно класифікованих зразків від загальної кількості зразків. Вона особливо корисна при завданні класифікації зображень, де необхідно визначити правильну категорію для кожного зображення.

Повнота. Метрика вимірює здатність моделі виявляти всі позитивні зразки. Вона корисна, коли важливо мінімізувати хибні результати і виявити якнайбільше істинно позитивних зразків.

F1-міра. Середнє значення між точністю та повнотою. F1-міра враховує точність і повноту та дозволяє оцінити баланс між ними.

Середня абсолютна помилка. Метрика використовується у задачах регресії та показує середню абсолютну різницю між прогнозованими та фактичними значеннями. Вона дозволяє оцінити, наскільки точно модель передбачає значення.

Залежно від конкретного завдання та типу даних, які обробляються, можливо вибрати відповідні метрики для оцінки моделі. Деякі метрики можуть бути релевантнішими в одних випадках, ніж в інших, тому важливо вибрати ті, які відображають особливості завдання.

В таблиці 2 представлено порівняльний аналіз відомих моделей обробки візуальної інформації з використанням декількох параметрів, включаючи точність, повноту, F1-міру, швидкість навчання та час навчання. Параметр «Швидкість навчання» (Learning Rate) вказує, наскільки швидко модель адаптується до даних під час навчання. Більш високе значення може призвести до швидкої збіжності, але може викликати перенавчання моделі. Час навчання відображає оцінку часу, витраченого для навчання моделі на певних даних. Дані параметри є зразковими і можуть відрізнятися залежно від конкретних умов та налаштувань навчання моделей.

Порівняння основних метрик моделей обробки візуальної інформації

Модель	Точність	Повнота	F1-міра	Швидкість навчання	Час навчання
Згорткові нейронні мережі (CNN)	0.85	0.87	0.86	0.001	12 годин
Глибокі нейронні мережі (DNN)	0.89	0.84	0.86	0.0001	16 годин
Рекурентні нейронні мережі (RNN)	0.82	0.78	0.80	0.0005	10 годин
Генеративні змагальні мережі (GAN)	0.76	0.81	0.78	0.0002	20 годин
Трансформери (Transformers)	0.91	0.92	0.91	0.0003	8 годин

Опишемо модель обробки візуальної інформації, яка складається з кількох компонентів, кожен із яких виконує певні операції з обробці даних та їх взаємодії:

1. *Вхідні дані.* Вхідними даними є зображенням або пакет зображень, які надходять на вхід моделі. Розмір вхідних зображень визначається у вигляді (ширина, висота, канали), де канали позначають колірні канали зображення (наприклад, RGB).

2. *Препроцесинг.* Препроцесинг виконує попередню обробку вхідних даних для підготовки їх до обробки моделлю. Даний етап може включати нормалізацію даних, зміну розмірів зображень, приведення до необхідного формату та інші операції обробки.

3. *Згорткові шари CNN.* Згорткові шари є ключовим компонентом моделі та виконують операцію згортки і активації на вхідних даних. Кожен згортковий шар має свій набір фільтрів, які застосовуються до вхідних даних для отримання ознак. Після згортки активація застосовується для введення нелінійності у вихідні шари.

4. *Пулінг шари.* Пулінг виконує операцію зменшення розміру вихідних даних зі згорткових шарів. Зазвичай використовується операція пулінгу максимуму, яка вибирає найбільший елемент вікна пулінга. Пулінг допомагає змен-

шити розмір даних, покращити інваріантність до масштабу та знизити обчислювальну складність моделі.

5. Повнозв'язані шари. Повнозв'язані шари приймають вихідні дані з пулінг шарів і виконують операції лінійного перетворення та активації. Кожен нейрон у пов'язаному шарі зв'язаний із усіма нейронами попереднього шару. Повнозв'язані шари виконують більш високорівневу обробку ознак і представляють їх у вигляді вектору.

6. Вихідний шар. Вихідний шар моделі приймає вектор ознак із пов'язаних шарів і перетворює його на ймовірність для різних класів. Зазвичай використовується функція активації Softmax для отримання нормованих ймовірностей. Вихідний шар представляє фінальні прогнози моделі для класифікації візуальної інформації.

7. Навчання моделі. Для навчання моделі необхідно визначити функцію втрат, оптимізатор та параметри навчання. Функція втрат визначає різницю між прогнозами моделі та дійсними мітками даних. Оптимізатор оновлює ваги моделі на основі градієнтного спуску, мінімізуючи функцію втрат. Параметри навчання, такі як швидкість навчання та кількість епох визначають процес навчання моделі.

8. Використання моделі. Після навчання модель може бути використана для класифікації, розпізнавання об'єктів та інших завдань обробки візуальної інформації. Подача вхідних даних на вхід моделі дозволяє отримати передбачення або вихідні результати, які можуть бути використані далі для аналізу та прийняття рішень [7].

Представимо модель автоматизованої обробки візуальної інформації для класифікації зображень з використанням нейронних мереж:

Вхідні дані. Довільний набір зображень представляється як вектори ознак або матриці пікселів. Позначимо зображення як $x = [x_1, x_2, \dots, x_n]$, де n – розмірність вектору ознак або розмір зображення.

Параметри моделі. K та d – параметри моделі, які потрібно визначити в процесі навчання. K є ваговою матрицею, а d – вектор зсувів.

Архітектура моделі. Використовуємо просту нейронну мережу з декількома прихованими шарами. Позначимо приховані шари як $h_1, h_2, \dots, h_m$, де m – кількість прихованих шарів. Кожен прихований шар є лінійною комбінацією вхідних даних і ваг, які застосовуються до нелінійної функції активації.

Математична модель. $f(x, K, d)$ – математична модель, яка приймає вхідні дані x , параметри моделі K і d , та обчислює передбачені значення y . Тоді модель набуває наступного вигляду:

$$\begin{aligned} c_1 &= \text{activation}(K_1 * x + d_1) \\ c_2 &= \text{activation}(K_2 * x + d_2) \\ &\dots \\ c_m &= \text{activation}(K_m * x + d_m) \\ y &= \text{softmax}(K_y * c_m + d_m) \end{aligned}$$

де *activation* – функція активації, така як ReLU, сигмоїд або гіперболічний тангенс, *softmax* – багатокласові класифікації, $*$ – означає операцію множення матриці.

Навчання моделі. Мета полягає в тому, щоб налаштувати параметри моделі K та d , щоб мінімізувати функцію втрат між передбаченими значеннями та фактичними значеннями. Для цього застосовуються методи оптимізації, такі як градієнтний стохастичний спуск або його варіації.

Таблиця 3

Результати роботи автоматизованої моделі на різних вхідних даних

Вхідні дані	Точність	Повнота	F1-міра	MAE	MSE
Вхідні дані 1	0.85	0.92	0.88	0.15	0.28
Вхідні дані 2	0.91	0.87	0.89	0.12	0.24
Вхідні дані 3	0.78	0.94	0.86	0.18	0.32
Вхідні дані 4	0.92	0.88	0.89	0.11	0.22

Модель може розширюватися та містити додаткові компоненти, такі як шари згортки, пулінг, рекурентні зв'язки, залежно від конкретної задачі та алгоритму. В таблиці 3 представлені результати роботи моделі на різних вхідних даних за наступними метриками: точність, повнота, F1-міра, середня абсолютна помилка (MAE) та середньоквадратична помилка (MSE) Таблиця може бути розширена з додатковими метриками залежно від конкретного завдання обробки візуальної інформації.

Висновки. В роботі досліджено комплексну модель обробки візуальної інформації, яка заснована на комбінованому підході, що включає застосування різних алгоритмів та методів. Це дозволило досягти високої точності та ефективності обробки візуальних даних. Автоматизована модель має широкий потенціал застосування у різних галузях, включаючи медицину, робототехніку,

автоматичне водіння та інші. Її здатність автоматично обробляти та аналізувати візуальні дані може значно покращити процеси прийняття рішень та підвищити ефективність роботи в різних галузях.

Слід зазначити, що розроблена модель також має обмеження. Наприклад, вона може бути чутливою до якості вхідних даних і вимагати великого обсягу обчислювальних ресурсів. Також необхідно провести подальші дослідження для оптимізації та адаптації моделі під конкретні програми та завдання. В цілому, результати цього дослідження підтверджують значущість розробки автоматизованих моделей обробки візуальної інформації та їх потенціал для покращення процесів аналізу та прийняття рішень. Подальший розвиток у цій галузі може призвести до нових проривів у різних додатках, де візуальна інформація відіграє важливу роль.

ЛІТЕРАТУРА

1. Ajeet Ram Pathak. Application of Deep Learning for Object Detection / Ajeet Ram Pathak, Manjusha Pandey, Siddharth Rautaray // International Conference on Computational Intelligence and Data Science (ICCIDS 2018). – 2018. – P. 1706-1717.
2. Christ P.F. Automatic liver and lesion segmentation in CT using cascaded fully convolutional neural networks and 3D conditional random fields / P. F. Christ, M. E. A. Elshaer, F. Ettlinger, S. Tatavarty, M. Bickel, P. Bilic, M. Rempfler, M. Armbruster, F. Hofmann, M. D'Anastasi, W. H. Sommer, S.-A. Ahmadi, B. H. Menze // MICCAI, Cham. – 2016. – P. 415–423.
3. Gonal J. S. Morphological Segmentation of the Brain Tumors by Using Image Processing and Lab-VIEW /J. S. Gonal, V. V. Kohir // X International Journal of Engineering Research & Technology (IJERT) – Volume 4, Special Issue. – 2015. – P. 334-341.
4. Hanchi Liu. Application of Deep Learning-Based Object Detection Techniques in Fish Aquaculture: A Review / Hanchi Liu, Xin Ma, Yining Yu, Liang Wang, Lin Hao // Journal of Marine Science and Engineering. – 2023 – №11(4):867.
5. Kapil Kumar Gupta. A Comparative Study of Medical Image Segmentation Techniques for Brain Tumor Detection / Kapil Kumar Gupta, Namrata Dhanda, Upendra Kumar // 4th International Conference on Computing Communication and Automation (ICCCA). – 2018. – P. 1-4.
6. Krasilenko Vladimir G. Application of nonlinear correlation functions and equivalence models in ad-vanced neuronets / Vladimir G. Krasilenko, Oleg K. Kolesnitsky, Anatoly K. Bogukhvalsky // Proceedings of SPIE – Vol. 3317. – 2017. – P. 211-222.

7. Pushpajit A. Khaire. An Overview Of Image Segmentation Algorithms / Pushpajit A. Khaire, Nileshsingh V. Thakur // International Journal of Image Processing and Vision Science. – 2013. – Vol. 1: Iss. 3, Article 1. – P. 150-156.

DOI: 10.47893/IJIPVS.2013.1028

8. Rituparna Sarma. A comparative study of new and existing segmentation techniques / Rituparna Sarma, Yogesh Kumar Gupta // IOP Conference Series: Materials Science and Engineering, Volume 1022, 1st International Conference on Computational Research and Data Analytics (ICCRDA). – 2020. – P.1022-1033.

DOI 10.1088/1757-899X/1022/1/012027

9. Ronneberger O. U-net: Convolutional networks for biomedical image segmentation / Ronneberger O., Fischer P., Brox T. // MICCAI, Vol. 9351. – 2015. – P. 234–241.

10. Sharmila T. Impact of applying pre-processing techniques for improving classification accuracy / Sharmila T., Ramar K., Thangaswamy Sree Renga Raja // Signal, Image and Video Processing. – 2014. – №8(1). – P. 149-157. DOI:8.10.1007/s11760-013-0505-7.

REFERENCES

1. Ajeet Ram Pathak. Application of Deep Learning for Object Detection / Ajeet Ram Pathak, Manjusha Pandey, Siddharth Rautaray // International Conference on Computational Intelligence and Data Science (ICCIDS 2018). – 2018. – P. 1706-1717.

2. Christ P.F. Automatic liver and lesion segmentation in CT using cascaded fully convolutional neural networks and 3D conditional random fields / P. F. Christ, M. E. A. Elshaer, F. Ettlinger, S. Tatavarty, M. Bickel, P. Bilic, M. Rempfler, M. Armbruster, F. Hofmann, M. D’Anastasi, W. H. Sommer, S.-A. Ahmadi, B. H. Menze // MICCAI, Cham. – 2016. – P. 415–423.

3. Gonal J. S. Morphological Segmentation of the Brain Tumors by Using Image Processing and Lab-VIEW / J. S. Gonal, V. V. Kohir // X International Journal of Engineering Research & Technology (IJERT) – Volume 4, Special Issue. – 2015. – P. 334-341.

4. Hanchi Liu. Application of Deep Learning-Based Object Detection Techniques in Fish Aquaculture: A Review / Hanchi Liu, Xin Ma, Yining Yu, Liang Wang, Lin Hao // Journal of Marine Science and Engineering. – 2023 – №11(4):867.

5. Kapil Kumar Gupta. A Comparative Study of Medical Image Segmentation Techniques for Brain Tumor Detection / Kapil Kumar Gupta, Namrata Dhanda, Upendra Kumar // 4th International Conference on Computing Communication and Automation (ICCCA). – 2018. – P. 1-4.

6. Krasilenko Vladimir G. Application of nonlinear correlation functions and equivalence models in ad-vanced neuronets / Vladimir G. Krasilenko, Oleg K. Kolesnitsky, Anatoly K. Bogukhvalsky // Proceedings of SPIE – Vol. 3317. – 2017. – P. 211-222. 444
7. Pushpajit A. Khaire. An Overview Of Image Segmentation Algorithms / Pushpajit A. Khaire, Nileshsingh V. Thakur // International Journal of Image Processing and Vision Science. – 2013. – Vol. 1: Iss. 3 , Article 1. – P. 150-156. DOI: 10.47893/IJIPVS.2013.1028
8. Rituparna Sarma. A comparative study of new and existing segmentation techniques / Rituparna Sarma, Yogesh Kumar Gupta // IOP Conference Series: Materials Science and Engineering, Volume 1022, 1st International Conference on Computational Research and Data Analytics (ICCRDA). – 2020. – P.1022-1033. DOI 10.1088/1757-899X/1022/1/012027
9. Ronneberger O. U-net: Convolutional networks for biomedical image segmentation / Ronneberger O., Fischer P., Brox T. // MICCAI, Vol. 9351. – 2015. – P. 234–241.
10. Sharmila T. Impact of applying pre-processing techniques for improving classification accuracy / Sharmila T., Ramar K., Thangaswamy Sree Renga Raja // Signal, Image and Video Processing. – 2014. – №8(1). – P. 149-157. DOI:8. 10.1007/s11760-013-0505-7.

Received 11.05.2023.
Accepted 21.05.2023.

Automated models of visual information processing

The article presents a study devoted to the development and research of an automated model of visual information processing. The goal of the research was to create a comprehensive model capable of automatically processing and analyzing various forms of visual data, such as images and videos. The model is developed on the basis of a combined approach that combines various algorithms and methods of visual information processing. The literature review conducted within the scope of this study allowed us to study the existing methods and algorithms for visual information processing. Various image processing approaches were analyzed, including segmentation, pattern recognition, object classification and detection, video analysis, and other aspects. As a result of the review, the advantages and limitations of each approach were identified, as well as the areas of their application were determined. The developed model showed high accuracy and efficiency in visual data processing. It can successfully cope with the tasks of segmentation, recognition and classification of objects, as well as video analysis. The results of the study confirmed the superiority of the proposed model. Potential applications of

the automated model are considered, such as medicine, robotics, security, and many others. However, limitations of the model such as computational resource requirements and quality of input data are also noted. Further development of this research can be aimed at optimizing the model, adapting it to specific tasks and expanding its functionality. In general, the study confirms the importance of automated models of visual information processing and its important place in modern technologies. The results of the research can be useful for the development of new systems based on visual data processing and contribute to progress in the field of computer vision and artificial intelligence.

Keywords: image processing, pattern recognition, object detection, machine learning, image classification, feature extraction.

Могильний Олександр Анатолійович – аспірант кафедри робототехніки та спеціалізованих комп’ютерних систем Черкаського державного технологічного університету.

Mohylnyi Oleksandr – post-graduate student at the Department of Robotics and Specialized Computer Systems at Cherkasy State Technological University.

Д.В. Солдатенко, Вік.В. Гнатушенко

ПОКРАЩЕННЯ ЕФЕКТИВНОСТІ РОЗПІЗНАВАННЯ СУПУТНИКОВИХ ЗОБРАЖЕНЬ ШЛЯХОМ ВИЗНАЧЕННЯ ОБСЯГУ НАВЧАЛЬНИХ ДАНИХ

Анотація. Розпізнавання супутникових зображень є життєво важливим застосуванням комп'ютерного зору з потенційними варіантами використання в таких сферах, як боротьба зі стихійними лихами, землеробство та міське планування. Це дослідження спрямоване на визначення оптимальної кількості вхідних даних, та підбору оптимальних методів їх аугментації, необхідних для навчання нейронної мережі CNN для розпізнавання супутникових зображень. З цієї метою проводиться серія експериментів, щоб дослідити вплив кількості вхідних даних на кілька показників продуктивності, включаючи точність, конвергенцію та узагальнення моделі. Дослідження пропонує кілька методів для визначення точки насичення та пом'якшення наслідків перенавчання. Результати, отримані в цьому дослідженні, можуть допомогти в розробці більш ефективних моделей розпізнавання супутникових зображень.

Ключові слова: нейронна мережа, розпізнавання зображень, супутникові знімки, обробка зображень, аугментація даних, штучний інтелект, згортова нейромережа, класифікація зображень.

Постановка проблеми. При підготовці даних, які використовуються для навчання нейромережі, в деяких випадках можна зіштовхнутися з їх недоліком, низькою якістю або необхідності в симуляції різноманітних вторинних умов, наприклад, погіршення або зміна погоди. Не завжди є можливість запросити оновлені дані або збільшити їх кількість. Аугментація даних - є одним з варіантів розв'язання проблеми недостатчі даних, але занадто покладатися на цей метод також не треба. Занадто сильне збільшення вхідних даних за допомогою аугментації може спричинити перенавчання моделі, що вплине на кінцевий результат. Під час аугментації даних необхідно враховувати, що нейромережі можуть бути чутливі до різних методів аугментації. Наприклад, деякі методи аугментації можуть підвищити різноманітність даних, а інші можуть спричинити перенавчання моделі або погіршення її точності. У деяких випадках може бути корисно використовувати комбінації різних типів аугментації для збільшення різноманітності даних та підвищення робастності моделі до різних вторинних умов.

© Солдатенко Д.В., Гнатушенко Вік.В., 2023

Враховуючи вищезгадане, використання аугментації даних є корисним інструментом при навчанні нейромереж, але його потрібно застосовувати з обережністю та дотримуватися оптимального балансу між збільшенням різноманітності даних та уникненням перенавчання моделі.

Таким чином, підбір оптимальної кількості аугментованих даних, для визначення точки насичення та пом'якшення наслідків перенавчання нейромережі є актуальною задачею.

Аналіз останніх досліджень і публікацій. Дослідження показують, що використання аугментації даних може підвищити точність моделі та зменшити вплив перенавчання. Проте ми спостерігаємо деякий дисбаланс у розвитку різних підходів. Хоча для різних задач комп'ютерного зору було запропоновано різноманітні архітектури нейромереж і обчислювальна потужність графічних процесорів (GPU) стрімко збільшується, менше уваги приділяється застосуванню методів аугментації даних для генерації якісних тренувальних даних. Основна ідея аугментації даних полягає в підвищенні достатності та різноманітності тренувальних даних шляхом створення синтетичного набору даних. Аугментовані дані можуть бути розглянуті як ті, що взяті з розподілу, який близький до реального. У цьому випадку, розширений набір даних може представляти повніші характеристики. Але деякі проблеми залишаються в методах аугментації даних, які використовуються до вхідних зображень.

По-перше, методи доповнення даних можуть бути застосовані в різних типах задач комп'ютерного зору, таких як виявлення об'єктів [1], семантична сегментація [2] та класифікація зображень [3]. Проте виклик полягає в тому, що методи доповнення даних є незалежними від задач. Оскільки операції виконуються одночасно на даних зображень та їх класах, а типи класів різні для різних завдань, методи аугментації даних для задачі виявлення об'єктів не можуть бути безпосередньо застосовані до задачі семантичної сегментації. Це призводить до неефективності та низької масштабованості.

По-друге, не існує теоретичного дослідження з аугментації даних. Наприклад, не існує кількісних стандартів для достатнього розміру тренувальних наборів даних. Розмір згенерованих тренувальних даних зазвичай проєктується згідно з особистим досвідом та розлогими експериментами. Крім того, може існувати парадокс, коли розмір початкового набору даних настільки малий [4], що ми стикаємося з викликом, як згенерувати якісні дані на основі дуже малої кількості даних.

Таким чином, невирішеним є завдання підбору методів аугментації та потрібної кількості даних, які можуть бути найбільш релевантними під базові класи супутникових зображень.

Мета дослідження. Це дослідження спрямоване на визначення оптимальної кількості вхідних даних, для підбору оптимальних методів аугментації даних, які підходять для різних видів вхідних даних. З цією метою проводиться серія експериментів, щоб дослідити вплив кількості вхідних даних на кілька показників продуктивності, включаючи точність, конвергенцію та узагальнення моделі.

Викладення основного матеріалу дослідження. Аугментація даних - це процес штучного збільшення об'єму даних шляхом застосування різноманітних операцій до наявних зразків даних. Для зображень слід виділити два з найбільш популярні методи аугментації даних, та підгрупи що належать до них:

- **Позиційна аугментація:**
 - Масштабування (scaling);
 - Обрізка (cropping);
 - Відображення (flipping);
 - Додавання краю (padding);
 - Повороту (rotation);
 - Переміщення (translation).
- **Кольорова аугментація:**
 - Яскравість (brightness);
 - Контрастність (contrast);
 - Насиченість (saturation);
 - Відтінок (hue);
 - Додавання шуму (noise).

У цьому дослідженні використовуються деякі з наведених вище методів аугментації, а саме: масштабування, відображення, поворот та додавання шуму, ці методи обрані як одні з найбільш використовуваних. Як нейромережа використовується CNN[5] (згорткова нейронна мережа). Також буди використанні вхідні зображення (RGB), з розмірами 128 на 128 пікселів, до яких застосовується аугментація. Кількість зображень була 16 для навчання та 8 для валідації, для кожного окремого класу. Давайте детальніше розберемо кожен окремий метод аугментації та потім перейдемо до самих експериментів з підбору відповідний, які є сенс використовувати під різні класи зображень.

Розпочнемо з одного з найпростіших в реалізації, зі сторони математичного алгоритму, методів аугментації даних, який можливо використовувати для зображення - відображення, який полягає в дзеркальному відображенні зображення вздовж вертикальної або горизонтальної осей. Цей метод може бути корисним, коли об'єкти на зображенні симетричні відносно потрібної осі, або коли нам потрібні додаткові зображення для підвищення точності моделі.

Математично відображення зображення можна описати так: нехай $I(x, y)$ - початкове зображення з розмірами (H, W) , тоді відображене зображення $I'(x', y')$ з розмірами (H, W) можна отримати за допомогою наступної формули (1), для горизонтального відображення:

$$I'(x', y') = I(W - x' - 1, y'), \quad 1)$$

де x' та y' - координати пікселя на відображеному зображенні, W - ширина початкового зображення. Ця формула говорить нам, що кожен піксель на горизонтально відображеному зображенні $I'(x', y')$ можна отримати, обертаючи початкове зображення $I(x, y)$ вздовж вертикальної осі. Приклад горизонтального відображення зображення можна побачити на рис. 1.

