

Національний технічний університет
«Дніпровська політехніка»

INFORMATION TECHNOLOGY: COMPUTER SCIENCE, SOFTWARE ENGINEERING AND CYBER SECURITY

Випуск 2



Видавничий дім
«Гельветика»
2023

РЕДАКЦІЙНА КОЛЕГІЯ:

ГОЛОВНИЙ РЕДАКТОР:

Удовик Ірина Михайлівна, кандидат технічних наук, доцент, декан факультету інформаційних технологій, Національний технічний університет «Дніпровська політехніка», Україна.

ЧЛЕНИ РЕДАКЦІЙНОЇ КОЛЕГІЇ:

Алексєєв Михайло Олександрович, доктор технічних наук, професор, завідувач кафедри програмного забезпечення комп'ютерних систем, Національний технічний університет «Дніпровська політехніка», Україна;

Бердник Михайло Геннадійович, кандидат технічних наук, доцент, доцент кафедри програмного забезпечення комп'ютерних систем, Національний технічний університет «Дніпровська політехніка», Україна;

Кабак Леонід Віталійович, кандидат технічних наук, доцент, доцент кафедри програмного забезпечення комп'ютерних систем, Національний технічний університет «Дніпровська політехніка», Україна;

Корнієнко Валерій Іванович, доктор технічних наук, професор, завідувач кафедри безпеки інформації та телекомунікацій, Національний технічний університет «Дніпровська політехніка», Україна;

Корченко Анна Олександрівна, доктор технічних наук, доцент, доцент кафедри безпеки інформаційних технологій, Національний авіаційний університет, Україна;

Лактіонов Іван Сергійович, доктор технічних наук, доцент, професор кафедри програмного забезпечення комп'ютерних систем, Національний технічний університет «Дніпровська політехніка», Україна;

Литвин Василь Володимирович, доктор технічних наук, професор, завідувач кафедри інформаційних систем та мереж, Національний університет «Львівська політехніка», Україна;

Любченко Віра Вікторівна, доктор технічних наук, професор, професор кафедри системного програмного забезпечення, Одеський національний політехнічний університет, Україна;

Маттіас Реч, Ph.D, професор, кафедра мехатроніки, Університет Ройтлінгену, Німеччина;

Молоканова Валентина Михайлівна, доктор технічних наук, професор, професор кафедри системного аналізу та управління, Національний технічний університет «Дніпровська політехніка», Україна;

Мороз Борис Іванович, доктор технічних наук, професор, професор кафедри програмного забезпечення комп'ютерних систем, Національний технічний університет «Дніпровська політехніка», Україна;

Мулеса Оксана Юріївна, доктор технічних наук, професор, професор кафедри інформаційних систем та мереж, Ужгородський національний університет, Україна;

Рак Тарас Євгенович, доктор технічних наук, доцент, професор кафедри інформаційних технологій, ПЗВО «ІТ СТЕП Університет», Україна;

Савенко Олег Станіславович, доктор технічних наук, професор, професор кафедри комп'ютерної інженерії та системного програмування, декан факультету програмування та комп'ютерних і телекомунікаційних систем, Хмельницький національний університет, Україна;

Семенов Сергій Геннадійович, доктор технічних наук, професор, професор кафедри кібербезпеки та інформаційних технологій, Харківський національний економічний університет імені Семена Кузнеця, Україна;

Сироткіна Олена Ігорівна, кандидат технічних наук, доцент, доцент школи комп'ютерних наук, Віндзорський університет, Канада;

Швачич Геннадій Григорович, доктор технічних наук, професор, професор кафедри програмного забезпечення комп'ютерних систем, Національний технічний університет «Дніпровська політехніка», Україна.

Журнал схвалено до друку Вченою радою

Національного технічного університету «Дніпровська політехніка»

31 травня 2023 р., протокол № 6

Журнал включений до Переліку наукових фахових видань України

(категорія «Б») за спеціальностями 121 – Інженерія програмного забезпечення, 122 – Комп'ютерні науки, 123 – Комп'ютерна інженерія, 124 – Системний аналіз, 125 – Кібербезпека відповідно до наказів МОН України № 735 (додаток № 4) від 29.06.2021 р. та № 491 (додаток № 3) від 27.04.2023 р.

Науковий журнал «Information Technology: Computer Science, Software Engineering and Cyber Security» зареєстрований Міністерством юстиції України

(свідоцтво про державну реєстрацію друкованого засобу масової інформації серія KB № 24879-14819P від 17.06.2021 р.)

Офіційний сайт видання: www.journals.