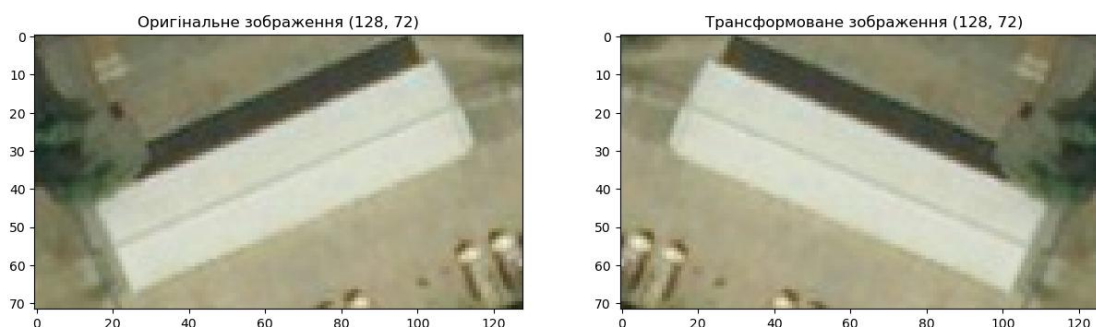


Рисунок 1 - Приклад оригінального та горизонтально відображеного зображень

Формула (2) для вертикального відображення зображення не сильно відрізняється від горизонтального і буде наступною:

$$I'(x', y') = I(x', H - y' - 1), \quad 2)$$

де (x', y') - нові координати пікселя, (x, y) - початкові координати пікселя, H - висота початкового зображення. Як і у першому випадку, з формулою горизонтального відображення, кожен піксель на вертикально відображеному зображенні $I'(x', y')$ можна отримати, обертаючи початкове зображення $I(x, y)$

вздовж горизонтальної осі. Приклад використання вертикального відображення зображення можна побачити на рис. 2.

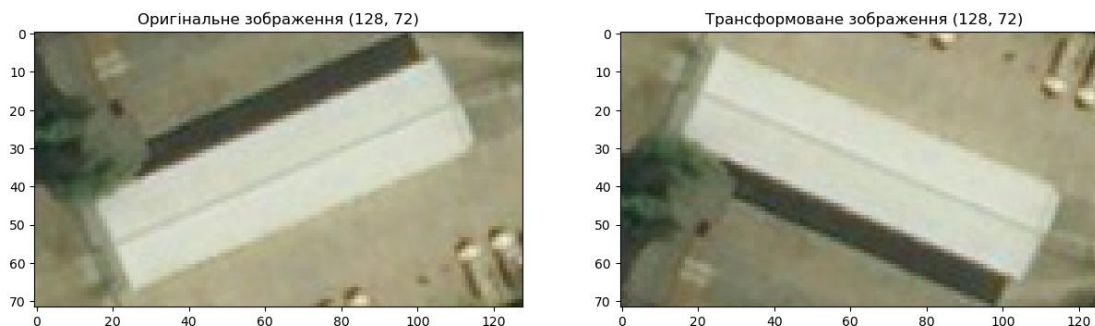


Рисунок 2 - Приклад оригінального та вертикально відображеного зображень

Наведені вище формули відображення використовується для зображень з одним каналом (чорно-білих) або зображень з кількома каналами [6] (RGB, наприклад). У випадку зображень з кількома каналами, ці формули потрібно застосовувати до кожного каналу окремо.

Наступний з методів аугментації даних - поворот. Цей метод також використовується для збільшення розміру навчальної вибірки та покращення здатності моделі до узагальнення, шляхом створення нових зображень з різним кутом огляду.

Математично поворот зображення можна описати так: нехай $I(x, y)$ буде зображенням з координатами пікселів (x, y) , де x - номер стовпця, а y - номер рядка. Нехай $I'(x', y')$ буде повернутим зображенням з координатами пікселів (x', y') , де x' - номер стовпця, а y' - номер рядка.

Для обчислення повороту зображення на кут θ за годинниковою стрілкою навколо центру зображення ми можемо використовувати наступні формули:

$$x' = (x - x_c) \cos \theta - (y - y_c) \sin \theta + x_c, \quad (3)$$

$$y' = (x - x_c) \sin \theta + (y - y_c) \cos \theta + y_c, \quad (4)$$

де (x_c, y_c) є координатами центру зображення, (x, y) - координати пікселя на зображенні, а θ - кут обертання в радіанах. Формули (3-4) враховують, що пікселі повинні бути перенесені на нові координати (x', y') , щоб отримати нове зображення, обернуте на кут θ .

Для застосування методу повороту потрібно вибрати випадковий кут обертання θ з певного діапазону, наприклад, від -15 до 15 градусів. Отже, якщо ми маємо зображення x і випадковий кут обертання θ , то зображення з застосованим методом повороту можна записати наступним чином:

$$x_{\text{rot}} = \text{rotate}(x, \theta), \quad (5)$$

де *rotate* - функція, яка застосовує формулу (5) обертання до зображення x на кут θ . Приклад повороту зображення можна побачити на рис. 3.

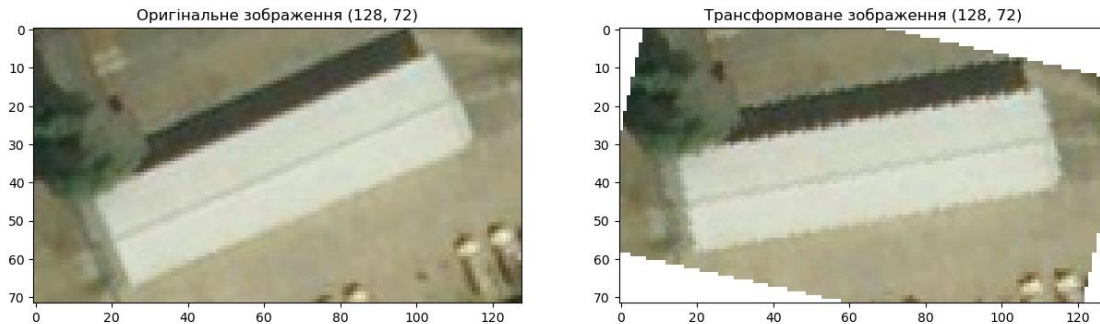


Рисунок 3 - Приклад оригінального та повернутого зображень

Важливо пам'ятати, що після повороту зображення ми можемо втратити частину інформації. Якщо ми повертаємо зображення на більше ніж 90 градусів, ми можемо втратити деякі деталі, такі як текст чи обличчя на зображенні, але у більшості випадків це не стосується супутникових знімків. Тому, при застосуванні методу повороту, ми повинні бути уважні, щоб не втратити важливу інформацію на зображенні.

Зміна масштабу[7] (або ресайзінг) зображення - це метод аугментації даних, що полягає в зміні розміру зображення без зміни його вмісту. Цей метод також може бути корисним для збільшення розміру даних або для адаптації зображення до потрібного розміру, другий варіант більше підходить якщо потрібно зробити дані однакового розміру.

Для зміни масштабу зображення можна використовувати формули білінійної інтерполяції. Для спрощення пояснення, ми розглянемо зміну масштабу зображення вдвічі.

Нехай $I(x, y)$ - початкове зображення з розмірами (H, W) , а $I'(x', y')$ - змінене зображення з розмірами $(\frac{H}{2}, \frac{W}{2})$. Для того, щоб отримати значення пікселя в точці (x', y') , нам потрібно взяти значення пікселів у чотирьох сусідніх точках (x_1, y_1) , (x_1, y_2) , (x_2, y_1) та (x_2, y_2) , де $(x_1, y_1) = (\lfloor x' \rfloor, \lfloor y' \rfloor)$, $(x_1, y_2) = (\lfloor x' \rfloor, \lfloor y' \rfloor + 1)$, $(x_2, y_1) = (\lfloor x' \rfloor + 1, \lfloor y' \rfloor)$ та $(x_2, y_2) = (\lfloor x' \rfloor + 1, \lfloor y' \rfloor + 1)$. Значення пікселів у цих точках можна обчислити (6) за допомогою лінійної інтерполяції:

$$I'(x', y') = \frac{(x_2 - x')(y_2 - y')I(x_1, y_1) + (x' - x_1)(y_2 - y')I(x_2, y_1) + (x_2 - x')(y' - y_1)I(x_1, y_2) + (x' - x_1)(y' - y_1)I(x_2, y_2)}{(x_2 - x_1)(y_2 - y_1)}, \quad (6)$$

де $\lfloor \cdot \rfloor$ - округлення до меншого цілого, а $\lceil \cdot \rceil$ - округлення до більшого цілого. Приклад зміну масштабу зображення, але зі збереженою висотою, для більш наочного приклада, можна побачити на рис. 4.

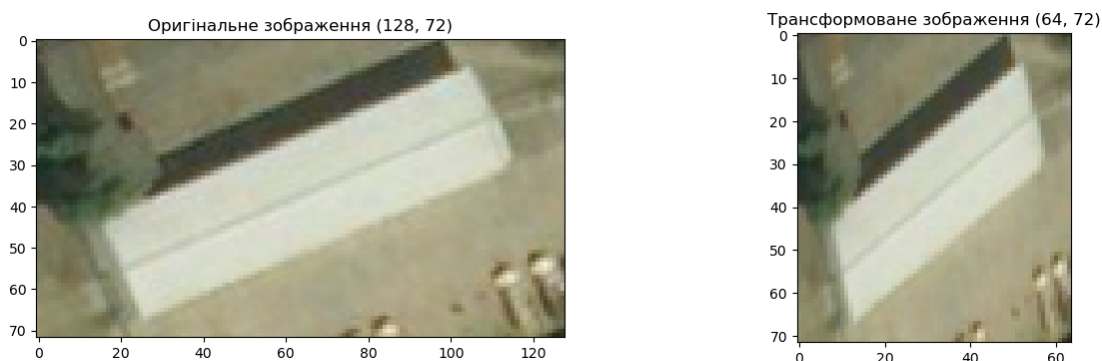


Рисунок 4 - Приклад оригінального та зображення зі зміненим масштабом

Додавання шуму є одним з методів аугментації даних, який може бути використаний для покращення результатів моделей машинного навчання. Гаусівський шум - це один з типів шуму, який може бути доданий до даних. Формула для додавання гаусівського шуму виглядає наступним чином:

$$x_{\text{noisy}} = x + \epsilon, \quad (7)$$

де x - оригінальний сигнал, ϵ - гаусівський шум з середнім значенням 0 та стандартним відхиленням σ .

Вибір стандартного відхилення σ залежить від домену задачі та рівня шуму, який потрібно додати до даних. При занадто високому значенні σ , може статися так, що шум буде переважати сигнал, і модель може втратити здатність до правильної класифікації.

Стандартне відхилення гаусівського шуму (8) визначається наступною формулою:

$$\sigma = \sqrt{\frac{1}{N} \sum_{i=1}^N (x_i - \mu)^2} \quad (8)$$

де x_i - значення сигналу, N - кількість значень у сигналі, μ - середнє значення сигналу.

Для кольорових зображень, гаусівський шум додається до кожного пікселя зображення з окремими значеннями $\epsilon_{i,j,c}$, де i та j - координати пікселя, а c - номер кольорового каналу (червоний, зелений, синій).

Отже, формула (8) для додавання гаусівського шуму до кольорового зображення може мати вигляд:

$$x_{\text{noisy}}(i, j, c) = x(i, j, c) + \epsilon_{i, j, c} \quad (9)$$

де $x(i, j, c)$ - значення пікселя на позиції (i, j) та кольоровому каналі c оригінального зображення, а $\epsilon_{i, j, c}$ - гаусівський шум зі стандартним відхиленням σ для кожного кольорового каналу.

У випадку кольорових зображень, формула для генерації гаусівського шуму залишається такою ж, як і для чорно-білих зображень. Тобто, значення шуму $\epsilon_{i, j, c}$ генеруються незалежно (10) для кожного кольорового каналу ($c \in [0, 2]$) і для кожного пікселя (i, j) за формулою:

$$\epsilon \sim \mathcal{N}(0, \sigma_c^2) \quad (10)$$

де $\mathcal{N}$ - нормальний розподіл з середнім значенням 0 та стандартним відхиленням σ_c для кожного кольорового каналу.

При виборі значення σ_c для кожного кольорового каналу, слід також враховувати рівень шуму, який ви хочете додати до даних, та специфіку домену задачі.

Додавання гаусівського шуму [8] до даних може бути корисним для зменшення перенавчання моделей, підвищення стійкості до шуму вхідних даних та збільшення різноманітності даних для тренування моделей. Приклад додавання гаусівського шуму до зображення можна побачити на рис. 5.

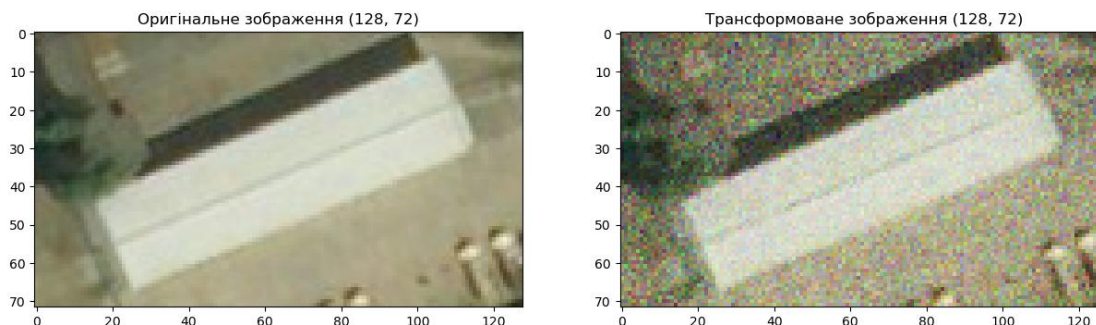


Рисунок 5 - Приклад оригінального та зображення з доданим шумом

Використовуючи перелічені вище методи аугментації, була подвоєна кількість вхідних даних для кожного окремого класу, при використанні кожного окремого метода, та у п'ять разів при використанні усіх одразу. Результати з точністю розпізнавання CNN після навчання на даних, отриманих за допомогою різних методів аугментації наведені у табл. 1.

Перш ніж почнемо розбирати результати - зауважте, що точність може коливатися залежно від того, які методи аугментації використовуються, а також від розміру мережі та інших факторів, які впливають на її навчання. Дані табл

1. є лише орієнтовним значенням, отриманими саму з цією нейромережею та вхідними даними.

Таблиця 1

Результати навчання CNN з та без використання аугментованих даних

Метод	Точність розпізнання			
	Вода	Дороги	Будинки	Дерева
Тільки початкові дані	0.85	0.92	0.78	0.81
Масштабування	0.80	0.94	0.76	0.79
Відображення	0.86	0.95	0.82	0.83
Поворот	0.80	0.92	0.81	0.78
Шум	0.85	0.89	0.84	0.80
Найбільш результативні	0.90	0.94	0.86	0.88

Оцінки кожного окремого метода аугментації, та варіанта з використанням усіх методів та їх поєднання:

- Масштабування - не показав значного покращення точності для жодного з класів, у деяких випадках навіть спричинив значну рецесію, тож можна вважати, що використання цього методу не є доцільним у нашому випадку. Однак, в інших випадках результат може бути іншим, тому варто спробувати різні методи аугментації та оцінити їх вплив на точність розпізнавання;
- Відображення - показав покращення точності для кожного з класів, особливо для класу "Дерева". Тому можна вважати, що використання цього методу є доцільним для даної задачі класифікації зображень;
- Поворот - показав невелике покращення точності для деяких класів, але не показав суттєвого впливу на точність для інших класів. Тому використання цього методу може бути доцільним при використанні з іншими методами, але не має особливої доцільності в окремому використанні;
- Шум - показав значне покращення точності тільки для класу "Будинки", в інших випадках не було результату, або була незначна рецесія. Тому використання цього методу, у нашому випадку, може бути доцільним для класу "Будинки", але не має особливого сенсу для інших;
- Найбільш результативні - використання найбільш результативних методів аугментації для кожного окремого класу, визначених у минулих експериментах - показало значне покращення точності для всіх класів порівняно з використанням окремих методів. Тому використання цього способу може бути

доцільним, оскільки це допоможе збільшити вибірку та точність розпізнавання моделі.

В ході проведених експериментів з використанням різних методів аугментації (Масштабування, відображення, поворот, шум та всі разом) для різних класів (вода, дерева, будинки, дороги), було визначено, що використання різних методів може позитивно або негативно впливати на точність розпізнавання нейромережі.

Висновок. Шляхом проведення експериментів з використанням різних методів аугментації було встановлено, що точність розпізнавання моделі може позитивно або негативно залежати від методу, використаного для підвищення кількості даних. Використання метода відображення показало покращення точності для всіх класів, тоді як використання методу масштабування та шуму не дали значних результатів у нашому випадку. Використання повороту може бути доцільним, якщо він поєднується з іншими методами.

Подальше покращення результатів передбачає проведення додаткових експериментів зі збільшеною кількістю методів аугментації та їх комбінацій.

Результати, отримані в цьому дослідженні, можуть допомогти у розробці ефективніших і результативних моделей CNN для розпізнавання супутникових зображень.

ЛІТЕРАТУРА

1. Christian L, Lucas T, Ferenc H, Jose C, Andrew C, Alejandro A, Andrew A, Alykhan T, Johannes T, Zehan W, Wenzhe S. Photo-realistic single image super-resolution using a generative adversarial network. arXiv preprint. 2016.
2. Shervin Minaee, Yuri Y Boykov, Fatih Porikli, Antonio J Plaza, Nasser Kehtarnavaz, and Demetri Terzopoulos. Image segmentation using deep learning: A survey. IEEE Transactions on Pattern Analysis and Machine Intelligence, 2021.
3. Chengyue Gong, Dilin Wang, Meng Li, Vikas Chandra, and Qiang Liu. Keypaugment: A simple information-preserving data augmentation approach. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition, pages 1055–1064, 2021.
4. Buda Mateusz, Maki Atsuto, Mazurowski Maciej A. A systematic study of the class imbalance problem in convolutional neural networks. Neural Networks. 2018;106:249–59.
5. Trishul C, Yutaka S, Johnson A, Karthik K. Project adam: building an efficient and scalable deep learning training system. In: Proceedings of OSDI. 2014. P. 571–82.

6. Xiaofeng Z, Zhangyang W, Dong L, Qing L. DADA: deep adversarial data augmentation for extremely low data regime classification. arXiv preprint. 2018.
7. Dong W, Zhou N, Paul JC, Zhang X. Optimized image resizing using seam carving and scaling. ACM Transactions on Graphics (TOG). 2009 Dec 1;28(5):1-0.
8. Russo F. A method for estimation and filtering of Gaussian noise in images. IEEE Transactions on Instrumentation and Measurement. 2003 Sep 23;52(4):1148-54.

REFERENCES

1. Christian L, Lucas T, Ferenc H, Jose C, Andrew C, Alejandro A, Andrew A, Alykhan T, Johannes T, Zehan W, Wenzhe S. Photo-realistic single image super-resolution using a generative adversarial network. arXiv preprint. 2016.
2. Shervin Minaee, Yuri Y Boykov, Fatih Porikli, Antonio J Plaza, Nasser Kehtarnavaz, and Demetri Terzopoulos. Image segmentation using deep learning: A survey. IEEE Transactions on Pattern Analysis and Machine Intelligence, 2021.
3. Chengyue Gong, Dilin Wang, Meng Li, Vikas Chandra, and Qiang Liu. Keepaugment: A simple information-preserving data augmentation approach. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition, pages 1055–1064, 2021.
4. Buda Mateusz, Maki Atsuto, Mazurowski Maciej A. A systematic study of the class imbalance problem in convolutional neural networks. Neural Networks. 2018;106:249–59.
5. Trishul C, Yutaka S, Johnson A, Karthik K. Project adam: building an efficient and scalable deep learning training system. In: Proceedings of OSDI. 2014. P. 571–82.
6. Xiaofeng Z, Zhangyang W, Dong L, Qing L. DADA: deep adversarial data augmentation for extremely low data regime classification. arXiv preprint. 2018.
7. Dong W, Zhou N, Paul JC, Zhang X. Optimized image resizing using seam carving and scaling. ACM Transactions on Graphics (TOG). 2009 Dec 1;28(5):1-0.
8. Russo F. A method for estimation and filtering of Gaussian noise in images. IEEE Transactions on Instrumentation and Measurement. 2003 Sep 23;52(4):1148-54.

Received 11.05.2023.

Accepted 22.05.2023.

Improving deep learning performance by augmenting training data

Satellite image recognition is a crucial application of computer vision that has the potential to be applied in various fields such as disaster management, agriculture, and urban planning. The objective of this study is to determine the optimal amount of input data required and select the most effective methods of augmentation necessary for training a convolutional neural network (CNN) for satellite image recognition.

To achieve this, we perform a series of experiments to investigate the effect of input data quantity on several performance metrics, including model accuracy, convergence, and generalization. Additionally, we explore the impact of various data augmentation techniques, such as rotation, scaling, and flipping, on model performance. The study suggests several strategies for identifying the saturation point and mitigating the effects of overtraining, including early stopping and dropout regularization.

The findings from this study can significantly contribute to the development of more efficient satellite recognition models. Furthermore, they can help improve the performance of existing models, in addition to providing guidance for future research. The study emphasizes the importance of carefully selecting input data and augmentation methods to achieve optimal performance in CNNs, which is fundamental in advancing the field of computer vision.

In addition to the above, the study investigates the potential of transfer learning by pretraining the model on a related dataset and fine-tuning it on the satellite imagery dataset. This approach can reduce the amount of required data and training time and increase model performance.

Overall, this study provides valuable insights into the optimal amount of input data and augmentation techniques for training CNNs for satellite image recognition, and its findings can guide future research in this area.

Солдатенко Дмитро Володимирович – аспірант кафедри інформаційних технологій і систем, Український державний університет науки і технологій.

Гнатушенко Вікторія Володимирівна - д.т.н., професор, завідувача кафедрою інформаційних технологій і систем, Український державний університет науки і технологій.

Soldatenko Dmytro – PhD Aspirant, Department of information technology and systems of the Ukrainian state University of science and technologies.

Hnatushenko Viktorija - Doctor of Engineering's Sciences, Professor, Head of Department of information technology and systems of the Ukrainian state University of science and technologies.

УПРАВЛІННЯ ПОТОКАМИ ДАНИХ В СУЧАСНІЙ ПРОМИСЛОВОСТІ ЗА ДОПОМОГОЮ БЛОКЧЕЙНУ

Анотація. Сучасна світова промисловість проживає трансформацію того, як компанії розробляють, виробляють і розповсюджують товари та послуги, зосереджуючись на більшій ефективності, гнучкості та підлаштуванню під потреби сучасного світу. Ці процеси відбуваються під впливом впровадження в промисловість інформаційних технологій, та часто описуються терміном "Індустрія 4.0", в якому описуються концепції цифровізації, автоматизації та взаємозв'язку у промислових секторах, появи розумних фабрик і об'єднаних ланцюжків поставок, що створює великі об'єми даних для аналізу та обробки, що створює нові виклики у питаннях адаптації інформаційних технологій. Тому є актуальною задача дослідження нових моделей та методів управління потоками даних у інформаційних системах, які можуть поліпшити взаємодію та прискорити адаптацію інформаційних технологій в різних промислових секторах економіки. У роботі розглянуто підхід до побудови інформаційних систем в промисловості та бізнесі за допомогою технології блокчейну та потенціал цієї технології у вирішенні проблем управління потоками даних у сучасній промисловості.

Ключові слова: потоки даних, інформаційні системи, блокчейн, інформатизація, промисловість, ланцюги постачань, Blockchain.

Постановка проблеми. Впровадження інформатизації у промисловості з метою автоматизації та налагодження взаємозв'язку у різних промислових секторах створює нові виклики в побудові інформаційних систем. Інтеграція передових технологій у виробництво, таких як штучний інтелект, Інтернет речей (IoT), хмарні обчислення, аналітика великих даних та інші, створює великі потоки даних, які звичайні інформаційні системи промисловості часто не здатні своєчасно обробити та зреагувати на них, що може створювати затримки та помилки, які можуть призводити до неефективності ланцюгів поставок, затримок у виробництві та комунікації між промисловими секторами.

Крім цього, також існують проблеми прозорості та підзвітності інформаційних систем управління у промисловості, проблеми з безпекою, цілісністю даних, проблем відстеження процесів у реальному часі тощо.

Аналіз останніх досліджень і публікацій. "Індустрія 4.0" – це концепція промислової революції, яка ґрунтується на використанні сучасних технологій та цифрових інновацій у виробничих процесах. Введення концепції "Індустрії 4.0" було покликане покращити конкурентоспроможність європейської промисловості та підвищити продуктивність та якість продукції [1].

Для реалізації інформаційних систем та поліпшення управління потоками даних, які працюють за концепціями «Індустрії 4.0» в цій роботі пропонується використовувати технологію блокчейну.

Блокчейн – це розподілена структура даних, яка реплікується та поширюється між членами мережі. Перша специфікація блокчейну була запропонована разом із цифровою валютою Bitcoin у 2008 році людиною під псевдонімом Сатосі Накамото, щоб розв'язати проблему централізації фінансів навколо банків [2].

Зараз блокчейни використовують переважно у сфері децентралізованих фінансів (DeFi) у вигляді криптовалют та інструментів до них у вигляді обмінників, бірж, аукціонів, банків тощо. Також існують вузькоспеціалізовані іноземні дослідження на тему використання блокчейну у інших сферах інформатизації, крім фінансів. Наприклад, використання блокчейну у риболовецькій галузі [3].

На сьогоднішній день, серед вітчизняних вчених майже немає тих, хто займався б вивченням використання технології Blockchain у інформаційних системах промисловості та логістики.

Тому дослідження методів управління потоками даних в інформаційних системах промисловості та логістики є актуальним, яке спрямоване на економічну та управлінську доцільність використання блокчейну у інформаційних системах промисловості, ланцюгів постачання, логістики тощо.

Метою дослідження є удосконалення процесів автоматизації, підвищення ефективності, зменшення затримок та помилок в інформаційних системах промисловості та ланцюгах постачання за допомогою використання технологій блокчейну в побудові інформаційних систем.

Основна частина. «Індустрія 4.0» або «Промисловість 4.0» — це термін, який використовується для опису так званої «четвертої промислової революції», яка характеризується інтеграцією передових технологій у виробництво та інші галузі. «Індустрію 4.0» часто описують як трансформаційну парадигму, яка стимулює цифровізацію, автоматизацію та взаємозв'язок у промислових сек-

торах, що призводить до появи розумних фабрик і об'єднаних ланцюжків поставок [4].

У класичних інформаційних системах у промисловості та логістиці використовуються реляційні та нереляційні бази даних, такі як MySQL або MongoDB [5]. Доступ до даних отримується через центральний сервер, на якому ця база даних знаходиться. Сервер обробляє запити клієнтів, комп'ютерів, які хочуть отримати або змінити дані у базі даних, на отримання та запису даних. Такий підхід носить назву “клієнт-серверна архітектура” і є загальноновживаним у індустрії [6].

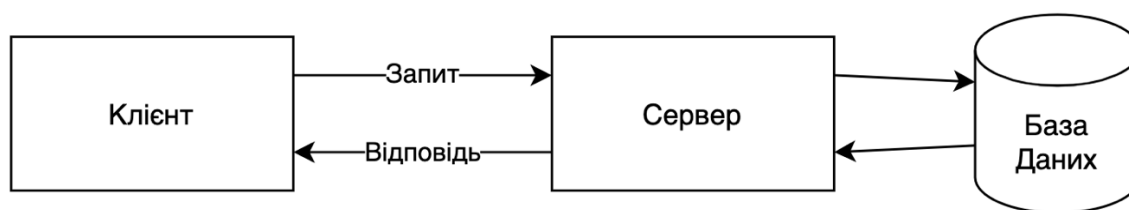


Рисунок 1 - Схема клієнт-серверної архітектури інформаційної системи

Але у великих інформаційних системах ланцюгів постачань та системах промисловості, де використовується постійний обмін даними між різними системами та їхніми клієнтами, для взаємодії у виробництві та логістиці, звичайна клієнт-серверна архітектура може виконувати свою роботу неефективно.

Цей збільшений потік даних створює низку проблем:

1. Ефективна аналітика даних.
2. Безпека, цілісність даних та конфіденційність даних в інформаційній системі.
3. Взаємодія з іншими системами та масштабування потоків даних у взаємодії між різними компонентами промислової екосистеми.

Ці проблеми можуть не дозволяти промисловості та бізнесу повністю реалізувати переваги концепту «Індустрії 4.0». Для рішення цих проблем в цій роботі пропонується взяти технологію блокчейну для управління потоками даних у інформаційних системах.

Технологія блокчейну на сьогоднішній день отримала розповсюдження у сфері електронних фінансів, але блокчейн також може знайти використання у логістиці, документообігу та засобі управління даними в інформаційних системах. Розглянемо які механізми дозволяють це зробити.

Блокчейн – це збудований за певними правилами безперервний послідовний ланцюжок блоків (зв'язковий список), що містять інформацію [7]. Зв'язок

між блоками забезпечується не тільки нумерацією, але й тим, що кожен блок містить власну хеш-суму і хеш-суму попереднього блоку. Зміна будь-якої інформації в блоці змінить його хеш-суму. Щоб відповідати правилам побудови ланцюжка, зміни хеш-суми потрібно буде записати в наступний блок, що викличе зміни його власної хеш-суми.

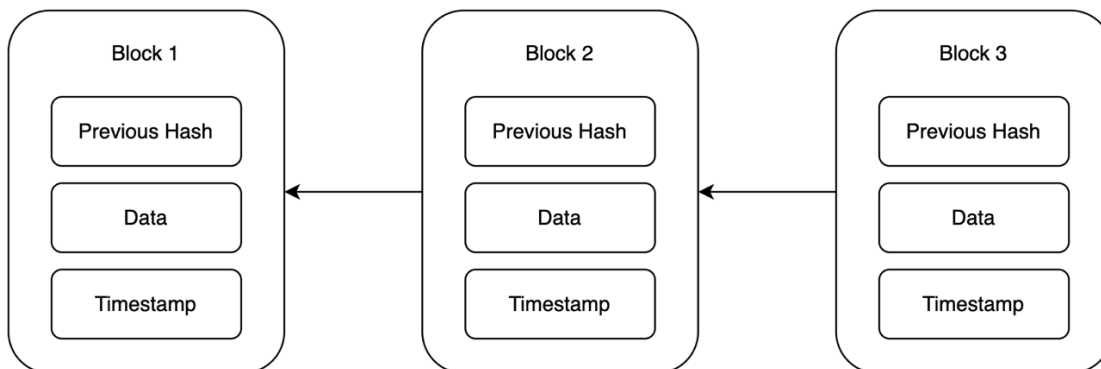


Рисунок 2 - Структурна схема блокчейну

Технологічною основою, якою будується блокчейн, є Distributed Ledger Technology (DLT). Ця технологія є розподіленою базою даних, яка зберігає записи транзакцій у кількох вузлах мережі, а чи не в єдиній централізованій базі даних. DLT не має центрального апарату, який контролює базу даних, і всі вузли мають однаковий статус.

DLT використовує криптографію для забезпечення безпеки та надійності бази даних. Кожна транзакція записується в блок ланцюжка блоків (блокчейн) і містить хеш попереднього блоку, що забезпечує безперервність ланцюжка. Крім того, DLT забезпечує можливість створення розумних контрактів, які автоматично виконуються під час виконання певних умов.

Одним із основних елементів DLT є хеш-функція, яка використовується для створення унікального ідентифікатора для блоку даних. Хеш-функція може бути представлена такою формулою:

$$H(x) = y,$$

де $H(x)$ — хеш-функція, застосована до даних x , а y — результуюче хеш-значення.

Іншим важливим елементом DLT є консенсус-протокол, який використовується для досягнення згоди між учасниками мережі щодо поточного стану блокчейну. Консенсус-протоколи можуть бути представлені різними математичними моделями, такими як модель прийняття рішень або ігор.

Алгоритми консенсусу використовуються в технології блокчейн, щоб гарантувати, що всі вузли в мережі погоджуються щодо стану блокчейну. Формула для алгоритму консенсусу може змінюватися залежно від конкретного алгоритму, який використовується, але, як правило, вона включає набір правил, які визначають, як вузли в мережі взаємодіють і узгоджують стан блокчейну.

Так, наприклад, блокчейн Ethereum перейшов з алгоритму Proof-of-Work (PoW), який вимагав неекологічний процес «майнінгу» блоків за допомогою математичних розрахунків, на алгоритм Proof-of-Stake (PoS).

В алгоритмі консенсусу PoS ймовірність того, що вузол буде обрано для створення нового блоку, пропорційна сумі частки, яку він займає в мережі.

Формула для розрахунку ймовірності того, що вузол буде обраний творцем блоку:

$$P(node) = \left(\frac{S(node)}{\sum S(i)} \right)^k,$$

де $P(node)$ – це ймовірність того, що вузол і буде обрано творцем блоку, $S(node)$ – це розмір частки, яку має вузол i , а $\sum S(i)$ – це загальна сума частки, яку має мати мережа, змінна k є константою, яка визначає рівень випадковості в процесі вибору та зазвичай встановлюється протоколом мережі.

Для перевірки цілісності транзакцій та зберігання даних у блокчейні використовується структура даних Дерево Меркла. Воно отримало свою назву на честь Ральфа Меркла, який першим описав цю структуру у 1987 році.

У блокчейні, дерево Меркла використовується для забезпечення цілісності блоків, що становлять ланцюжок блоків. Кожен блок містить набір транзакцій, а дерево Меркла дозволяє швидко перевіряти, чи не були змінені транзакції в блоці.

Для побудови дерева Меркла кожен блок ділиться на кілька частин, наприклад, на транзакції, які входять в блок. Потім кожна частина хешується з використанням криптографічної хеш-функції, наприклад, SHA-256 і результати об'єднуються в пари. Кожна пара потім хешується знову, і цей процес повторюється до тих пір, поки не залишиться тільки один хеш, який називається кореневим хешом дерева Меркла.

В «Індустрії 4.0» дерево Меркла може бути використане для забезпечення цілісності даних у виробничих системах, таких як системи управління якістю або системи моніторингу виробничих процесів. Також воно може бути використане для забезпечення безпеки даних у системах керування постачанням, наприклад, для перевірки справжності документів та контрактів.

Припустимо, що ми маємо блок даних, який потрібно захистити від змін, і ми хочемо використовувати дерево Меркла для цієї мети. Ми починаємо з хешування кожного окремого елемента даних у блоці (наприклад, транзакцій), щоб отримати їхнє хеш-значення. Потім ми об'єднуємо хеш-значення попарно і хешуємо їх разом, щоб отримати нові хеш-значення, які є корінням піддерев дерева Меркла.

Цей процес повторюється до тих пір, поки всі хеш-значення не будуть об'єднані в одне хеш-значення, яке є коренем всього дерева Меркла.

Математично це може бути представлено такою формулою:

$$H = H_n(H_{n-1}(H_{n-2}(\dots(H_2(H_1(T_{x_1}, T_{x_2})).\dots)T_{x_{n-1}}), T_{x_n})),$$

де H – підсумкове хеш-значення кореня дерева Меркла, H_n – функція хешування, T_{x_n} – окремий елемент даних, які потрібно захистити.

Таким чином, дерево Меркла дозволяє ефективно захищати великі обсяги даних, забезпечуючи цілісність їх зберігання та передачі. У блокчейні дерево Меркла використовується для забезпечення цілісності блоків та підтвердження правильності хеш транзакцій, що є важливим елементом безпеки системи.

Для авторизації та додаткового захисту транзакцій в блокчейн, кожна транзакція додатково захищена цифровим підписом.

Цифровий підпис – це математичний метод, який використовується для автентифікації особи відправника повідомлення або транзакції в DLT. Цифрові підписи базуються на криптографії з відкритим ключем, створюються за допомогою закритого ключа та перевіряються за допомогою відкритого ключа.

Описані вище властивості технології Blockchain та Distributed Ledger Technology дозволяють змінити підхід до побудови інформаційних систем у логістиці та промисловості. У контексті «Індустрії 4.0» блокчейн можна використовувати різними способами, щоб створити нові форми довіри та співпраці між різними сторонами в екосистемі виробництва та ланцюгів постачання.

Так, використовуючи алгоритми дерева Меркла, цифрового підпису, розподіленої програмованої бази даних та валідації транзакцій за допомогою алгоритмів консенсусу, тобто технологій, на яких базується блокчейн та Distributed Ledger Technology, можна створити інформаційні системи для промисловості та ланцюгів постачання, які можуть використовуватися в такий спосіб:

1. Управління ланцюгом поставок. Блокчейн можна використовувати для створення захищеного від несанкціонованого втручання запису кожної транза-

кції, що відбувається в ланцюзі поставок, від сировини до готової продукції. Це може забезпечити більшу прозорість і відстежуваність, зменшити ризик шахрайства та підробок, а також допомогти компаніям швидше виявляти та вирішувати проблеми.

2. Контроль якості та управління гарантіями. Блокчейн можна використовувати для відстеження якості продуктів протягом усього життєвого циклу, а також для автоматичного підтвердження гарантії та обробки претензій на основі попередньо визначених правил смарт-контрактами.

3. Децентралізоване виробництво. Блокчейн можна використовувати для створення децентралізованих виробничих мереж, де різні сторони можуть співпрацювати над проектуванням і виробництвом, використовуючи спільні дані та ресурси, без необхідності центрального органу чи посередника.

Висновки. За допомогою впровадження алгоритмів, на яких базується технологія блокчейну, а саме за допомогою використання Дерева Меркла для забезпечення цілісності та безпеки великих обсягів даних, та їх подальшої передачі та зберігання, цифрового підпису для автентифікації транзакцій, з використанням відкритого та закритого ключів, та за допомогою алгоритмів консенсусу у рамках децентралізованого зберігання даних у DLT, удосконалюються процеси автоматизації та ефективності в управлінні даними в інформаційних системах промисловості та ланцюгів постачання шляхом створення нових форм довіри та співпраці в «Індустрії 4.0», надаючи безпечний і прозорий спосіб до зберігання та обміну даними, який зменшує затримки та помилки в інформаційних системах промисловості та ланцюгах постачання.

Результати дослідження управління ланцюгами постачання на основі використання алгоритмів блокчейну корисні при проектуванні інформаційних систем промисловості та систем моніторингу, що дозволить створювати спільні реєстри та бази даних для потоків даних, до яких мають доступ та довіру усі сторони промислового виробництва та логістики, що зменшує помилки, затримки та неправомірні дії з потоками даних в інформаційних системах, та покращує процеси автоматизації та ефективність в управлінні потоками даних в інформаційних системах, для їхньої обробки, моніторингу, аналізу та подальшого зберігання.

ЛІТЕРАТУРА / REFERENCE

1. H. Lasi, P. Fettke, H.G. Kemper, T. Feld, M. Hoffmann. Industry 4.0. Business & information systems engineering, 6, pp.239-242, 2014.

2. Satoshi Nakamoto, Bitcoin: A Peer-to-Peer Electronic Cash System, 2008, [Електронний ресурс]. – Режим доступу: <https://bitcoin.org/bitcoin.pdf>
3. Magnus Mathisen, The Application of Blockchain Technology in Norwegian Fish Supply Chains. A Case Study // Norwegian University of Science and Technology, 2018, [Електронний ресурс]. – Режим доступу: <https://ntnuopen.ntnu.no/ntnu-xmlui/handle/11250/2561323>
4. L. Gamio, P. S. Goodman, New York Times. How the Supply Chain Crisis Unfolded. [Електронний ресурс]. – Режим доступу: <https://www.nytimes.com/interactive/2021/12/05/business/economy/supply-chain.html>
5. H. Matallah, G. Belalem, K. Bouamrane. Comparative study between the MySQL relational database and the MongoDB NoSQL database. International Journal of Software Science and Computational Intelligence (IJSSCI) 13.3 (2021): 38-63.
6. H.S. Oluwatosin, Client-server model. IOSR Journal of Computer Engineering 16.1 (2014): 67-71.
7. M. Nofer, P. Gomber, O. Hinz, D. Schiereck. Blockchain. Business & Information Systems Engineering, 2017, 59, 183-187.

Received 11.05.2023.

Accepted 23.05.2023.

Management of data flows in modern industry using blockchain

Recent research and publications. "Industry 4.0" is a concept of the industrial revolution, which is based on the use of modern technologies and digital innovations in production and distribution processes. The introduction of the concept of "Industry 4.0" was designed to improve the competitiveness of European industry and increase productivity and product quality. A blockchain is a distributed data structure that is replicated and distributed among network members.

The purpose of the study is to improve automation processes, increase efficiency, reduce delays and errors in information systems of industry and supply chains by using blockchain technologies in the construction of information systems.

Main material of the study. The paper makes an analysis of approaches and algorithms to data management in "Industry 4.0" information systems. Blockchain algorithms are compared to classical approach with other databases in the client-server architecture.

Conclusions. By implementing algorithms based on blockchain technology, namely by using the Merkle Tree, digital signature technology, and by using consensus algorithms in the framework of decentralized data storage in Distributed Ledger Technology, the processes of automation and efficiency in data flow management are improved,

providing a secure and transparent way to store and share data that reduces delays and errors in industry information systems and supply chains.

Гнатушенко Вікторія Володимирівна – доктор технічних наук, професор, завідувач кафедри інформаційних технологій і систем Українського державного університету науки і технологій.

Ситник Роман Сергійович – аспірант кафедри інформаційних технологій і систем Українського державного університету науки і технологій.

Hnatushenko Viktoriia – Doctor of engineering's sciences, Professor, Head of Department of Information Technologies and Systems, Ukrainian State University of Science and Technology.

Sytnyk Roman – Postgraduate Student, Department of Information Technologies and Systems, Ukrainian State University of Science and Technology.

Х.В. Лип'яніна-Гончаренко

ІНТЕЛЕКТУАЛЬНИЙ МЕТОД ВИБОРУ ЛОКАЦІЇ ДЛЯ СТАРТУ БІЗНЕСУ В РОЗУМНОМУ МІСТІ

Анотація. Вибір оптимальної локації для бізнесу в розумних містах є складною задачею. У цьому дослідженні розроблено інтелектуальний метод, використовуючи машинне навчання, для швидкого та точного вибору локації. Результати цього методу, зможе допомогти підприємцям знаходити оптимальні місця для старту бізнесу, забезпечуючи задоволення клієнтів і збільшення прибутку. Інтелектуальний метод є потужним інструментом для вирішення проблем вибору локації в розумних містах.

Ключові слова: інтелектуальний метод, машинне навчання, сегментація, розпізнавання образів, старт бізнесу, розумне місто.

Постановка проблеми полягає в тому, що вибір правильної локації для старту бізнесу є критичним фактором успіху. Вибір неправильної локації може призвести до недостатнього трафіку клієнтів, високої конкуренції, невідповідності демографії та інших проблем, які можуть призвести до невдачі бізнесу.

В контексті розумного міста, де використовуються технології для підвищення якості життя мешканців та ефективності міських служб, використання цих технологій для вибору локації бізнесу може бути важливим інструментом.

Проте, використання цих технологій вимагає знань та навичок в області аналітики даних, машинного навчання та штучного інтелекту. Крім того, потрібно зібрати та обробити великі обсяги даних з різних джерел, що може бути викликом.

Таким чином, проблема полягає в розробці інтелектуального методу вибору локації для старту бізнесу в розумному місті, який би використовував доступні дані та технології для визначення найбільш оптимальної локації для нового бізнесу.

Аналіз останніх досліджень і публікацій. Дослідження [1] досліджує взаємозв'язок між розвитком штучного інтелекту та розвитком міст. Воно розглядає поняття "урбаністичного штучного інтелекту", а також досліджує вплив розвитку штучного інтелекту на управління міськими сервісами та урбаністичним управлінням.

У цьому дослідженні [2] розглядається підхід еволюційного планування для розвитку розумного міста. Воно базується на зусиллях міста Салоніки для створення стратегії розумного міста, яка сприяє економічній, екологічній та соціальній сталості міста. Процес планування розумного міста відрізняється від традиційного планування, оскільки він поєднує технології, залучення користувачів та можливості, які на початку процесу планування не є чіткими. У дослідженні [3] розглядаються фактори, які відрізняють політику розвитку розумних міст і надають зрозумілий огляд стратегічних виборів, які виникають при розробці такої стратегії. Дослідження включає аналіз чотирьох стратегічних виборів з просторовим посиленням та представляє переваги і недоліки кожного з них. Також розглядаються приклади стратегій розумних міст з усього світу. У заключній частині дослідження надаються рекомендації щодо розвитку розумних міст на основі отриманих висновків.

Дослідження [4] розглядає концепцію "розумного міста" як область, де існують різні моделі та підходи. Основною метою впровадження розумного міста є стимулювання економічного зростання та соціального розвитку за допомогою інноваційних технологій. Дослідження аналізує ролі різних акторів екосистеми, зацікавлених сторін та соціально-економічних та політичних агентів у створенні економічної цінності та вирішенні суспільних проблем.

Це дослідження [5] пропонує інтелектуальний бізнес-додаток, який використовує методи машинного навчання для оцінки локації бізнесу. Система збирає ключові характеристики, навчає модель прогнозування та пропонує оптимальні локації для бізнесу.

На відміну від аналога [5] розроблений інтелектуальний метод вибору локації для старту бізнесу в розумному місті включає використання відеоаналітики, мобільної аналітики та сервісів геолокації для збору даних.

Метою дослідження є розробка інтелектуального методу вибору локації для старту бізнесу в розумному місті. Цей метод має на меті використовувати великі обсяги даних, зібраних з різних джерел, для визначення найбільш оптимальних місць для відкриття нового бізнесу. Дослідження також має на меті

вивчити, як різні технології, такі як відеоаналітика, мобільна аналітика та сервісів геолокації, можуть бути використані для збору та аналізу цих даних.

Викладення основного матеріалу дослідження. Інтелектуальний метод вибору локації для старту бізнесу в розумному місті - це комплексний підхід, який використовує технології великих даних, штучний інтелект та машинне навчання для визначення оптимального місця для відкриття нового бізнесу. Ось основні кроки цього методу та його структура (Рис.1):

Крок 1. Збір даних. Цей процес включає збір великих обсягів даних з різних джерел:

1.1. Відеокамери: Камери відеоспостереження, розміщені в магазинах та на вулицях міста, можуть надавати цінну інформацію про поведінку покупців та пішоходів, їх стать та вікова категорія. Вони можуть виявити, які магазини або вулиці приваблюють найбільше людей, в який час дня та дні тижня.

1.2. Мобільні додатки та Сервіси геолокації: З мобільного додатку розумного міста, можна використовувати вбудовані інструменти аналітики для збору даних про поведінку користувачів. Деякі сервіси, такі як Google Maps чи відповідні мобільні додатки, надають API чи геолокацію, які дозволяють відстежувати місцезнаходження користувачів (з їхньою згодою). Це може бути корисно для визначення популярних місць та маршрутів.

Крок 2. Обробка та аналіз даних. Зібрані дані потім обробляються та аналізуються за допомогою алгоритмів машинного навчання. Це може включати розпізнавання образів, визначення шаблонів, виявлення трендів та визначення ключових впливових факторів.

2.1. Обробка даних.

2.1.1. Розпізнавання образів [6,7]: Використовуючи алгоритми машинного навчання, дані з відеоаналітики проходять процес розпізнавання образів. Це може включати в себе визначення облич, визначення об'єктів, відстеження руху та інші техніки, які допомагають визначити стать та вік людей.

2.1.2. Попередня обробка числових масивів даних [8]. Обробка даних включає в себе перевірку на наявність відсутніх або некоректних даних, нормалізацію даних для забезпечення їх сумісності та видалення шуму або викидів, які можуть спотворити результати аналізу.

2.2. Сегментація даних [9]: Після обробки даних вони аналізуються за допомогою різних алгоритмів машинного навчання, що включає в себе класифікацію, кластеризацію.

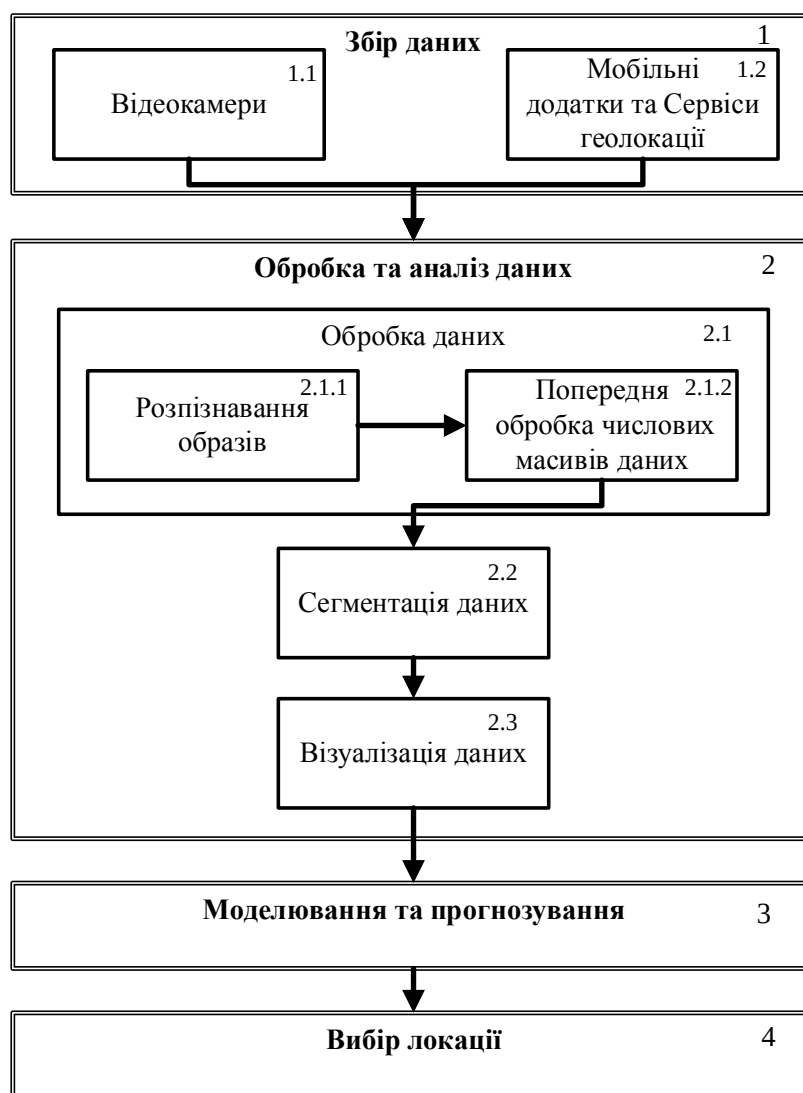


Рисунок 1 - Схематичне представлення інтелектуального методу вибору локації для старту бізнесу в розумному місті

2.3. Візуалізація даних: Для кращого розуміння результатів аналізу даних їх часто візуалізують. Візуалізація може включати в себе графіки, діаграми, теплові карти та інші графічні зображення, які допомагають інтерпретувати складні набори даних.

Крок 3. **Моделювання та прогнозування.** На основі аналізу даних створюються моделі, які можуть прогнозувати потенційну успішність бізнесу в різних локаціях для певних сегментів [10,11]. Це може включати визначення потенційного обсягу продажів, оцінку конкуренції та інші важливі метрики.

Крок 4. **Вибір локації.** На основі результатів моделювання та прогнозування вибирається найбільш оптимальна локація для відкриття нового бізнесу.

Цей метод використовує передові технології для визначення найкращої локації для нового бізнесу, що забезпечує високу ймовірність успіху.

Отже, розроблений інтелектуальний метод вибору локації для старту бізнесу в розумному місті, як і аналог [5] використовує інтелектуальний підхід для вибору оптимальної локації бізнесу в розумному місті, також обидва методи використовують аналіз даних та оцінювання характеристик для рекомендацій локацій. Однак, розроблений метод фокусується безпосередньо на виборі локації для старту бізнесу, враховуючи особливості підприємства, тоді як аналог [5] зосереджується на загальному розумному бізнес-додатку, що оцінює оптимальні локації для ресторанів.

Висновки. Використання інтелектуального методу вибору локації для старту бізнесу в розумному місті може допомогти зробити складну та довгострокову рішення більш об'єктивним та ефективним. Застосування методів машинного навчання та аналізу даних дозволяє зібрати ключові характеристики для певного бізнесу та навчити прогностичну модель для рекомендацій локацій.

Розроблений інтелектуальний метод вибору локації для старту бізнесу в розумному місті пропонує використання інтелектуального підходу для вибору оптимальної локації бізнесу в розумному місті. Метод, на відміну від аналогу [5], фокусується на безпосередньому виборі локації для старту бізнесу, враховуючи особливості підприємства та використовуючи відеоаналітику, мобільну аналітику та сервіси геолокації для збору даних.

Однак, для подальшого розвитку цього методу та вдосконалення його точності, необхідні подальші дослідження та випробування на різних кейсах. Це дозволить розширити розуміння процесу вибору локації для бізнесу в розумних містах і покращити ефективність таких рішень для підприємців.

ЛІТЕРАТУРА

1. Cugurullo F. Urban Artificial Intelligence: From Automation to Autonomy in the Smart City [Electronic resource] / Federico Cugurullo // Frontiers in Sustainable Cities. — 2020. — Vol. 2. — Mode of access: <https://doi.org/10.3389/frsc.2020.00038> (date of access: 26.05.2023). — Title from screen.
2. Smart City Planning from an Evolutionary Perspective [Electronic resource] / N. Komninos [et al.] // Journal of Urban Technology. — 2018. — Vol. 26, no. 2. — P. 3—20. — Mode of access: <https://doi.org/10.1080/10630732.2018.1485368> (date of access: 26.05.2023). — Title from screen.

3. Angelidou M. Smart city policies: A spatial approach [Electronic resource] / Margarita Angelidou // Cities. — 2014. — Vol. 41. — P. S3—S11. — Mode of access: <https://doi.org/10.1016/j.cities.2014.06.007> (date of access: 26.05.2023). — Title from screen.
4. Sarma S. Civic entrepreneurial ecosystems: Smart city emergence in Kansas City [Electronic resource] / Sumita Sarma, Sanwar A. Sunny // Business Horizons. — 2017. — Vol. 60, no. 6. — P. 843—853. — Mode of access: <https://doi.org/10.1016/j.bushor.2017.07.010> (date of access: 26.05.2023). — Title from screen.
5. A Smart City Application: Business Location Estimator Using Machine Learning Techniques [Electronic resource] / Tugce Bilen [et al.] // 2018 IEEE 20th International Conference on High Performance Computing and Communications; IEEE 16th International Conference on Smart City; IEEE 4th International Conference on Data Science and Systems (HPCC/SmartCity/DSS), Exeter, United Kingdom, 28—30 June 2018. — [S. l.], 2018. — Mode of access: <https://doi.org/10.1109/hpcc/smartcity/dss.2018.00219> (date of access: 26.05.2023). — Title from screen.
6. Recognizing Age-Separated Face Images: Humans and Machines [Electronic resource] / Daksha Yadav [et al.] // PLoS ONE. — 2014. — Vol. 9, no. 12. — P. e112234. — Mode of access: <https://doi.org/10.1371/journal.pone.0112234> (date of access: 26.05.2023). — Title from screen.
7. Ng C. B. Recognizing Human Gender in Computer Vision: A Survey [Electronic resource] / Choon Boon Ng, Yong Haur Tay, Bok-Min Goi // Lecture Notes in Computer Science. — Berlin, Heidelberg, 2012. — P. 335—346. — Mode of access: https://doi.org/10.1007/978-3-642-32695-0_31 (date of access: 26.05.2023). — Title from screen.
8. Software System for Processing and Visualization of Big Data Arrays [Electronic resource] / Fedorova Nataliia [et al.] // Advances in Computer Science for Engineering and Education. — Cham, 2022. — P. 324—336. — Mode of access: https://doi.org/10.1007/978-3-031-04812-8_28 (date of access: 26.05.2023). — Title from screen.
9. Intelligent Method for Forming the Consumer Basket [Electronic resource] / Khrystyna Lipianina-Honcharenko [et al.] // Communications in Computer and Information Science. — Cham, 2022. — P. 221—231. — Mode of access: https://doi.org/10.1007/978-3-031-16302-9_17 (date of access: 26.05.2023). — Title from screen.
10. Method of choosing a competitive product based on the emotional color of the calls [Electronic resource] / Khrystyna Lipianina-Honcharenko [et al.] // Herald of

Khmelnitskyi National University. — 2021. — Vol. 303, no. 6. — P. 86—88. — Mode of access: <https://doi.org/10.31891/2307-5732-2021-303-6-86-88> (date of access: 26.05.2023). — Title from screen.

11. Multiple Regression Method for Analyzing the Tourist Demand Considering the Influence Factors [Electronic resource] / Viktor Krylov [et al.] // 2019 10th IEEE International Conference on Intelligent Data Acquisition and Advanced Computing Systems: Technology and Applications (IDAACS), Metz, France, 18—21 September 2019. — [S. l.], 2019. — Mode of access:

<https://doi.org/10.1109/idaacs.2019.8924461> (date of access: 26.05.2023). — Title from screen.

REFERENCES

1. Cugurullo F. Urban Artificial Intelligence: From Automation to Autonomy in the Smart City [Electronic resource] / Federico Cugurullo // Frontiers in Sustainable Cities. — 2020. — Vol. 2. — Mode of access:

<https://doi.org/10.3389/frsc.2020.00038> (date of access: 26.05.2023). — Title from screen.

2. Smart City Planning from an Evolutionary Perspective [Electronic resource] / N. Komninos [et al.] // Journal of Urban Technology. — 2018. — Vol. 26, no. 2. — P. 3—20. — Mode of access: <https://doi.org/10.1080/10630732.2018.1485368> (date of access: 26.05.2023). — Title from screen.

3. Angelidou M. Smart city policies: A spatial approach [Electronic resource] / Margarita Angelidou // Cities. — 2014. — Vol. 41. — P. S3—S11. — Mode of access: <https://doi.org/10.1016/j.cities.2014.06.007> (date of access: 26.05.2023). — Title from screen.

4. Sarma S. Civic entrepreneurial ecosystems: Smart city emergence in Kansas City [Electronic resource] / Sumita Sarma, Sanwar A. Sunny // Business Horizons. — 2017. — Vol. 60, no. 6. — P. 843—853. — Mode of access:

<https://doi.org/10.1016/j.bushor.2017.07.010> (date of access: 26.05.2023). — Title from screen.

5. A Smart City Application: Business Location Estimator Using Machine Learning Techniques [Electronic resource] / Tugce Bilen [et al.] // 2018 IEEE 20th International Conference on High Performance Computing and Communications; IEEE 16th International Conference on Smart City; IEEE 4th International Conference on Data Science and Systems (HPCC/SmartCity/DSS), Exeter, United Kingdom, 28—30 June 2018. — [S. l.], 2018. — Mode of access:

<https://doi.org/10.1109/hpcc/smartcity/dss.2018.00219> (date of access: 26.05.2023). — Title from screen.

6. Recognizing Age-Separated Face Images: Humans and Machines [Electronic resource] / Daksha Yadav [et al.] // PLoS ONE. — 2014. — Vol. 9, no. 12. — P. e112234. — Mode of access: <https://doi.org/10.1371/journal.pone.0112234> (date of access: 26.05.2023). — Title from screen.
7. Ng C. B. Recognizing Human Gender in Computer Vision: A Survey [Electronic resource] / Choon Boon Ng, Yong Haur Tay, Bok-Min Goi // Lecture Notes in Computer Science. — Berlin, Heidelberg, 2012. — P. 335—346. — Mode of access: https://doi.org/10.1007/978-3-642-32695-0_31 (date of access: 26.05.2023). — Title from screen.
8. Software System for Processing and Visualization of Big Data Arrays [Electronic resource] / Fedorova Nataliia [et al.] // Advances in Computer Science for Engineering and Education. — Cham, 2022. — P. 324—336. — Mode of access: https://doi.org/10.1007/978-3-031-04812-8_28 (date of access: 26.05.2023). — Title from screen.
9. Intelligent Method for Forming the Consumer Basket [Electronic resource] / Khrystyna Lipianina-Honcharenko [et al.] // Communications in Computer and Information Science. — Cham, 2022. — P. 221—231. — Mode of access: https://doi.org/10.1007/978-3-031-16302-9_17 (date of access: 26.05.2023). — Title from screen.
10. Method of choosing a competitive product based on the emotional color of the calls [Electronic resource] / Khrystyna Lipianina-Honcharenko [et al.] // Herald of Khmelnytskyi National University. — 2021. — Vol. 303, no. 6. — P. 86—88. — Mode of access: <https://doi.org/10.31891/2307-5732-2021-303-6-86-88> (date of access: 26.05.2023). — Title from screen.
11. Multiple Regression Method for Analyzing the Tourist Demand Considering the Influence Factors [Electronic resource] / Viktor Krylov [et al.] // 2019 10th IEEE International Conference on Intelligent Data Acquisition and Advanced Computing Systems: Technology and Applications (IDAACS), Metz, France, 18—21 September 2019. — [S. l.], 2019. — Mode of access: <https://doi.org/10.1109/idaacs.2019.8924461> (date of access: 26.05.2023). — Title from screen.

Received 11.05.2023.

Accepted 24.05.2023.

Intellectual method for business location selection in smart cities

The relevance of the topic lies in the complexity of selecting a location for starting a business in smart cities, as it requires analyzing a large amount of data and considering various factors such as population, competition, infrastructure, and other parameters. The use of an intelligent method based on machine learning enables the collection, processing, and analysis of large volumes of data for accurate location assessment and

providing recommendations to entrepreneurs. This enhances the decision-making process, ensures more informed choices, and increases the chances of business success in a smart city.

The problem statement involves the need to expedite the process of selecting an optimal location for business placement in a smart city. This task is challenging and long-term, requiring the analysis of extensive data and consideration of various factors that impact business success, such as geographical position, competition, potential customer base, and other relevant aspects. It is also crucial to provide entrepreneurs with fast access to information and precise recommendations to make informed decisions regarding their business location. Solving this problem will facilitate efficient resource utilization and ensure business success in a smart city.

The purpose of the study is to develop an intelligent method for choosing a location for starting a business in a smart city. This method aims to use large amounts of data collected from various sources to determine the most optimal locations for starting a new business. The method is based on existing machine learning techniques such as image recognition, data preprocessing, classification, and clustering of numerical data.

Results and key conclusions. A method has been developed, the implementation of which will allow recommending optimal locations for business in smart cities. This will help to increase customer satisfaction, improve the quality of life and increase the profit of entrepreneurs. The intelligent method is a powerful tool for solving the problems of choosing a location for starting a business in smart cities.

Keywords: intelligent method, machine learning, segmentation, pattern recognition, business startup, smart city.

Ліп'яніна-Гончаренко Христина Володимирівна - кандидат технічних наук, доцент, доцент кафедри інформаційно-обчислювальних систем і управління Західноукраїнського національного університету.

Khrystyna Lipianina-Honcharenko - Associate professor, Ph.D. in information technologies, Department for Information Computer Systems and Control, West Ukrainian National University.

МЕХАНІЗМИ ТА МЕТОДИ ФІШИНГУ ЯК ПЕРШОГО КРОКУ ДО ОТРИМАННЯ ДОСТУПУ

Анотація. Розглянуто фішинг – техніку надсилання фішингових повідомлень. Аналіз зроблено на підставі даних у відкритому доступі. Проаналізовано процес фішингової атаки, та досліджено технічні вектори того, як користувачі стають жертвами атаки. Також розглянуто існуючі параметри фішингових атак та відповідні підходи до запобігання.

Ключові слова: фішинг, кібербезпека, багатofакторна аутентифікація, соціальна інженерія.

1. Проблема фішингу. Найпоширенішим підходом для запуску фішингової атаки є надсилання фішингового електронного листа. Згідно зі звітом Verizon [1] про розслідування витоку даних за 2023 рік, соціальна інженерія — це тактика, використання якої збільшилось з 14% 2022 року до 21% у 2023 році (рисунок 1).

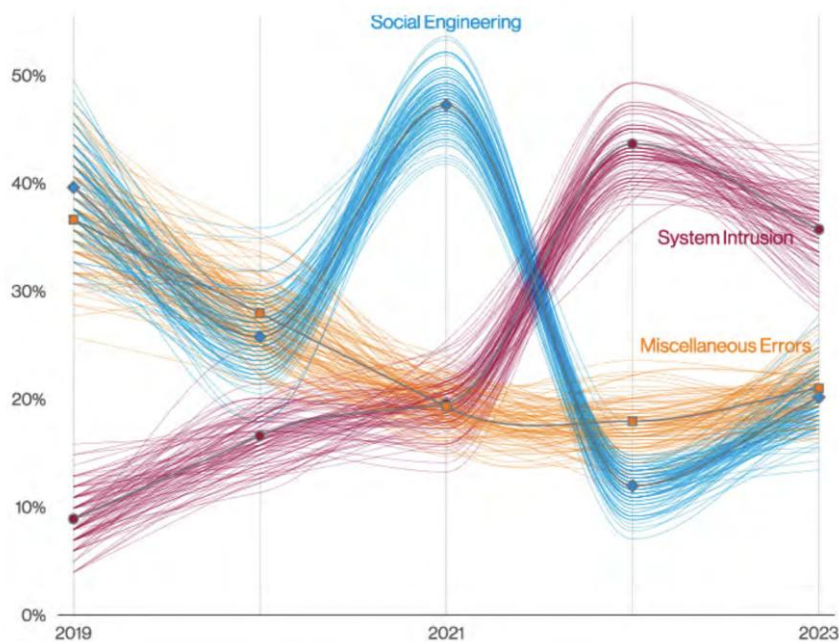


Рисунок 1 - Соціальна інженерія поміж інших атак [1]

Це зростання в першу чергу представлено фішинговими атаками, які виявилися в 18% зломів, і сценаріями претексту (4%). Половина випадків витоку

даних, у 2023 році була зафіксована з використанням життєвого циклу атаки соціальної інженерії: «розслідування», «побудова стосунків», «гра» та «вихід». У більшості випадків фішингові атаки соціальної інженерії відбувалися через електронну пошту, а саме 98% з зареєстрованих випадках, тоді як лише 2% за допомогою інших комунікацій, такі як телефон, соціальні мережі або внутрішній додаток для обміну повідомленнями, в якому деякі люди можуть розслабитися [1].

Атаки соціальної інженерії часто дуже ефективні та надзвичайно прибуткові для кіберзлочинців. Можливо, саме тому атаки компрометації бізнес-електронної пошти (BEC, Business Email Compromise) що, по суті, є атаками з використанням претекстів, майже подвоїлися з усього набору даних про інциденти, як можна побачити на малюнку нижче, і тепер становлять понад 50% інцидентів у шаблоні соціальної інженерії [1].

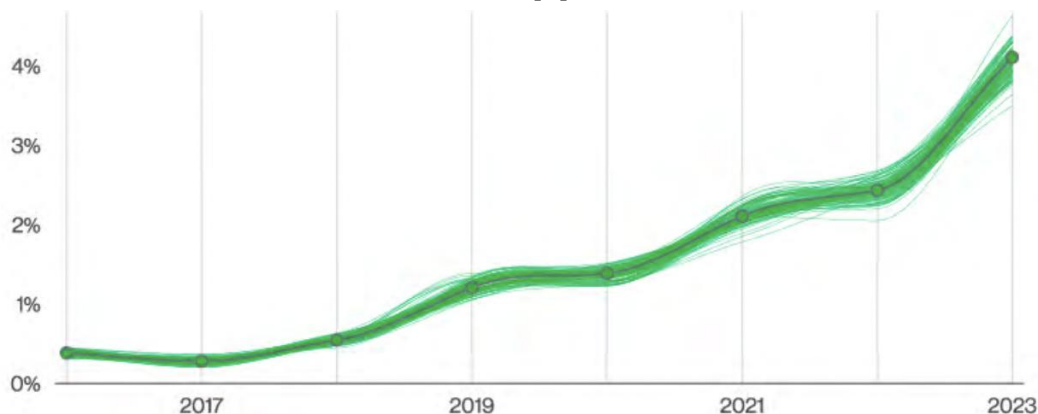


Рисунок 2 - Pretexting інциденти з часом [1]

Фішинг залишається дуже успішним способом для зловмисників отримати доступ до приватних даних і комп'ютерних систем. Що відбувається після цього першого електронного листа, розвиток може бути різним. Атаки зазвичай здійснюються двома основними шляхами. Найчастіше, якщо зловмисники вимагають або виманюють облікові дані та отримують їх, далі використовуватимуть ці облікові дані для доступу до папки «Вхідні» користувача (в 32% випадків). Інший розвиток зловмисники можуть вигадати достовірну історію (хоча й вигадану). Ціллю є переконання когось виконати бажання зловмисника. Переконання когось змінити банківський рахунок для заявленого одержувача, зустрічається в 56% випадків. Звичайно, також можна використовувати комбінацію тактик. Зловмисники можуть використати отриманий доступ до папки "Вхідні" користувача, щоб знайти ланцюжок електронної пошти, який вони

можуть захопити, або шукати в адресній книзі жертви людей, які можуть стати ціллю для подальших дій.

Фішинговий електронний лист часто додає зовнішнє посилання, яке перенаправляє жертв на підроблений веб-сайт і вимагає від жертви ввести свою особисту інформацію. Робочий процес фішингової атаки проілюстрували Джайн і Гупта [2], як показано на рисунку нижче. У цьому випадку фішинговий веб-сайт є незаконним клоном законного веб-сайту та має високу візуальну схожість. Ці фішингові сайти завжди включають деякі поля введення, наприклад текстове поле, щоб вимагати від жертви введення конфіденційної інформації. Інформація надсилається фішеру, коли жертви надсилають свої дані.

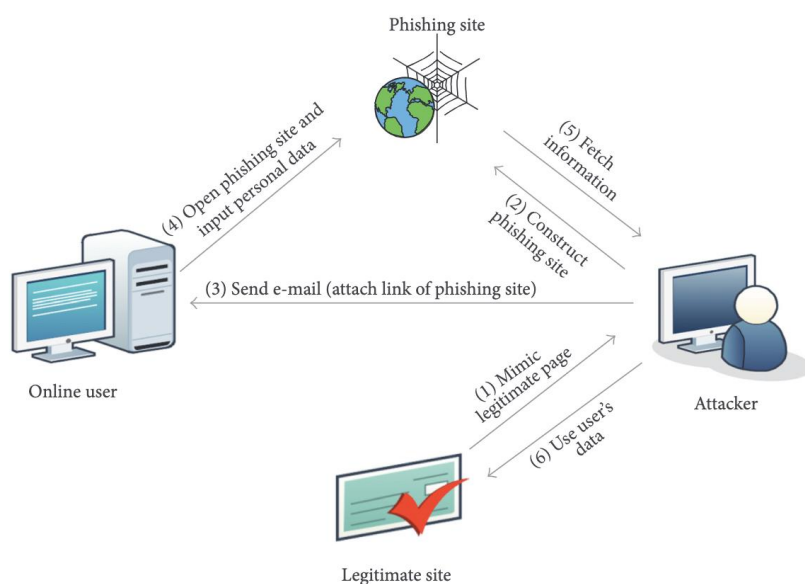


Рисунок 3 - Ілюстрація фішинг механізму [2]

Повна процедура того, як фішер запускає фішингову атаку, розділена на шість кроків.

- Крок 1. Фішери збирають і дублюють дані з добре відомого законного (цільового) веб-сайту.
- Крок 2. Фішери розробляють власний фішинговий веб-сайт відповідно до даних, зібраних на кроці 1.
- Крок 3. Фішери надсилають цей шкідливий веб-сайт кільком жертвам електронною поштою, зазвичай із посиланням у тексті електронного листа.
- Крок 4. Жертви клацають шкідливе посилання в електронному листі, переходять на фішинговий веб-сайт, вводять і надсилають свою особисту інформацію.
- Крок 5: дані жертв надсилаються фішеру з фішингового веб-сайту.

- Крок 6: Зрештою, фішери отримують доступ до цільового веб-сайту, використовуючи ідентифікаційну інформацію жертви.

2. Методи, які мають вплив на успішність фішингові атаки. Фішинг вивчається з початку 2000-х років, але проблема не була повністю вирішена, оскільки фішери постійно вдосконалюють свій підхід до атак [3]. Фішинговий веб-сайт відносно легко ідентифікувати шляхом спостереження за URL-шляхом, і уважний і досвідчений користувач тепер знає про ці підозрілі веб-сайти [4].

Фішер вводить в оману жертву фішингової атаки, видаючи себе за візуально подібний веб-сайт, використовуючи таку тактику [2]:

1. Візуальний вигляд. Фішинговий веб-сайт має високу візуальну схожість із законним веб-сайтом. Це пов'язано з тим, що фішери дублюють HTML-код законного веб-сайту, щоб створити свій фішинговий веб-сайт.

2. Адресний рядок. Фішери приховують URL-адресу за допомогою сценарію або зображення, у результаті чого жертва вважає, що вона перебуває на правильному веб-сайті.

3. Вбудовані об'єкти. Фішери використовують вбудовані об'єкти, такі як зображення, сценарії тощо, щоб приховати свій шкідливий текстовий вміст і HTML-код.

4. Подібність Favicon. Favicon – це піктограма зображення, пов'язана з певним веб-сайтом. Фішери можуть дублювати фавікон цільового веб-сайту, у результаті чого жертва буде вірити, що вони перебувають на правильному веб-сайті.

Крім того, особа, яка не обізнана з кібер безпекою, є основною причиною того, що вона потрапляє під фішингову атаку. Відповідно до репорту KnowBe4 Phishing By Industry Benchmarking Report 2023 Edition [5]:

- Користувачам бракує детальних знань про (справжню) URL-адресу.
- Користувачі не знають, якій веб-сторінці можна довіряти.
- Користувачі не перевіряють повну URL-адресу через переспрямування сторінки або приховані URL-адреси.
- Користувачі не мають багато часу, щоб перевірити URL-адреси, або вони випадково заходять на веб-сторінку
- Користувачі не можуть відрізнити фішингові веб-сторінки від законних.

3. Сучасні загрози фішингу. Наразі фішингові атаки є більш складними, оскільки фішери розвивають свій підхід до атак за допомогою різних методів. Наприклад, загальну фішингову атаку можна відносно легко ідентифікувати шляхом спостереження за URL-шляхом, і багато користувачів зараз знають про такий вид атаки через підозрілі URL-адреси або очевидну інформацію попередження з браузерів. Однак такі шкідливі URL-адреси можна приховати за допомогою різноманітних передових технологій, таких як використання скороченої URL-адреси чи QR-коду.

Крім того, запобігання фішинговим атакам на мобільних платформах є більш складним, ніж очікувалося. Поряд із тими ж проблемами, що й настільні комп'ютери, мобільні пристрої все ще стикаються з додатковими проблемами. Згідно з нашим підсумком, більшість фішингових посилань надходять із фішингових електронних листів, а мобільні платформи не підтримують безпечну ідентифікацію. Мобільний користувач не може знати, чи URL-адреса, до якої він отримав доступ, є безпечною, і, зокрема, чи користувачеві бракує достатньої обізнаності щодо безпеки. Наприклад, Google Chrome забезпечує набагато кращі механізми захисту від фішингових атак, ніж інші веб-браузери [6]. Він друкує сторінку попередження, яка показує, що URL-адреса, до якої відкривається доступ, містить потенційний ризик для користувачів у їхніх веб-переглядачах, якщо ця URL-адреса є зловмисною. Також Google анонсував про додавання нового захисту стандартної функції безпечного перегляду Google Chrome, увімкнувши захист від фішингу в реальному часі для всіх користувачів [7]. Хоча функція розширеного безпечного перегляду залишається незмінною та пропонує найкращий захист у Chrome, Google тепер додає захист у реальному часі до стандартної функції безпечного перегляду для підвищення безпеки.

Розробник браузера каже, що робить це, оскільки локальний список безпечного перегляду оновлюється лише кожні 30–60 хвилин, але 60% усіх фішингових доменів залишаються активними лише 10 хвилин. Це створює значний часовий проміжок, який залишає людей незахищеними від нових шкідливих URL-адрес. «Щоб заблокувати ці небезпечні сайти в момент їх запуску, ми оновлюємо безпечний перегляд, щоб тепер він перевіряв сайти на відомі шкідливі сайти Google у режимі реального часу», — каже Google.

«Завдяки скороченню часу між ідентифікацією та запобіганням загрозам ми очікуємо на 25% покращеного захисту від шкідливих програм і фішингових загроз».

Google повідомила, що функція Enhanced Safe Browsing за вмиканням зв'язується безпосередньо з протоколом Safe Browsing і надсилає додаткові дані. Хоча конфіденційність трохи менша, вона забезпечує найкращий захист, оскільки може виявляти шкідливі URL-адреси ще до того, як Google їх побачить.

Оскільки стандартна функція безпечного перегляду є опцією за замовчуванням, менеджер із продуктів Google Chrome Джасіка Бава повідомила BleepingComputer, що вони запроваджують захист у режимі реального часу з більшою мірою збереження конфіденційності через Fastly Oblivious HTTP Relays.

Протокол Oblivious передає частково хешовані URL-адреси користувачів системі безпечного перегляду Google, не розкриваючи особисту інформацію користувачів, таку як IP-адреси та заголовки запитів.

Однак ця стандартна функція безпечного перегляду в режимі реального часу, що забезпечує конфіденційність, має недолік. Оскільки він не надсилає стільки метаданих системі, він не зможе евристично визначити, чи є URL-адреса зловмисною, якщо її попередньо не помітить Google.

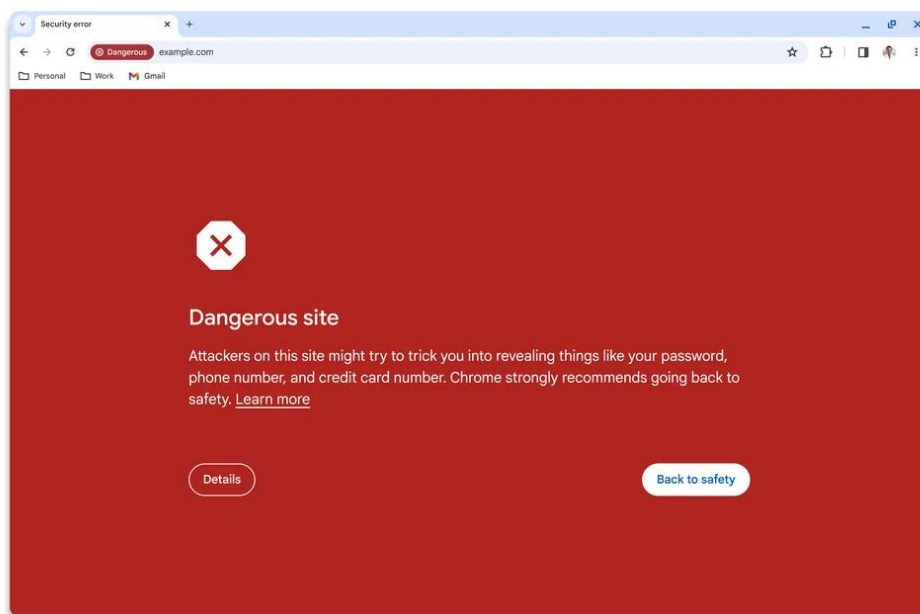


Рисунок 4 - Приклад роботи механізму захисту в Google Chrome [7]

4. Дослідження методів пом'якшення та захисту від фішингу. Виявлення фішингових атак є складною проблемою [5], оскільки ця атака використовує слабкі сторони людських характеристик, а не недоліки мережі. Для покращення захисту від фішингових атак можна виділити два напрями це або підвищення

ефективності технології виявлення фішингу, або розвиток освіти у користувачів. Загалом технологія виявлення фішингу зосереджена на ідентифікації зловмисних веб-сайтів, і існує два основні підходи, прийняті для пом'якшення фішингових атак: підхід на основі перевірки URL (за списком) та підхід на основі вмісту повідомлення. Схема на основі списків передбачає два типи списків: білий і чорний список. Ця схема є статичним підходом, у якому цільова URL-адреса порівнюється зі списками фішингу перед доступом до URL-адреси. У підході білого списку законні домени зберігаються в списку, і будь-який відповідний результат показує, що цільовий веб-сайт є законним веб-сайтом, і тому користувач може безпечно отримати доступ до цього веб-сайту. У підході до чорного списку шкідливі домени збираються в список, і будь-який відповідний результат показує, що цільовий веб-сайт, ймовірно, є фішинговим веб-сайтом, і користувачу буде надіслано сповіщення з попередженням у його браузері. Схема, заснована на вмісті, виявляє фішингові атаки шляхом вилучення певних функцій із цільової URL-адреси. Прикладом одного із методів є OpenPhish [8], що аналізує десятки мільйонів URL-адрес, щоб виявити фішинговий вміст. Цей звіт розбиває зміни в цільових брендах, галузях та фішинговій інфраструктурі. Дані нижче генеруються за допомогою їх бази даних фішингу.

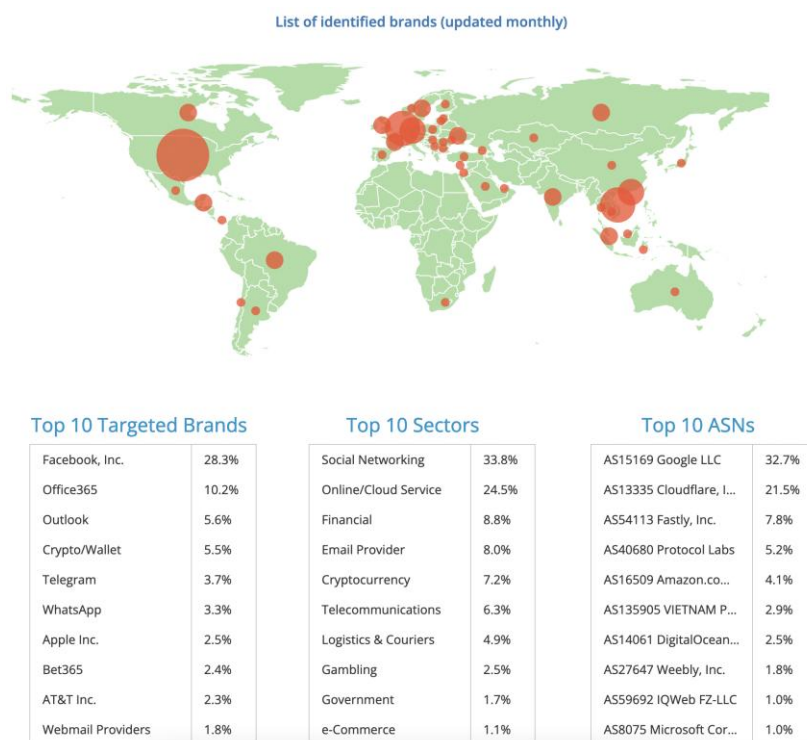


Рисунок 5 - Статистика брендів для impersonation [8]

5. Метод white black лістингу як захист від фішингу. В даний час різні дослідження базуються на цих двох підходах. Прихильники методи white black списків вважають, що перспективним рішенням є покращення продуктивності виявлення шляхом об'єднання чорних списків, перевірених вручну, з обчислювальними методами не лише для підвищення точності виявлення, але й для зменшення витрат часу на перевірку атак. Це рішення принесе користь сайтам, які підтримують служби фішингових чорних списків, наприклад PhishTank [9]. Комбінована техніка розглядалася у дослідженні «Smartening the Crowds: Computational Techniques for Improving Human Verification to Fight Phishing Scams» і складалась з чотирьох основних підходів; роблячи проблему непомітною для кінцевих користувачів, покращуючи дизайн інтерфейсів користувача, покращуючи навчання кінцевих користувачів [10]. Вони розробили систему, подібну до краудсорсингу, для виявлення фішингових веб-сайтів. Система вилучала неперевірені URL-адреси з PhishTank і фільтрувала URL-адреси, які не були в білому списку.

Вони розгорнули відповідні алгоритми (DBSCAN і shingling) для кластеризації схожих сторінок, щоб підвищити ефективність виявлення. Ці кластери були надіслані в Amazon Mechanical Turk для перевірки. Зрештою URL-адресу було визначено як фішингову або звичайною відповідно до ваги голосів від кожного учасника. Їхня система була швидшою за інші існуючі чорні списки. Крім того, їх рішення можна легко прийняти будь-яким існуючим чорним списком, перевіреним вручну, таким як ті, що пропонуються Google, Microsoft і PhishTank.

В дослідженні вони виявили, що велика кількість поточних фішингових веб-сайтів створюється за допомогою схожих інструментів, що призводить до створення схожих фішингових результатів через високу схожість вмісту. Таким чином, вони запропонували розширений метод ієрархічного чорного списку для виявлення фішингових атак шляхом використання існуючих зворотних списків. у їхньому підході для аналізу вмісту фішингових веб-сторінок було застосовано техніку n-grams на перевірених вручну чорних списках URL-адрес, а також одну майже дубльовану фішингову веб-сторінку було ідентифіковано ймовірнісним способом за допомогою методу shingling. Крім того, щоб зменшити кількість помилкових спрацьовувань, вони використали модуль фільтра, який додатково визначив легітимність потенційного фішингового веб-сайту за допомогою методів пошуку інформації в пошукових системах.

Підхід на основі списку не індексує лише цільову URL-адресу. Деякі дослідники [11] вважають, що більшість сучасних веб-сайтів складаються з кількох ресурсів, таких як CSS, JS та зображення. Однак ці законні (фірмові) веб-сайти зазвичай отримують вміст ресурсів з іншого домену через обмеження браузера щодо максимальної кількості одночасних підключень до того самого домену. Наприклад, законний веб-сайт PayPal отримує ресурси CSS, JS і файли зображень із paypalobjects.com. Таким чином, аналіз зв'язку запиту ресурсів є ефективним методом запобігання появі фішингових веб-сайтів. Geng та ін. запропонували новий підхід, який базується на зв'язках запитів ресурсів бренду майнінгу для виявлення фішингових атак. Їхній підхід не тільки ефективний і простий у застосуванні, але й є ефективним доповненням до існуючих методів. У їхньому рішенні цільова URL-адреса та всі запитувані домени перевірятимуться з чорного та білого списків.

6. Метод аналізу вмісту як захист від фішингу. З досліджень «Machine learning based phishing detection from URLs» [12] бачимо що було розроблено систему виявлення фішингу на основі навчання. Згідно з їхнім підходом, навчальні дані включають багато функцій, які стосуються як фішингу, так і законних класів веб-сайтів. Було вибрано сорок різних функцій на основі NLP, таких як підрахунок необроблених слів, перевірка домена, найбільша довжина слова, найкоротша довжина слова, TLD тощо. Крім того, сім різних алгоритмів були проведені в процесі навчання з метою вибору оптимального алгоритму, який має найвищу точність. Згідно з результатами цього підходу, алгоритм Random Forest показав найкращу продуктивність із найвищим рівнем точності виявлення серед цих семи алгоритмів для виявлення фішингових веб-сайтів. Niakanlahiji, Chi та Al-Shaer [13] запропонували масштабований багатофункціональний підхід машинного навчання для виявлення фішингових веб-сайтів під назвою «PhishMon». На відміну від інших підходів до виявлення фішингу на основі машинного навчання, вони вибрали п'ятнадцять нових функцій, які можна ефективно обчислити, використовуючи функції, які не вимагають взаємодії зі сторонніми службами, такими як пошукові системи та сервери WHOIS, таким чином зменшуючи час прийняття рішень. Функції були витягнуті з чотирьох аспектів: response HTTP, сертифікати SSL, документ HTML і файл JavaScript відповідно. Їхнє рішення забезпечує високу точність виявлення фішингових веб-сайтів, оскільки вибрані функції фіксують різні характеристики легітимних веб-додатків, а також їхні базові веб-інфраструктури.

7. MFA як додатковий рівень захисту. Поряд із вищезазначеними дослідженнями деякі дослідники вважали, що двофакторну автентифікацію можна застосувати для пом'якшення фішингових атак. Двофакторна автентифікація (також називається 2FA, MFA) — це механізм безпеки, який реалізує два вектори для автентифікації безпеки, і він вважається більш безпечним, ніж традиційна система автентифікації. Три загальновизнаних фактора автентифікації: те що ви знаєте або something that you know (наприклад, паролі), що у вас є або something that you have (наприклад, токени) і що ви є або what you are (біометрія) [14]. За останні кілька років 2FA реалізовано на більшості популярних веб-сайтів, таких як Google, Microsoft. 2FA допомагає користувачам захистити облікові записи; додатковий запит автентифікації надсилається користувачеві на іншому векторі, щоб перевірити та підтвердити додаткову ідентифікацію. Для пом'якшення фішингових атак 2FA може захистити обліковий запис користувача, навіть якщо фішери збирають ідентифікаційні дані користувача, фішер все одно не може отримати доступ до цільового веб-сайту, оскільки ідентифікацію користувача потрібно перевіряти за допомогою двох векторів. По суті, інші особи не можуть увійти в обліковий запис користувача без згоди користувача.

Однак лише 2FA не запобіжить успіху всіх фішингових атак. Існують різні підходи, які можуть обійти механізм 2FA. Наприклад, під час фішингової атаки, якщо жертва переходить на фішингову сторінку та вводить свої облікові дані, фішер може використовувати ці дані в режимі реального часу для входу на законний сайт [15]. Ілюстрацію цього сценарію показано на рисунку 6 нижче. Запит на код 2FA надсилається жертві (на кроці 4), потім ця жертва вводить код на фішинговому веб-сайті (на кроці 5). Згодом цей фішер використовує цей код для входу на справжній сайт як ідентифікований користувач. Зрештою, цей фішер може надіслати будь-що назад жертві, переконавши її, що веб-сайт, на який він отримав доступ, є законним, щоб зменшити її обізнаність про безпеку.

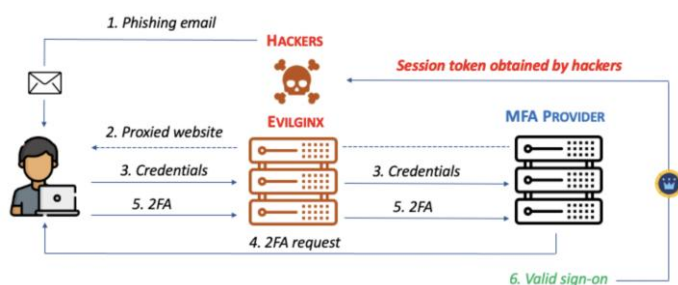


Рисунок 5 - Man-in-the-middle attack with evilginx [16]

Висновки. Пом'якшення фішингових атак є важливою темою дослідження, яку варто вивчити. Хоча було проведено багато досліджень, ця загроза все ще існує в реальному світі, поширеність якої постійно зростає. Згідно з результатами досліджень виявлення фішингових атак є складною проблемою. Для пом'якшення фішингових атак використовуються дві основні стратегії; або підвищення продуктивності технології виявлення фішингу, або покращення обізнаності людей про ці атаки. Розвиток людської досвідченості є основним способом подолання фішингових атак, оскільки фішингові атаки використовують слабкі сторони людських якостей, а не недоліки мережі. Крім того, люди завжди є найслабшою ланкою в атаках соціальної інженерії. Належне навчання є важливим, щоб гарантувати, що користувачі розуміють, як розпізнавати фішингові атаки, а якість навчання може вплинути на успіх запобігання фішингу. Таким чином, експеримент з тренінгами та навчанням вимагає від користувачів безпосередньої участі, оскільки якість навчання залежить від порівняння показників успішної ідентифікації фішингу в учасників до та після відповідного тренінгу. Однак для покращення продуктивності технології виявлення фішингу дослідники зосереджуються на двох аспектах пом'якшення фішингових атак; один базується на виявленні фішингових електронних листів, оскільки електронна пошта є найбільш вразливим засобом для запуску фішингової атаки. Таким чином, фішингові атаки будуть заблоковані з джерела, якщо буде виявлено фішинговий електронний лист. Інший спосіб зосереджений на виявленні фішингових веб-сайтів, оскільки більшість фішингових атак використовують фішинговий веб-сайт для незаконного збору даних жертви. Порівняно з виявленням фішингових веб-сайтів, виявлення фішингової електронної пошти може потребувати участі користувачів, щоб отримати кращий результат виявлення. Тому що успіх фішингової електронної пошти залежить від її контексту. Зокрема, коли передумова фішингового електронного листа узгоджується з робочим контекстом користувача (або поточною ситуацією). Таким чином, у цьому випадку досвід користувача та поточна ситуація є релевантними характеристиками, і ці змінні матимуть високу вагу, якщо дослідники запровадять машинне навчання для навчання відповідних даних. Крім того, для оцінювання потрібні учасники з різним досвідом для охоплення різноманітних ситуацій і тестів. Більшість антифішингових рішень впроваджено для пом'якшення загальних фішингових атак, але вони ігнорують деякі конкретні ситуації, наприклад розширені фішингові атаки. Щоб запобігти

розширеним фішинговим атакам, фішингові веб-сайти важко виявити, якщо жертва піддається атаці з використання викрадених DNS даних, оскільки вміст URL-адреси та вміст веб-сайту є такими самими, як і законні веб-сайти. Більшість підходів, заснованих на вмісті, можуть не спрацювати, оскільки вміст URL-адреси, до якої здійснюється доступ, є важливим фактор у прийнятті рішення. Однак дійсний і надійний сертифікат SSL неможливо підробити. Таким чином, ми розглядаємо визначення узгодженості сертифіката SSL між веб-сайтом, до якого ви отримали доступ, і законним веб-сайтом, щоб ідентифікувати фішинговий веб-сайт, який перебуває під атакою викрадення DNS. Крім того, під час атаки викрадення Інтернет-провайдера (ISP hijacking) фішери можуть не фішингувати жертву безпосередньо з законного веб-сайту, але вони можуть заманити жертв, скеровуючи їх на фішинговий веб-сайт, вставляючи переконливий контекст на законному веб-сайті. Щоб запобігти атакам захоплення субдоменів, фішинговий веб-сайт важко виявити, якщо фішери розмістили цей веб-сайт у субдомені, взятому з законного веб-сайту. Незалежно від веб-вмісту, URL-адреси та інформації сертифіката SSL, усі вони будуть такими самими, як і законний веб-сайт. Більше того, підхід до перебору субдоменів потребує вдосконалення, оскільки більшість поточних інструментів базується на грубому переборі, існуючі словники можуть не охоплювати всі випадки субдоменів, оскільки деякі субдомени можуть бути безглуздими. Таким чином, ми бачимо що ще є місце для покращення та вдосконалення захисту від фішингових атак.

ЛІТЕРАТУРА / REFERENCES

1. Verizon, "2023 Data Breach Investigations Report". 2023. [Online]. Access: <https://www.verizon.com/business/en-gb/resources/reports/dbir/>
2. A. K. Jain and B. B. Gupta, "Phishing detection: Analysis of visual similarity based approaches". In Journal of Security and Communication Networks, vol. 2017, Article ID. 5421046, pp. 1-20, Jan 2017. DOI: 10.1155/2017/5421046.
3. Zainab Alkhalil, Chaminda Hewage, Liqaa Nawaf and Imtiaz Khan, "Phishing Attacks: A Recent Comprehensive Study and a New Anatomy" 2021. DOI: 10.3389/fcomp.2021.563060
4. F. Castaño et al., "Evaluation of state-of-art phishing detection strategies based on machine learning". 2021. DOI: 10.18239/jornadas_2021.34.06.
5. Report Phishing by industry benchmarking report 2023. Access: <https://info.knowbe4.com/en-us/phishing-by-industry-benchmarking-report>

6. HKCERT, Browser's Anti-phishing feature: What is it and how it helps to block phishing attack? Access: <https://www.hkcert.org/blog/browser-s-anti-phishing-feature-what-is-it-and-how-it-helps-to-block-phishing-attack>
7. Google is enabling Chrome real-time phishing protection for everyone. Доступ: <https://www.cnet.com/news/privacy/google-chrome-can-now-warn-you-in-real-time-if-youre-getting-phished/>
8. Phishing feeds. Access: <https://openphish.com/>
9. PhishTank. Access: <https://phishtank.org/>
10. G. Liu, G. Xiang, B. A. Pendleton, J. I. Hong, and W. Liu, "Smartening the crowds: Computational techniques for improving human verification to fight phishing scams". 2011. DOI: 10.1145/2078827.2078838.
11. G. G. Geng, Z. W. Yan, Y. Zeng, and X. B. Jin, "RRPhish: Anti-phishing via mining brand resources request". 2018. DOI: 10.1109/ICCE.2018.8326085.
12. O. K. Sahingoz, E. Buber, O. Demir, and B. Diri, "Machine learning based phishing detection from URLs." 2019. DOI: 10.1016/j.eswa.2018.09.029
13. A. Niakanlahiji, B. T. Chu, and E. Al-haer, "PhishMon: A machine learning framework for detecting phishing webpages". 2018. DOI: 10.1109/ISI.2018.8587410.
14. M. Papathanasaki, L. Maglaras, N. Ayres, "Modern Authentication Methods: A Comprehensive Survey" 2022 DOI:10.5772/acrt.08
15. Evilginx - Bypassing MFA, phishing is back on the menu. Access: <https://bleekseeks.com/blog/evilginx-bypassing-mfa-phishing-is-back-on-the-menu>
16. How hackers beat MFA at-scale. Access: <https://www.mantra.ms/blog/beating-mfa>

Received 11.05.2023.

Accepted 16.05.2023.

Phishing like the first step to gaining access

Phishing as a term that means the technique of sending phishing messages will be researched based on findings in public access and using the listed links. The process of a phishing attack will be analyzed, and then we will pay attention to the technical vectors of how users become victims of the attack. Finally, existing research on phishing attacks and related prevention approaches will be reviewed.

Mitigating phishing attacks is an important research topic worth exploring. Although a lot of research has been done, this threat still exists in the real world, and its prevalence is constantly increasing. According to research results, detecting phishing attacks is a difficult problem. There are two main strategies used to mitigate phishing attacks; or improving the performance of phishing detection technology or improving peo-

ple's awareness of these attacks. Developing human expertise is a key way to defeat phishing attacks, as phishing attacks exploit human weaknesses rather than network weaknesses. Also, humans are always the weakest link in social engineering attacks.

Compared to phishing website detection, phishing email detection may require user involvement to achieve better detection results. Because the success of a phishing email depends on its context. Specifically, when the premise of the phishing email is consistent with the user's work context (or current situation).

Most anti-phishing solutions are implemented to mitigate general phishing attacks, but they ignore some specific situations, such as advanced phishing attacks. To prevent advanced phishing attacks, phishing websites are difficult to detect if a victim is attacked using stolen DNS data because the URL content and website content are the same as legitimate websites. Most content-based approaches may not work because the content of the accessed URL is an important factor in the decision.

To prevent subdomain hijacking attacks, it is difficult to detect a phishing website if the phishers have hosted the website on a subdomain taken from a legitimate website. Regardless of the web content, URL, and SSL certificate information, they will all be the same as the legitimate website. Moreover, the approach to enumeration of subdomains needs improvement, as most current tools are based on rough enumeration, existing dictionaries may not cover all instances of subdomains, as some subdomains may be meaningless.

Key words: phishing, cyber security, multifactor authentication, social engineering.

Гуда Антон Ігорович – д.т.н, проф., професор кафедри ІТС, ІПБТ УДУНТ.

Кліщ Сергій Михайлович – асистент кафедри ІТС, ІПБТ УДУНТ.

Guda Anton – doctor of engineering's sciences, professor, IIBT.

Klishch Sergey – post-graduate student, assistant, IIBT.

К.Ю. Островська, І.В. Стовпченко, Д.С. Печений

ДОСЛІДЖЕННЯ МЕТОДІВ НА ОСНОВІ НЕЙРОННИХ МЕРЕЖ ДЛЯ АНАЛІЗУ ТОНАЛЬНОСТІ КОРПУСУ ТЕКСТІВ

Анотація Об'єктом дослідження є методи з урахуванням нейронних мереж для аналізу тональності корпусу текстів. Для досягнення поставленої в роботі мети необхідно вирішити такі завдання: дивити теоретичний матеріал для навчання глибоких нейронних мереж та їх особливості стосовно обробки природної мови; вивчити документацію бібліотеки Tensorflow; розробити моделі згорткової та рекурентної нейронних мереж; розробити реалізацію лінійних та нелінійних методів класифікації на моделях мішка слів та Word2Vec; порівняти точність та інші показники якості реалізованих нейромережових моделей із класичними методами. Для візуалізації навчання використовується Tensorboard. У роботі показано перевагу класифікаторів на основі глибоких нейронних мереж над класичними методами класифікації, навіть якщо для векторних уявлень слів використовується модель Word2Vec. Найвищу точність для даного корпусу текстів має модель рекурентної нейронної мережі з LSTM блоками.

Ключові слова: штучні нейронні мережі, Глибокі нейронні. мережі, навчання з учителем, глибоке навчання, рекурентна нейронна мережа, LSTM, згорткова нейронна мережа, аналіз тональності тексту, мішок слів, Word2vec.

Вступ. Щодня в Інтернеті з'являється величезна кількість контенту: користувачі висловлюють свою думку про фільми та події, залишають відгуки про різні продукти та послуги. Для вирішення завдань, пов'язаних із аналізом емоційного забарвлення тексту, використовуються методи аналізу тональності тексту.

Завдання автоматичного аналізу тональності тексту є досить популярним [1]. Найчастіше користувачі при виборі чогось (наприклад, до якого університету вступити) керуються думками інших людей. Тому набір опрацьованих думок становить значний інтерес для соціологів, маркетологів та власників бізнесу. Також система для аналізу тональності тексту може стати в нагоді на сайтах, де люди залишають відгуки: наприклад, користувач може позитивно оцінити фільм, випадково поставивши негативну оцінку - система автоматичного аналізу тональності тексту виправить помилку. Тому аналіз тональності текстів є важливим та актуальним завданням.

© Островська К.Ю., Стовпченко І.В., Печений Д.С., 2023

Автоматичний аналіз тональності тексту зазвичай застосовується на корпусах текстів, які містять рецензії чи відгуки. Також його можна застосувати для даних з блогів та соціальних мереж, щоб отримати громадську думку з того чи іншого питання або товару.

Визначення емоційної оцінки авторів рецензій перестав бути очевидним завданням. Стандартні методи класифікації тексту ізолюють слова за ознаками та застосовують різні методи вибору ознак, щоб знайти найбільш “важливе слово” у тексті: наприклад, речення "Мені подобається бігати" і "Мені не подобається бігати" вважатимуться однаковими. У ситуації, коли рецензія представлена лише позитивно забарвленим набором слів, а насправді негативною, стандартні методи класифікації тексту перестають бути ефективними. Методи машинного навчання дозволяють алгоритмам “розуміти” структуру речення та його семантичну структуру. Пропозиція буде представлена у вигляді вектора, в якому збережена структура речення та те, як слова пов'язані один з одним.

Для проведення чисельних експериментів були використані рецензії сайту Rotten Tomatoes [2] — набір із 5331 позитивних та 5331 негативних рецензій.

Потрібно побудувати бінарний класифікатор, визначальною, позитивною чи негативною виявилася рецензія. Як методи розглядаються згортова нейронна мережа і рекурентна нейронна мережа з LSTM-блоками, наївний класифікатор байесовського, метод опорних векторів і логістична регресія. Також було використано два варіанти векторних моделей подання тексту: мішок слів (Bag of Words) та Word2Vec.

В результаті навчання були отримані моделі нейронних мереж, що дозволяють з високою точністю визначати тональність рецензій, а також порівняння ефективності використання реалізованих методів.

Аналіз останніх досліджень і публікацій. Завданням, поставленою у роботі, є реалізація різних алгоритмів глибинного навчання для аналізу тональності тексту та порівняння їх ефективності з класичними (лінійними та нелінійними) алгоритмами класифікації тексту.

У дослідженні [3] представлені результати ефективності застосування різних архітектур згорткових нейронних мереж при використанні попередньо навченої моделі векторного уявлення word2vec.

Результати застосування нейронних мереж з LSTM-блоками та їх модифікаціями були представлені в дослідженні [4], а також було показано, що ця архітектура є найбільш ефективною для побудови бінарного та багатокласового класифікаторів для аналізу тональності текстів.

Також у роботі [5] стверджується, що можна значно покращити ефективність класифікатора, використовуючи попередньо навчені моделі векторних уявлень слів (наприклад, Word2Vec).

У цій роботі для побудови бінарного класифікатора будуть використовуватися рецензії, що складаються з коротких речень.

В даний час завдання побудови класифікатора (бінарного або багатокласового) для визначення тональності тексту вирішується за допомогою нейромережових моделей, оскільки ефективність архітектур, використаних у згаданих роботах, значно вища, ніж у класичних лінійних алгоритмів.

Існує безліч бібліотек для реалізації алгоритмів машинного навчання. Існують два типи фреймворків: символічні та імперативні. У символічних фреймворках набагато більше можливостей використовувати пам'ять багаторазово, а оптимізація на основі граф залежностей здійснюється автоматично. Найпопулярнішими символічними (symbolic) фреймворками нині є TensorFlow та Theano.

На відміну від Theano, Tensorflow не орієнтований лише на навчання нейронних мереж, тому можна використовувати колекції графів та черги як складові для високорівневих компонентів.

Якщо необхідно навчати масштабні моделі і використовувати багато зовнішньої пам'яті, Theano буде дуже повільно працювати через необхідність компіляції коду C/CUDA в бінарний код.

TensorFlow має прозору модульну архітектуру з безліччю фронт-ендів. В архітектурі Theano розібратися досить непросто: весь код – це Python, де код C/CUDA упакований як рядок Python. У такому коді складно орієнтуватися, його непросто налагоджувати та проводити рефакторинг. Більше того, візуалізація графів у TensorFlow реалізована значно ефективніше, ніж у Theano [6].

Векторні уявлення слів для лінійних алгоритмів будуть представлені двома моделями: Word2Vec та мішок слів (bag of words), а для класифікаторів на основі нейронних мереж буде використано лише модель Word2Vec. За допомогою інструменту для побудови векторних моделей Gensim буде навчено модель Word2Vec. З використанням TensorFlow будуть реалізовані згортоква нейронна мережа та рекурентна нейронна мережа з LSTM-блоками.

Методи. Процес обробки даних складається з наступних кроків:

- видалити розмітку HTML;
- видалити всі символи, окрім літер та пробілів;

- з отриманого набору слів видалити стоп-слова.

Наступним завданням є перетворення кожної рецензії на векторне уявлення. Для оцінки ефективності кожного з методів буде використано дві моделі векторного подання слів: мішок слів та Word2Vec.

У роботі логістична регресія використовується для передбачення ймовірності виникнення певної події.

Наївний класифікатор Байеса - простий імовірнісний класифікатор, заснований на застосуванні Теорему Байеса зі суворими припущеннями про незалежність елементів вектора ознак. Перевагою наївного класифікатора Байеса є мала кількість даних для навчання, необхідних для оцінки параметрів, необхідних для класифікації.

Випадковий ліс (random forest) – алгоритм полягає у використанні комітету вирішальних дерев. Класифікація об'єктів проводиться шляхом голосування: кожне дерево комітету відносить об'єкт, що класифікується, до одного з класів, і перемагає клас, за який проголосувало найбільшу кількість дерев. Оптимальна кількість дерев підбирається таким чином, щоб мінімізувати помилку класифікатора на тестовій вибірці. Використовує усереднення для підвищення точності прогнозування та контролю надлишкового припасування. Розмір підвиборки завжди збігається з розміром оригінальної вибірки.

Метод опорних векторів (SVM) - Основна ідея методу - переведення вихідних векторів у простір більш високої розмірності та пошук роздільної гіперплощини з максимальним зазором у цьому просторі. Стандартна функція `scikitlearn (sklearn.svm.SVC)` не підходить, оскільки тимчасова складність даного алгоритму квадратична. Ефективність методу опорних векторів значно знижується, якщо кількість ознак дуже велика. Він має велику гнучкість у виборі штрафів та функцій втрат і має краще масштабуватись для великої кількості зразків. Оптимізація функції втрат SVM за допомогою градієнтного спуску:

$$L(\omega, D) = \frac{1}{2} \|\omega\|_2^2 + C \sum_{i=1}^N \max(0, 1 - y_i(\omega^T x_i + b)), \quad (1)$$

де

$$D = \{(x_i, y_i)\}_{i=1}^N, x_i \in R^d \text{ and } y_i \in \{-1, +1\}. \quad (2)$$

Оптимізація функції втрат - (regularization_loss + hinge_loss). Для випадку word2vec буде використано `SGDClassifier` зі `sklearn.linear_model` з помилкою l2.

Згортова нейронна мережа - Векторні подання даних здійснено за допомогою моделі Word2Vec (Skip-gram Model).

Як функція оптимізації використовується Adam (Adaptive Moment Estimation). Їого відмінними рисами є:

- оцінка першого моменту обчислюється як ковзне середнє;
- оскільки оцінки першого та другого моментів ініціалізуються нулями, використовується невелика корекція, щоб результуючі оцінки не були зміщені до нуля.

Рекурентна нейронна мережа з LSTM-блоками - векторні подання даних здійснено за допомогою моделі Word2Vec (Skip-gram Model). LSTM вважаються найкращою архітектурою для аналізу тональності тексту. Нейронні мережі, складені з LSTM-модулів, особливо добре обробляють заперечення (negation), якщо у осередку є projection unit (тобто більше пам'яті мережі).

Як досвідчений зразок були використані рецензії веб-сайту RottenTomatoes [1] — набір із 5331 позитивних та 5331 негативних рецензій. Для розробки використовувалися бібліотеки Pandas, Scikit-Learn та PyMorphy2, а також фреймворк TensorFlow як засіб аналізу.

Алгоритм аналізу даних у додатку.

1. Отримуємо дані з іншого джерела (бази даних, користувальницького інтерфейсу).

2. Видаляємо зайву інформацію з пропонованого тексту залишаючи тільки російські літери.

3. Проводимо морфологічний аналіз тексту, та лематизуємо текст.

4. Будуємо модель:

- схема n-грам: (1, 3) (уніграми + біграми + триграми);
- Метод векторизації: Word2Vec;
- Тип моделі: Рекурентна нейронна мережа з LSTM-блоками;

Параметри моделі: penalty – 12, alpha – 0.000001, loss – log.

5. Навчаємо нейронну мережу за отриманими даними.

Алгоритм навчання нейронної мережі. Для того, щоб нейронна мережа почала виконувати свої завдання її необхідно навчити, процес навчання відбувається за принципом, показаним на рисунку 1.

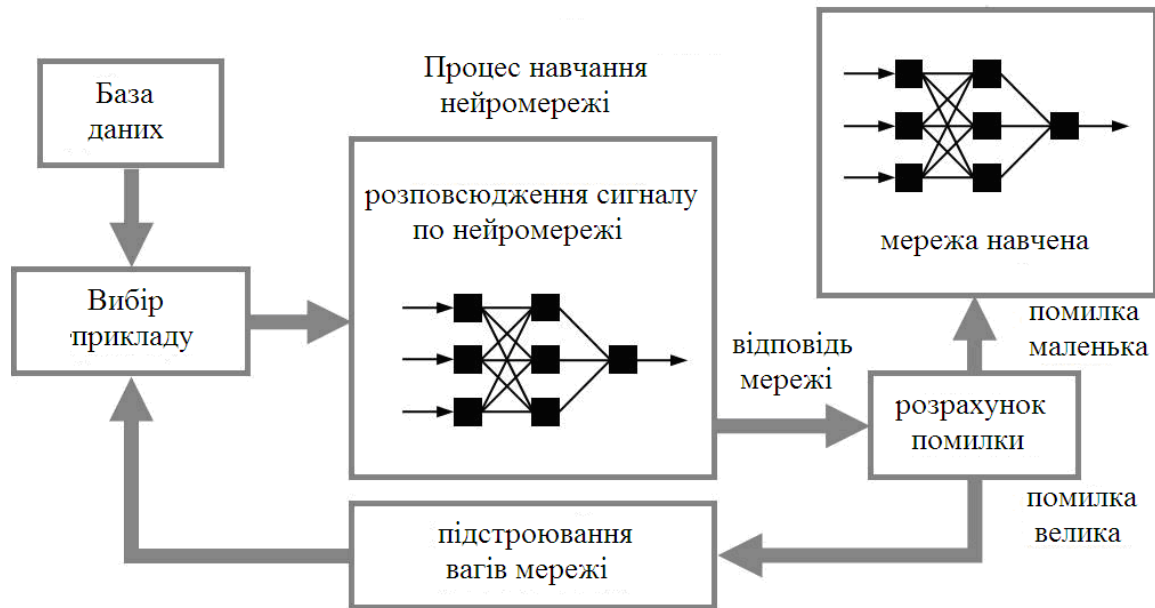


Рисунок 1 - Частка коректних прогнозів (ассигасу):
синій графік - навчання, червоний – перевірка

Показники якості нейронної мережі. Частка коректних прогнозів (ассигасу) - відсоток помилок, допустимих класифікатором.

Наступні показники будуть використані лише для класичних методів класифікації:

- міра точності (precision) - відношення tp до $(tp + fp)$, де tp - кількість істинних позитивних величин, а fp - кількість хибних позитивних величин. Тобто міра точності характеризує скільки одержаних від класифікатора позитивних рішень вважаються вірними;

- Міра повноти (recall) – відношення tp до $(tp + fn)$, де fn – кількість хибних негативних величин. Міра повноти встановлює вміння класифікатора дізнаватися так само як і найбільше позитивних рішень з прогнозованих;

- Міра F1 – середній гармонійний міри точності та міри повноти.

Визначає лімінальну властивість класифікатора;

- Носій міри (support) – кількість інформації будь-якого з класів.

Найбільш жорстке визначення: мінімальне закрите велике число, в якому сконцентровано ступінь.

Метод мішка слів (див. таблиці 1 - 4).

Таблиця 1

Логістична регресія

Клас	Мера точності	Мера повноти	Мера F1	Носій міри
0	0.75	0.76	0.75	1192
1	0.74	0.74	0.74	1139
total	0.74	0.74	0.74	2331
Достовірність 0.756				

Таблиця 2

Наївний байєсовський класифікатор

Клас	Мера точності	Мера повноти	Мера F1	Носій міри
0	0,73	0,73	0,72	1192
1	0,71	0,75	0,72	1139
total	0,72	0,71	0,71	2331
Достовірність 0.730				

Таблиця 3

Випадковий ліс

Клас	Мера точності	Мера повноти	Мера F1	Носій міри
0	0,71	0,74	0,72	1192
1	0,71	0,68	0,69	1139
total	0,71	0,71	0,71	2331
Достовірність 0.718				

Таблиця 4

Лінійний метод опорних векторів

Клас	Мера точності	Мера повноти	Мера F1	Носій міри
0	0,79	0,56	0,65	1192
1	0,64	0,85	0,73	1139
total	0,72	0,70	0,69	2331
Достовірність 0.689				

Таблиця 5

Логістична регресія

Клас	Мера точності	Мера повноти	Мера F1	Носій міри
0	0.77	0.77	0.77	1192
1	0.76	0.76	0.76	1139
total	0.86	0.86	0.86	2331
Достовірність 0.767				

Таблиця 6

Наївний байєсовський класифікатор

Клас	Мера точності	Мера повноти	Мера F1	Носій міри
0	0.73	0.72	0.72	1192
1	0.71	0.72	0.71	1139
total	0.72	0.72	0.72	2331
Достовірність 0.728				

Таблиця 7

Випадковий ліс

Клас	Мера точності	Мера повноти	Мера F1	Носій міри
0	0.73	0.74	0.75	1192
1	0.73	0.73	0.73	1139
total	0.73	0.73	0.74	2331
Достовірність 0.738				

Таблиця 8

Лінійний метод опорних векторів

Клас	Мера точності	Мера повноти	Мера F1	Носій міри
0	0.831	0.633	0.711	1092
1	0.691	0.864	0.771	1039
total	0.762	0.745	0.741	2131
Достовірність 0.743				

Згортова нейронна мережа. Достовірність 0.789. Графіки точності моделі та функції втрат представлені на рисунках 2 та 3 відповідно.

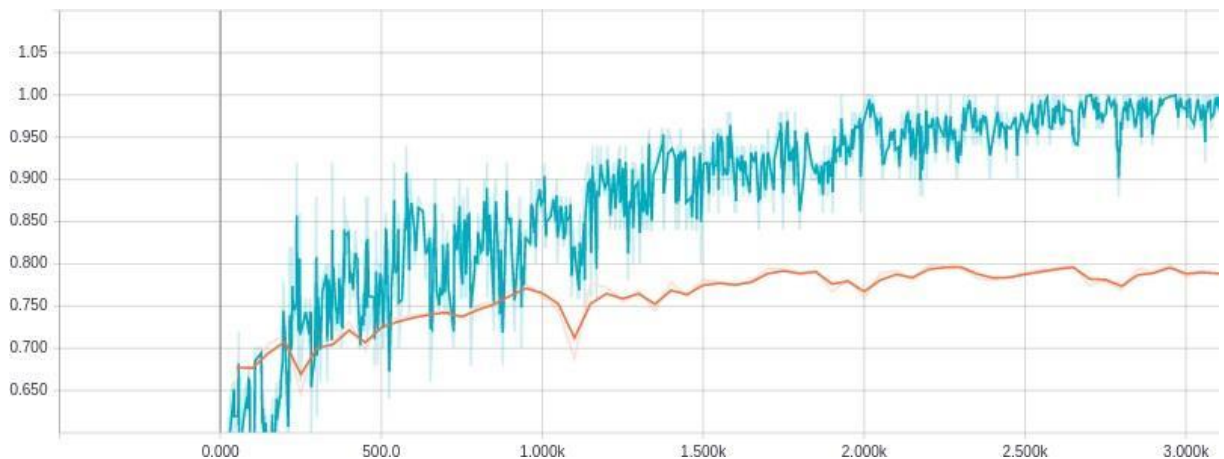


Рисунок 2 - Частка коректних прогнозів (ассурасу):
синій графік - навчання, червоний – перевірка

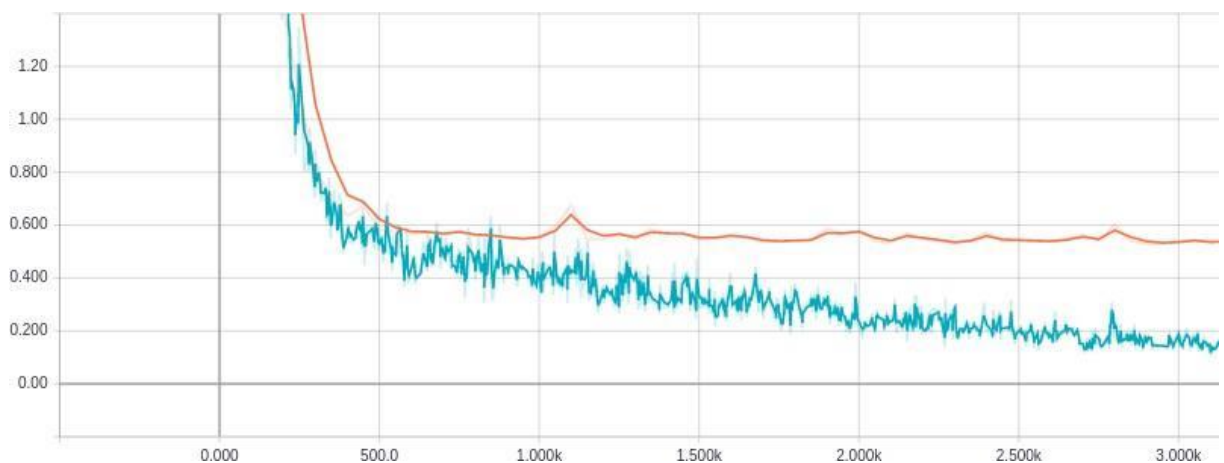


Рисунок 3 - Функція втрат: синій графік - навчання, червоний - перевірка

Рекурентна нейронна мережа з LSTM-блоками. Вирішальним у навчанні моделі виявився вибір функції мінімізації градієнтного спуску. Спочатку було використано алгоритм RMSProp (root mean square propogation), ідея якого полягає у масштабуванні градієнта.

Однак за інших рівних умов (розміру прихованого LSTM шару та темпі навчання) застосування алгоритму оптимізації Adam (який крім ідеї масштабування градієнта використовує ідею інерції) дозволило досягти максимальної точності щодо вже реалізованих методів цієї роботи - 83.1%.

Графіки точності моделі та функції втрат представлені на рисунках 4 та 5 відповідно.

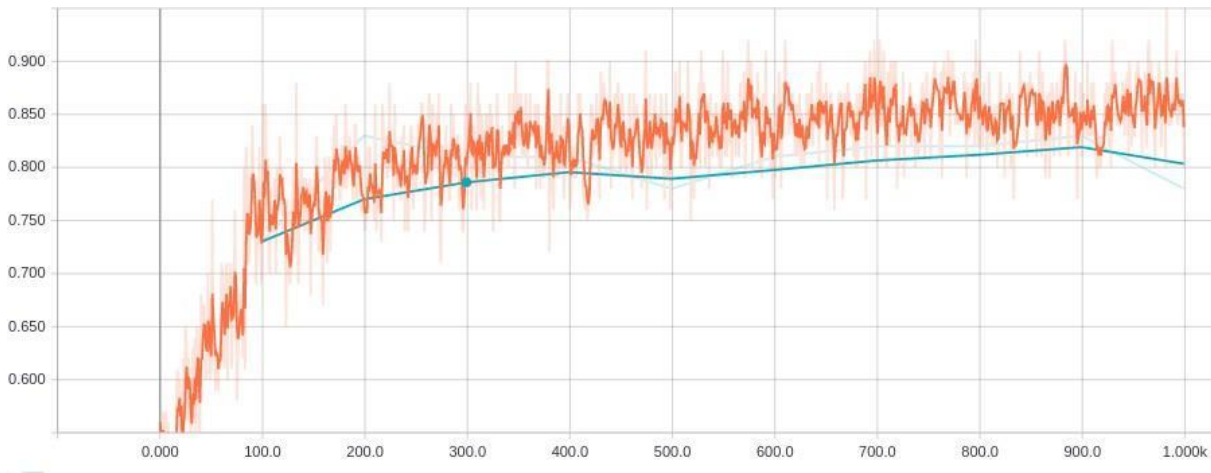


Рисунок 4 - Частка коректних прогнозів (ассурасы):
червоний графік-навчання, синій – перевірка

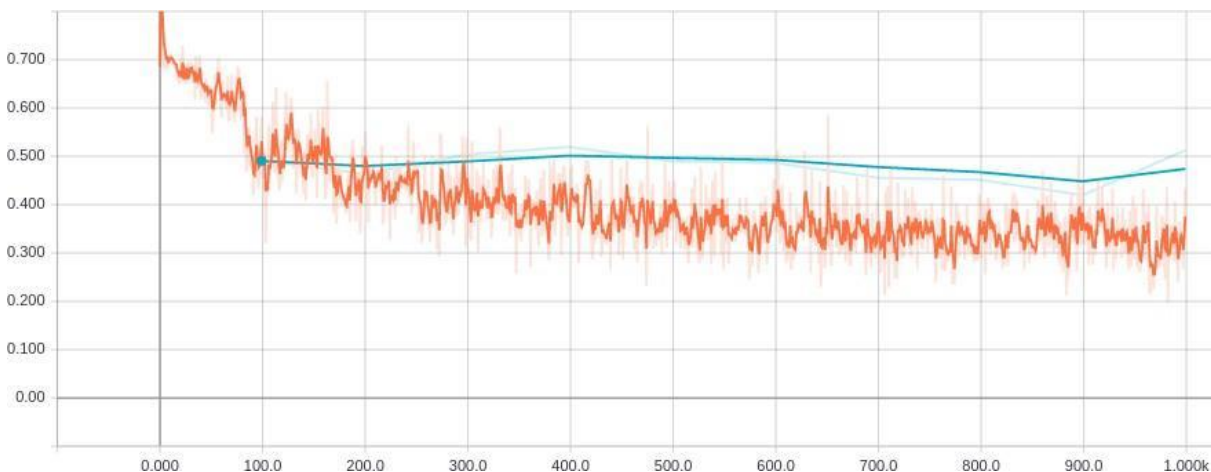


Рисунок 5 - Функція втрат: червоний графік - навчання, синій – перевірка

Результати. Для класичних алгоритмів класифікації використання мішка слів як модель векторного подання слів точність передбачень не перевищувала 75.4% (логістична регресія), мінімальна точність становила 69.9% (лінійний метод опорних векторів).

Завдяки використанню моделі Word2Vec вдалося покращити точність передбачень майже всім методів (крім наївного байесівського класифікатора): наприклад, точність лінійного методу опорних векторів стала 74.2%. Проте найефективнішим бінарним класифікатором виявилася логістична регресія: її точність із використанням моделі Word2Vec дорівнює 76,6%. За допомогою згорткових нейронних мереж з моделлю Word2Vec вдалося отримати точність 79.9%.

Найефективнішою архітектурою для аналізу тональності тексту виявилася рекурентна нейронна мережа з LSTM-блоками. Її максимальна точність становила 83.0%.

Отримані експериментальні дані показують вищу ефективність роботи глибоких нейронних мереж проти класичними алгоритмами для аналізу тональності тексту.

Результати дослідження свідчать, що використання глибоких нейронних мереж значно покращує точність аналізу тональності тексту.

Перевага рекурентної мережі на основі LSTM над згортковою нейронною мережею в галузі аналізу тональності вже була доведена в різних дослідженнях, проте важливо відзначити, що в даній роботі були реалізовані найпростіші архітектури глибоких нейронних мереж. Поліпшення параметрів моделі, використання більш розширеної моделі векторного уявлення слів Word2Vec, застосування attention-механізмів дозволить значно збільшити ефективність бінарного класифікатора для аналізу тональності на основі глибоких нейронних мереж.

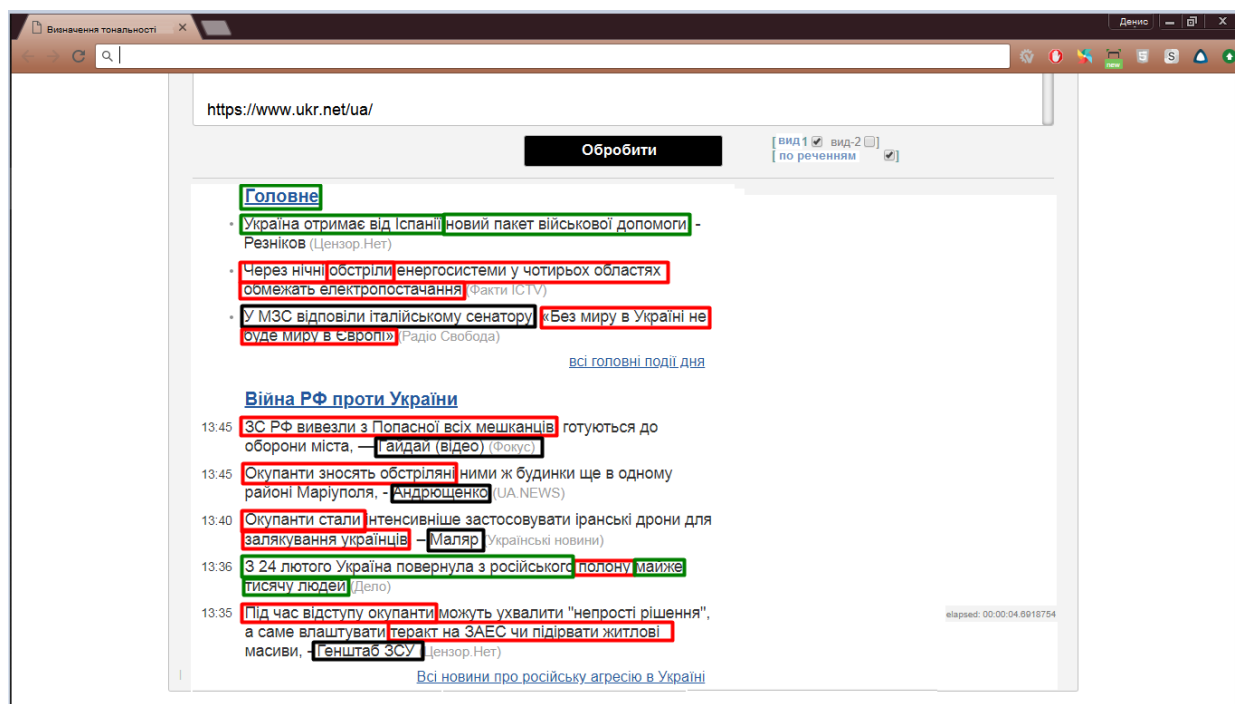


Рисунок 6 - Приклад роботи модуля

Перспективи подальших досліджень. Можливим напрямом для подальшої роботи може бути розпізнавання ключових слів, які роблять найбільший внесок у позитивний або негативний відгук. Введення модуля в експлуатацію.

Висновки. Відповідно до поставленої мети - розробка алгоритмів глибокого навчання для аналізу тональності тексту та порівняння їх ефективності з

іншими класифікаторами на основі алгоритмів машинного навчання – було реалізовано архітектуру згорткової нейронної мережі, рекурентної нейронної мережі з LSTM-блоками, а також проведено порівняння показників якості їх класифікації з іншими класифікаторами.

При використанні моделі мішка слів точність різних методів була значно вищою за випадкову (близько 70%), проте застосовуючи модель Word2Vec, вдалося значно покращити точність роботи алгоритмів (на кілька одиниць). Однак нейронні мережі показали найкращі результати.

Точність класифікатора на основі згорткової нейронної мережі виявилася 79.9%. Найвищу точність показав класифікатор на основі рекурентної мережі з LSTM-блоками – 83.3%.

ЛІТЕРАТУРА / REFERENCES

1. B. Pang and L. Lee. Opinion mining and sentiment analysis. Foundations and Trends in Information Retrieval archive, 2008.
2. Данные рецензий, используемых в работе, sentence polarity dataset v1.0. <http://www.cs.cornell.edu/people/pabo/movie-review-data/>
3. Y.Kim. Convolutional neural networks for sentence classification.arXiv:1408.5882 [cs.CL], 2014.
4. K.S. Tai et al. Improved semantic representations from tree-structured long short-term memory network. arXiv:1503.00075 [cs.CL], 2015.
5. Q. Le and T. Mikolov. Distributed representations of sentences and documents. arXiv:1405.4053 [cs.CL], 2014.
6. K.Tran et al. Evaluation of deep learning toolkits. <https://github.com/zer0n/deepframeworks/blob/master/README.md>

Received 19.05.2023.

Accepted 24.05.2023.

Research of methods based on neural networks for the analysis of the tonality of the corps of the texts

The object of the study is methods based on neural networks for analyzing the tonality of a corpus of texts. To achieve the goal set in the work, it is necessary to solve the following tasks: study the theoretical material for learning deep neural networks and their features in relation to natural language processing; study the documentation of the Tensorflow library; develop models of convolutional and recurrent neural networks; to develop the implementation of linear and non-linear classification methods on bag of words and Word2Vec models; to compare the accuracy and other quality indicators of implemented neural network models with classical methods. Tensorboard is used for

learning visualization. The work shows the superiority of classifiers based on deep neural networks over classical classification methods, even if the Word2Vec model is used for vector representations of words. The model of recurrent neural network with LSTM blocks has the highest accuracy for this corpus of texts.

Keywords: artificial neural networks, DEEP neural networks. networks, tutored learning, deep learning, recurrent neural network, LSTM, convolutional neural network, text tonality analysis, bag of words, Word2vec.

Островська Катерина Юріївна – к.т.н., доцент, доцент кафедри інформаційних технологій і систем, Український держаний університет науки та технологій Дніпро, Україна.

Стовпченко Іван Володимирович – ст. викладач кафедри інформаційних технологій і систем, Український держаний університет науки та технологій Дніпро, Україна.

Печений Денис Сергійович - магістрант кафедри інформаційних технологій і систем, Український держаний університет науки та технологій Дніпро, Україна.

Ostrovska Kateryna - Ph.D., Associate Professor, Associate Professor of the Department of Information Technologies and Systems, Ukrainian State University of Science and Technology of Dnipro, Ukraine.

Stovpchenko Ivan - senior teacher of the Department of Information Technologies and Systems, Ukrainian State University of Science and Technology, Dnipro, Ukraine.

Pechenyi Denys - master's student at the Department of Information Technologies and Systems, Ukrainian State University of Science and Technology, Dnipro, Ukraine.

ЗМІСТ

CONTENTS

Скалозуб В.В., Горячкін В.М., Мурашов О.В. Ре- ляційно-сепарабельні моделі про- цесів моніторингу при перемінних і нечітких інтервалах спостере- жень	3	Skalozub V.V., Horiachkin V.M., Murashov O.V. Re- lational-separable models of monitoring processes at variable and unclear observation intervals	3
Малієнко С.Є., Селівьорова Т.В. Огляд математичних моделей та інформаційних технологій бізнес аналізу великих web-даних	21	Maliienko S.E., Selivorstova T.V. Review of mathematical models and information technologies for business analysis of the big web data	21
Дмитрієва І.С., Бімалов Д.В. Роз- робка програмного модулю для ідентифікація емоційного стану користувача	29	Dmytriieva I.S., Bimalov D.V. De- velopment of a software module for the identification of the emo- tional state of the user	29
Ігнаткін В.У., Дудніков В.С., Лучишин Т.Р., Алексєєнко С.В., Юшкевич О.П., Карпова Т.П., Хохлова Т.С., Хомош Ю.С., Тіхо- нов В.А. Альтернатива методам середніх та найменших квадратів, які викори- стовуються при обробці результатів науково-технічних експериментів	35	Ignatkin V.U., Dudnikov V.S., Luchyshyn T.R., Alekseenko S.V., Yushkevich O.P., Karpova T.P., Khokhlova T.S., Khomosh Y. S., Tikhonov V.A. Alternative to mean and least squares methods used in processing the results of scientific and technical experiments	35
Лабуткіна Т.В., Ананко Р.В. Глобальне покриття навколоземного простору зонами використання пристроїв його спостереження: концепція і алгоритми	49	Tetyana V., Ananko R. Global near-earth space coverage by zones of the use of its observation devices: concept and algorithms	49

Атаманиук О.В., Легеза В.П. Метод побудови цифрового двійника віброзахисного процесу	71	Atamaniuk O.V., Legeza V.P. Method of creation a digital twin of a vibration protection process	71
Бубенко М.О., Герасимов В.В., Карпенко Н.В., Морозов О.С. Аналіз методів тестування веб-додатків	82	Bubenko M.O., Gerasimov V.V., Karpenko N.V., Morozov A.S. Analysis of web application testing methods	82
Тулуб В.О. Методи підвищення рівня ефективності роботи автоматизованих систем	91	Tulub V. Methods of increasing the level efficiency of automated systems	91
Могильний О.А. Автоматизовані моделі обробки візуальної інформації	100	Mohylnyi O. Automated models of visual information processing	100
Солдатенко Д.В., Гнатушенко Вік.В. Покращення ефективності глибокого навчання шляхом збільшення навчальних даних	111	Soldatenko D., Hnatushenko Vik. Improving deep learning perfor- mance by augmenting training da- ta	111
Ситник Р.С., Гнатушенко Вік. В. Управління потоками даних в сучасній промисловості за допомогою блокчейну	123	Sytnyk R., Hnatushenko Vik. Management of data flows in modern industry using blockchain	123
Ліп'яніна-Гончаренко Х.В. Інтелектуальний метод вибору локації для старту бізнесу в розумному місті	132	Lipianina-Honcharenko Kh. Intellectual method for business location selection in smart cities	132
Гуда А.І., Кліщ С.М. Фішинг як перший крок до отримання доступу	141	Guda A.I., Klishch S.M. Phishing like the first step to gaining access	141
Островська К.Ю., Стовпченко І.В., Печений Д.С. Дослідження методів на основі нейронних мереж для аналізу тональності корпусу текстів	155	Ostrovska K.Yu., Stovchenko I.V., Pechenyi D. Research of methods based on neural networks for the analysis of the tonality of the corps of the texts	155
ISSN 1562-9945 (Print) ISSN 2707-7977 (Online)			169

РЕФЕРАТИ

УДК 004.6: 007.5: 004.94

Скалозуб В.В., Горячкін В.М., Мурашов О.В. **Реляційно сепарабельні моделі процесів моніторингу при перемінних і нечітких інтервалах спостережень** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 4(147). – Дніпро, 2023. – С.3 – 20.

Стаття присвячена розвитку комбінованих моделей, методів і засобів, призначених для вирішення актуальних завдань моделювання та аналізу даних процесів моніторингу, які представлені часовими рядами і відрізняються перемінним або нечітким інтервалом спостережень (ЧРПНІ).

В результаті виконаних досліджень шляхом числового моделювання було встановлено, що запровадження комбінованих моделей процесів при ЧРПНІ являється раціональним та результативним. Приклади аналізу даних моніторингу процесів реабілітації хворих на діабет показали певні можливості забезпечення вимоги до точності результатів аналізу показників та їх короткострокового прогнозування.

Бібл. 14, іл. 5.

УДК 004.62

Малієнко С.Є., Селівьорстова Т.В. **Огляд математичних моделей та інформаційних технологій бізнес аналізу великих web даних** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 4(147). – Дніпро, 2023. – С.21 – 28.

Ця стаття представляє огляд математичних моделей та інформаційних технологій, що використовуються для бізнес аналізу великих web даних. У статті розглянуто такі методи аналізу даних, як машинне навчання, аналіз соціальних мереж, алгоритми графів та аналітика великих даних. Важливість використання цих методів у бізнесі обґрунтовується їхньою здатністю обробляти величезні обсяги даних, виділяти значні тренди та допомагати у прийнятті важливих рішень. Стаття також містить порівняльну таблицю різних методів та їх ефективність. В цілому цей огляд надає цінну інформацію для дослідників, фахівців з аналітики даних та бізнес аналітиків, які хочуть використовувати новітні технології для ефективного аналізу web даних.

Бібл. 9.

УДК 004.934

Дмитрієва І.С., Бімалов Д.В. **Розробка програмного модулю для ідентифікація емоційного стану користувача** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 4(147). – Дніпро, 2023. – С.29 – 34.

Дослідження ідентифікації емоцій у текстовому спілкуванні є актуальним напрямком досліджень в галузі обробки природної мови та машинного навчання. Основна мета роботи полягає в розробці програмного модулю, який реалізує алгоритми та моделі автоматичної ідентифікації емоцій людини у текстових повідомленнях. В роботі емоційні слова знаходилися за допомогою аналізу семантики речення та розглянуто два алгоритми для визначення

емоцій векторний та бульовий. У ході дослідження визначилося, що бульовий алгоритм найбільше підходить для пошуку емоційних слів.

Бібл. 2 ілл. 3.

УДК 519.6:519.85

Ігнаткін В.У., Дудніков В.С., Лучишин Т.Р., Алексєєнко С.В., Юшкевич О.П., Карпова Т.П., Хохлова Т.С., Хомош Ю.С., Тіхонов В.А. **Альтернатива методам середніх та найменших квадратів, які використовуються при обробці результатів науково технічних експериментів** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 4(147). – Дніпро, 2023. – С.35 – 48.

У статті приведена процедура обробки інформації у реальному масштабі часу та пропонується операція псевдозвороту, яка виконується за допомогою рекурентних формул. Ця процедура є процедурою послідовного оновлення (зі зсувом) по стовбцям матриці заданих розмірів та її псевдозвороту на кожному кроці зміни інформації. Цей підхід є прямим та використовує переваги властиві методу облямівки, при цьому контролюється правильність обчислень на кожному кроці, використовуючи умови Пенроуза. Необхідність псевдозвороту може виникнути при оптимізації, прогнозуванні тих чи інших параметрів і характеристик систем різного призначення, а також в різноманітних задачах лінійної алгебри, статистики, представленні структури одержаних рішень і зрозуміти зміст некоректності рішення, що виникає, в сенсі Адомара Тихонова і побачити шляхи регуляризації таких рішень.

Відмічено, що у існуючому методі найменших квадратів, при визначенні коефіцієнтів апроксимуючих поліномів, спостерігається зсув оцінок при збільшенні шумів у знайдених даних, так як сказався вплив шумів попередніх етапів обробки інформації.

Бібл. 17, іл.1.

УДК 629.7; 004:02

Лабуткіна Т.В., Ананко Р.В. **Глобальне покриття навколоземного простору зонами використання пристроїв його спостереження: концепція і алгоритми** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 4(147). – Дніпро, 2023. – С.49 – 70.

Представлені результати дослідження в рамках задачі забезпечення повного покриття заданої області висот над поверхнею Землі (області простору між двома сферами зі спільним центром у центрі Землі) миттєвими зонами можливого застосування пристроїв спостереження орбітального базування, розташованими на космічних апаратах в різновисоких орбітальних угрупованнях на колових орбітах. У загальному випадку вирішення задачі передбачає використання декількох різновисоких орбітальних угруповань на колових квазіполярних орбітах, які в спрощеній постановці задачі прийняті полярними. Миттєва зона можливого застосування пристрою спостереження спрощено прийнята у вигляді конусу. Розглянуті випадки застосування пристроїв спостережень «вверх» (над площиною миттєвого місцевого горизонту космічного апарату, що є носієм пристрою спостереження) і спостережень «вниз» (під цією площиною). Запропонована концепція вирішення задачі, яка базується на виборі (на основі розвинення методів застосування відомих алгоритмів) такої структури кожного орбітального угруповання, яка забезпечить неперервне покриття частини заданого простору спостереження

ня (області гарантованого спостереження), границі якої відсунуті від розташування пристроїв спостереження, а після того – заповнення простору цими областями. Робота присвячена космічній тематиці, але, узагальнивши постановку задачі, варіюючи низку умов цієї постановки та змінюючи «масштаб» вхідних даних, можна прийти до різноманіття технічних задач, де будуть доцільні і прийнятні (частково або повністю) запропонована концепція та алгоритми, застосовані при її реалізації.

Бібл. 17.

УДК 004.021

Атаманюк О.В., Легеза В.П. **Метод побудови цифрового двійника віброзахисного процесу** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 4(147). – Дніпро, 2023. – С.71 – 81.

Розглядаються різні підходи до побудови цифрових двійників, зокрема різні математичні моделі засновані на гасниках. Запропоновано новий метод створення цифрових двійників віброзахисного процесу, в основі якого лежить кульовий гасник. Також розглянуто питання оптимального налаштування гасника для зменшення руйнівної сили вимушених коливань та запропоновано аналітично графічний метод пошуку оптимальних налаштувань системи. Наведено блок схеми запропонованих методів, відображено результати аналізу ефективності застосування запропонованого методу, а також дослідження швидкої в залежності від реалізації.

Бібл. 9, іл. 5.

УДК 004.415.53

Бубенко М.О., Герасимов В.В., Карпенко Н.В., Морозов О.С. **Аналіз методів тестування веб додатків** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 4(147). – Дніпро, 2023. – С.82 – 90.

У статті розглядається тестування веб додатків через призму власного досвіду роботи в ІТ компанії. Показано, що під час тестування додатків різних видів слід звертати увагу на їх особливості та на критичні зони, які повинні бути перевірені. Авторами зроблено порівняння тестування web додатків та desktop програм за різними параметрами.

Бібл. 11, табл. 2.

УДК 004.02

Тулуб В.О. **Методи підвищення рівня ефективності роботи автоматизованих систем** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 4(147). – Дніпро, 2023. – С.91 – 99.

Автоматизовані системи відіграють ключову роль у сучасному світі, забезпечуючи ефективність та автоматизацію різних процесів. Однак, з постійним розвитком технологій та збільшенням складності завдань, потрібне безперервне вдосконалення та підвищення ефективності цих систем. У даній статті досліджуються методи, що дозволяють підвищити рівень ефективності роботи автоматизованих систем. Проаналізовані різні аспекти, такі як оптимізація роботи, покращення продуктивності, скорочення часу виконання завдань, зниження помилок та підвищення точності. Основний мета статті звернена на методології

підвищення рівня ефективності. У статті також пропонується новий алгоритм підвищення ефективності автоматизованих систем. Алгоритм заснований на використанні сучасних технологій та підходів, таких як аналіз даних та оптимізація процесів.

Бібл. 8, табл. 1.

УДК 004.82

Могильний О.А. **Автоматизовані моделі обробки візуальної інформації** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 4(147). – Дніпро, 2023. – С.100 – 110.

В роботі проаналізовано існуючі методи та алгоритми обробки візуальної інформації. Результати аналізу дозволили виявити переваги та недоліки, а також визначити сфери застосування. Розглянуто сегментацію зображень, розпізнавання образів, класифікацію і детекцію об'єктів та інші аспекти обробки візуальної інформації. Основна мета дослідження полягає у створенні комплексної моделі, здатної автоматично обробляти та аналізувати різні форми візуальних даних. Модель показала високу ефективність та точність в обробці візуальних даних, включаючи завдання сегментації, розпізнавання та класифікації об'єктів. Результати дослідження підтверджують ефективність та значущість автоматизованих моделей обробки візуальної інформації. Запропонована модель може бути корисною для розробки нових інформаційних систем, які базуються на обробці візуальних даних та сприяти розвитку комп'ютерного зору та штучного інтелекту.

Бібл. 10, табл. 3

УДК 004.896

Солдатенко Д.В., Гнатушенко Вік.В. **Покращення ефективності глибокого навчання шляхом збільшення навчальних даних** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 4(147). – Дніпро, 2023. – С.111 – 122.

Метою цього дослідження є визначення оптимальної кількості вхідних даних та методів аугментації, необхідних для навчання згорткової нейронної мережі (CNN) для розпізнавання супутникових зображень. Через серію експериментів досліджується вплив кількості вхідних даних на точність, збіжність та узагальнення моделі, а також ефект різноманітних методів аугментації даних. Дослідження також досліджує стратегії виявлення точки насичення та зменшення ефектів перенавчання. Висновки можуть сприяти розвитку більш ефективних моделей розпізнавання супутникових зображень, покращенню роботи існуючих моделей.

Бібл. 8, іл. 5, табл. 1.

УДК 004.623

Ситник Р.С., Гнатушенко Вік. В. **Управління потоками даних в сучасній промисловості за допомогою блокчейну** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 4(147). – Дніпро, 2023. – С.123 – 131.

Дослідження застосування блокчейну для управління потоками даних в сучасних інформаційних системах промисловості та логістики. Проведено аналіз підходів до управління даними в інформаційних системах "Індустрії 4.0", порівняно алгоритми блокчейну з класичним підходом з іншими базами даних в клієнт серверній архітектурі. Зроблено вис

новок про можливість покращення ефективності та прозорості інформаційних систем та руху даних в них за допомогою нових методів.

Бібл. 6, іл. 3.

УДК 004.02

Ліп'яніна Гончаренко Х.В. **Інтелектуальний метод вибору локації для старту бізнесу в розумному місті** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 4(147). – Дніпро, 2023. – С.132 – 140.

Постановка проблеми полягає у необхідності пришвидшити процес вибору оптимальної локації для розміщення бізнесу в розумному місті. Це завдання є складним та довгостроковим, оскільки вимагає аналізу великої кількості даних та врахування різних факторів, таких як географічне положення, наявність конкуренції, потенційна клієнтська база та інші аспекти, що впливають на успішність бізнесу. Також важливо забезпечити швидкий доступ до інформації та надання підприємцям точних рекомендацій, щоб вони могли зробити обґрунтоване рішення щодо вибору локації свого бізнесу. Вирішення цієї проблеми дозволить забезпечити ефективне використання ресурсів та забезпечити успіх підприємства в розумному місті.

Метою дослідження є розробка інтелектуального методу вибору локації для старту бізнесу в розумному місті. Цей метод має на меті використовувати великі обсяги даних, зібраних з різних джерел, для визначення найбільш оптимальних місць для відкриття нового бізнесу. При розробці методу використовуються існуючі методи машинного навчання, такі як розпізнавання зображень, попередня обробка даних, класифікація та кластеризація числових даних.

Розроблено метод, реалізація якого дозволить рекомендувати оптимальні локації для бізнесу в розумних містах. Що сприятиме підвищенню задоволення клієнтів, покращенню якості життя та збільшенню прибутку підприємців. Інтелектуальний метод є потужним інструментом для вирішення проблем вибору локації для старту бізнесу в розумних містах.

Бібл. 11.

УДК 004.056.5

Гуда А.І., Кліщ С.М. **Фішинг як перший крок до отримання доступу** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 4(147). – Дніпро, 2023. – С.141 – 154.

Фішинг як термін який означає техніку надсилання фішингових повідомлень буде розглянуто базуючись на здобутках у відкритому доступі та з використанням перерахованих посилань. Буде проаналізовано процес фішингової атаки, а потім приділимо увагу технічним векторам того, як користувачі стають жертвами атаки. Нарешті, буде розглянуто існуючі дослідження фішингових атак та відповідні підходи до запобігання.

Бібл. 16, іл. 6, табл. 0.

УДК 004.9

Островська К.Ю., Стовпченко І.В., Печений Д.С. **Дослідження методів на основі нейронних мереж для аналізу тональності корпусу текстів** // Системні технології.

Регіональний міжвузівський збірник наукових праць. Випуск 4(147). – Дніпро, 2023. – С.155 – 167.

Об'єктом дослідження є методи з урахуванням нейронних мереж для аналізу тональності корпусу текстів. Для досягнення поставленої в роботі мети необхідно вирішити такі завдання: дивити теоретичний матеріал для навчання глибоких нейронних мереж та їх особливості стосовно обробки природної мови; вивчити документацію бібліотеки Tensorflow; розробити моделі згорткової та рекурентної нейронних мереж; розробити реалізацію лінійних та нелінійних методів класифікації на моделях мішка слів та Word2Vec; порівняти точність та інші показники якості реалізованих нейромережових моделей із класичними методами. Для візуалізації навчання використовується Tensorboard. У роботі показано перевагу класифікаторів на основі глибоких нейронних мереж над класичними методами класифікації, навіть якщо для векторних уявлень слів використовується модель Word2Vec. Найвищу точність для даного корпусу текстів має модель рекурентної нейронної мережі з LSTM блоками.

Бібл. 6.

UDC 004.6: 007.5: 004.94

Skalozub V.V., Horiachkin V.M., Murashov O.V. **Relational separable models of monitoring processes at variable and unclear observation intervals** // System technologies. N 4(147) Dnipro, 2023. P.3 – 20.

The article is devoted to the development of combined models, methods and tools designed to solve the current problems of modeling and analysis of monitoring process data, which are represented by time series and differ in variable or fuzzy observation intervals (CHRPNI).

As a result of the conducted research by means of numerical modeling, it was established that the introduction of combined process models in the case of PNEU is rational and effective. Examples of data analysis of monitoring processes of rehabilitation of diabetic patients showed certain possibilities of ensuring the accuracy of the results of the analysis of indicators and their short term forecasting.

Ref. 14, ill. 5.

UDC 004.62

Maliienko S.E., Selivorstova T.V. **Review of mathematical models and information technologies for business analysis of the big web data** // System technologies. N 4(147) Dnipro, 2023. P.21 – 28.

This article provides an overview of mathematical models and information technologies used for business analysis of big web data. The article discusses data analysis methods such as machine learning, social network analysis, graph algorithms, and big data analytics. The importance of using these methods in business is justified by their ability to process huge amounts of data, identify significant trends, and assist in making important decisions. The article also includes a comparative table of different methods and their effectiveness. Overall, this review provides valuable information for researchers, data analytics professionals, and business analysts who want to use the latest technologies for effective analysis of web data.

Ref. 9.

UDC 004.934

Dmytriieva I.S., Bimalov D.V. **Development of a software module for the identification of the emotional state of the user** // System technologies. N 4(147) Dnipro, 2023. P.29 – 34.

The study of the identification of emotions in text communication is an actual direction of research in the field of natural language processing and machine learning. The main goal of the work is to develop a software module that implements algorithms and models for automatic identification of human emotions in text messages. In the work, emotional words were found using the analysis of sentence semantics, and two algorithms for determining emotions vector and Boolean were considered. During the research, it was determined that the Boolean algorithm is most suitable for searching for emotional words.

Ref. 2, fig. 3.

UDC 519.6:519.85

Ignatkin V.U., Dudnikov V.S., Luchyshyn T.R., Alekseenko S.V., Yushkevich O.P., Karpova T.P., Khokhlova T.S., Khomosh Y. S., Tikhonov V.A. **Alternative to mean and least squares methods used in processing the results of scientific and technical experiments** // System technologies. N 4(147) Dnipro, 2023. P.35 – 48.

The article presents a procedure for processing information in real time and proposes a pseudo reversal operation, which is performed using recurrent formulas. This procedure is a procedure for successive updating (with a shift) along the columns of a matrix of given sizes and its turnover at each step of changing information. This approach is direct and takes advantage of the hemming method, while checking the correctness of the calculations at each step using Penrose conditions. The need for pseudo inversion may arise during optimization, prediction of certain parameters and characteristics of systems for various purposes, as well as in various problems of linear algebra, statistics, the structures of the obtained solutions are presented and understand the content of the emerging incorrectness of the solution in the sense of Adomar Tikhonov and see the ways of regularization. such decisions.

It is noted that in the existing least squares method, when determining the coefficients of approximating polynomials, there is a shift in estimates with an increase in noise in the data, since the influence of the noise of the previous stages of information processing is affected.

Ref. 17.

UDC 629:7; 004:02

Tetyana V., Ananko R. **Global near earth space coverage by zones of the use of its observation devices: concept and algorithms** // System technologies. N 4(147) Dnipro, 2023. P.49 – 70.

The results of the study are presented within the framework of the task of ensuring full coverage of a given area of heights above the Earth's surface (the area of space between two spheres with a common center at the center of the Earth) by instantaneous zones of possible application of orbital based surveillance devices located on spacecraft in orbital groups of different heights in circular orbits. In the general case, the solution of the problem involves the use of several orbital groupings of different heights on circular quasi polar orbits, which in the simplified statement of the problem are assumed to be polar. The instantaneous zone of possible application of the surveillance device is simplified in the form of a cone. The cases of using observation devices "up" (above the plane of the instantaneous local horizon of the spacecraft, which is the carrier of the observation device) and observations "down" (below this plane) are considered. The concept of solving the problem is proposed, which is based on the selection (based on the development of methods of applying known algorithms) of such a structure of each orbital grouping, which will ensure continuous coverage of a part of the given observation space (area of guaranteed observation), the boundaries of which are moved away from the location of observation devices, and then filling the space with these areas. The work is devoted to the space theme, but by generalizing the statement of the problem,

varying a number of conditions of this statement and changing the "scale" of the input data, it is possible to arrive at a variety of technical problems where the proposed concept and algorithms used in its implementation will be appropriate and acceptable (in part or in full). In particular, when some surveillance systems or systems of complex application of technical operations devices are created.

Ref. 17.

UDC 004.021

Atamaniuk O.V., Legeza V.P. **Method of creation a digital twin of a vibration protection process** // System technologies. N 4(147) Dnipro, 2023. P.71 – 81.

Various approaches to building digital twins, including various mathematical models based on dampers, are considered. A new method for creating digital twins of the vibration protection process is proposed, based on a spherical damper. The issue of optimal tuning of the damper to reduce the destructive force of forced vibrations is also considered, and an analytical graphical method for finding optimal system settings is proposed. The block schemas of the proposed methods are presented, the results of the analysis of the effectiveness of the proposed method are shown, as well as the study of the speed depending on the implementation.

Bibl. 9, ill. 5.

UDC 004.415.53

Bubenko M.O., Gerasimov V.V., Karpenko N.V., Morozov A.S. **Analysis of web application testing methods** // System technologies. N 4(147) Dnipro, 2023. P.82 – 90.

The article deals with testing web applications through the prism of my own experience in an IT company. It is shown that when testing applications of different types, one should pay attention to their features and critical areas that must be checked. The authors made a comparison of testing web applications and desktop programs according to different criteria.

Bibl.11, table 2.

UDC 004.02

Tulub V. **Methods of increasing the level efficiency of automated systems** // System technologies. N 4(147) Dnipro, 2023. P.91 – 99.

Automated systems play a key role in the modern world, ensuring efficiency and automation of various processes. However, with the constant development of technology and the increasing complexity of tasks, continuous improvement and efficiency of these systems is required. This article explores methods that can improve the efficiency of automated systems. Various aspects are analyzed, such as optimization of work, improvement of productivity, reduction of task execution time, reduction of errors, and increase of accuracy. The main goal of the article is to focus on the methodologies for increasing the level of efficiency. The table shows the methodologies with a description of their advantages, disadvantages, and areas of application. In addition, additional parameters such as the degree of automation, the degree of system flexibility, and the level of autonomy are proposed. The article also proposes a new algorithm for improving the efficiency of automated systems. The algorithm is based on the

use of modern technologies and approaches, such as data analysis and process optimization. The proposed algorithm has the potential to improve the efficiency of automated systems and can be adapted many times over. The research represents a significant contribution to the field of improving the efficiency of automated systems. The algorithm can be useful for researchers, engineers, automation professionals, and managers interested in improving and optimizing their systems.

Bibl.8.

UDC 004.82

Mohylnyi O. **Automated models of visual information processing** // System technologies. N 4(147) Dnipro, 2023. P.100 – 110.

The article presents a study devoted to the development and research of an automated model of visual information processing. The goal of the research was to create a comprehensive model capable of automatically processing and analyzing various forms of visual data, such as images and videos. The model is developed on the basis of a combined approach that combines various algorithms and methods of visual information processing. The literature review conducted within the scope of this study allowed us to study the existing methods and algorithms for visual information processing. Various image processing approaches were analyzed, including segmentation, pattern recognition, object classification and detection, video analysis, and other aspects. As a result of the review, the advantages and limitations of each approach were identified, as well as the areas of their application were determined. The developed model showed high accuracy and efficiency in visual data processing. It can successfully cope with the tasks of segmentation, recognition and classification of objects, as well as video analysis. The results of the study confirmed the superiority of the proposed model. Potential applications of the automated model are considered, such as medicine, robotics, security, and many others. However, limitations of the model such as computational resource requirements and quality of input data are also noted. Further development of this research can be aimed at optimizing the model, adapting it to specific tasks and expanding its functionality. In general, the study confirms the importance of automated models of visual information processing and its important place in modern technologies. The results of the research can be useful for the development of new systems based on visual data processing and contribute to progress in the field of computer vision and artificial intelligence.

Bibl.8.

UDC 004.896

Soldatenko D., Hnatushenko Vik. **Improving deep learning performance by augmenting training data** // System technologies. N 4(147) Dnipro, 2023. P.111 – 122.

This study aims to determine the optimal input data quantity and augmentation techniques required for training a convolutional neural network (CNN) for satellite image recognition. Through a series of experiments, the study investigates the impact of input data quantity on model accuracy, convergence, and generalization, as well as the effect of various data augmentation techniques. The study also explores strategies for identifying the saturation point and mitigating overtraining effects. The findings can contribute to the development of

more efficient satellite imagery recognition models, improve the performance of existing models.

Refs. 8., imgs. 5, table 1.

UDC 004.623

Sytnyk R., Hnatushenko Vik. **Management of data flows in modern industry using blockchain** // System technologies. N 4(147) Dnipro, 2023. P.123 – 131.

Research on the use of blockchain to manage data flows in modern information systems of industry and logistics. An analysis of approaches to data management in "Industry 4.0" information systems was carried out. Blockchain algorithms are compared with a client server architecture and classical databases. A conclusion was made about the possibility of increasing the efficiency and transparency of information systems and the movement of data in them with the help of Blockchain algorithms.

Refs. 6., imgs. 3.

UDC 004.02

Lipianina Honcharenko Kh. **Intellectual method for business location selection in smart cities** // System technologies. N 4(147) Dnipro, 2023. P.132 – 140.

The problem statement involves the need to expedite the process of selecting an optimal location for business placement in a smart city. This task is challenging and long term, requiring the analysis of extensive data and consideration of various factors that impact business success, such as geographical position, competition, potential customer base, and other relevant aspects. It is also crucial to provide entrepreneurs with fast access to information and precise recommendations to make informed decisions regarding their business location. Solving this problem will facilitate efficient resource utilization and ensure business success in a smart city.

The purpose of the study is to develop an intelligent method for choosing a location for starting a business in a smart city. This method aims to use large amounts of data collected from various sources to determine the most optimal locations for starting a new business. The method is based on existing machine learning techniques such as image recognition, data preprocessing, classification, and clustering of numerical data.

A method has been developed, the implementation of which will allow recommending optimal locations for business in smart cities. This will help to increase customer satisfaction, improve the quality of life and increase the profit of entrepreneurs. The intelligent method is a powerful tool for solving the problems of choosing a location for starting a business in smart cities.

Refs. 11.

UDC 004.056.5

Guda A.I., Klishch S.M. **Phishing like the first step to gaining access** // System technologies. N 4(147) Dnipro, 2023. P.141 – 154.

Phishing as a term that means the technique of sending phishing messages will be researched based on findings in public access and using the listed links. The process of a phish

ing attack will be analyzed, and then we will pay attention to the technical vectors of how users become victims of the attack. Finally, existing research on phishing attacks and related prevention approaches will be reviewed.

Ref.16, pic.6, tabl.0.

UDC 004.9

Ostrovska K.Yu., Stovchenko I.V., Pechenyi D. **Research of methods based on neural networks for the analysis of the tonality of the corps of the texts** // System technologies. N 4(147) Dnipro, 2023. P.155 – 167.

The object of the study is methods based on neural networks for analyzing the tonality of a corpus of texts. To achieve the goal set in the work, it is necessary to solve the following tasks: study the theoretical material for learning deep neural networks and their features in relation to natural language processing; study the documentation of the Tensorflow library; develop models of convolutional and recurrent neural networks; to develop the implementation of linear and non linear classification methods on bag of words and Word2Vec models; to compare the accuracy and other quality indicators of implemented neural network models with classical methods. Tensorboard is used for learning visualization. The work shows the superiority of classifiers based on deep neural networks over classical classification methods, even if the Word2Vec model is used for vector representations of words. The model of recurrent neural network with LSTM blocks has the highest accuracy for this corpus of texts.

Ref.6.

Системні технології

ЗБІРНИК НАУКОВИХ ПРАЦЬ

Випуск 4 (147)

Головний редактор: к.т.н., доц. Т.В. Селівьорстова

Технічний редактор та секретар збірки: к.т.н., доц. К.Ю. Островська

Здано до набору 26.05.2023. Підписано до друку 31.05.2023.

Формат 60x84 1/16. Друк - різнограф. Папір типограф.

Умов. друк арк. – 12,93. Обл.–видавн. арк. – 11,313.

Тираж 300 прим. Замовл. – 04/23

Український державний університет науки і технологій,
ІНІ «Інститут промислових та бізнес технологій»,
кафедра Інформаційних технологій та систем: ІВК «Системні технології»
49600, Дніпро, а/с 493

<http://journals.nmetau.edu.ua/index.php/st>

Свідоцтво про державну реєстрацію друкованого засобу масової інформації:

Серія КВ № 8684 від 23 квітня 2004 рік

Редакційна колегія

Селівьорстова Тетяна Віталіївна
(головний редактор)

доцент, кандидат технічних наук

Алпатов Анатолій Петрович

Член-кореспондент НАН України,
професор, доктор технічних наук

Архипов Олександр Євгенійович

професор, доктор технічних наук

Бабічев Сергій Анатолійович

доцент, доктор технічних наук

Білозьоров Василь Євгенович

професор,

доктор фізико-математичних наук

Гече Федір Елемирович

професор, доктор технічних наук

Гуда Антон Ігорович

(заст. головного редактора)

професор, доктор технічних наук

Гнатушенко Вікторія Володимирівна

(вчений секретар)

професор, доктор технічних наук

Гнатушенко Володимир Володимирович

професор, доктор технічних наук

Гожий Олександр Петрович

професор, доктор технічних наук

Єрьомін Олександр Олегович

професор, доктор технічних наук

Кіріченко Людмила Олегівна

професор, доктор технічних наук

Світличний Дмитро Святозарович

професор, доктор технічних наук

Скалозуб Владислав Васильович

професор, доктор технічних наук

Хандецький Володимир Сергійович

професор, доктор технічних наук

Український державний університет науки і технологій, ННІ «Інститут промислових та бізнес технологій», Україна

Інститут технічної механіки

НАНУ і ДКАУ, Україна

Національний технічний університет

України «Київський політехнічний інститут»
імені Ігоря Сікорського», Україна

Jan Evangelista Purkyně University
in Ústí nad Labem

Університет імені Яна Євангеліста Пуркіне,
Усті над Лабем, Чеська Республіка

Дніпровський національний університет
імені Олеся Гончара, Україна

Ужгородський національний університет,
Україна

Український державний університет науки і технологій, ННІ «Інститут промислових та бізнес технологій», Україна

Український державний університет науки і технологій, ННІ «Інститут промислових та бізнес технологій», Україна

Національний технічний університет
«Дніпровська політехніка», Україна

Чорноморський національний університет
імені П.Могили, Україна

Український державний університет науки і технологій, ННІ «Інститут промислових та бізнес технологій», Україна

Харківський національний університет
радіоелектроніки, Україна

Akademia Górniczo-Hutnicza

Краківська гірничо-металургійна академія
ім. С. Сташіца, Польща

Український державний університет науки і технологій, ННІ «Дніпровський інститут інфраструктури і транспорту» Україна

Дніпровський національний університет
імені Олеся Гончара, Україна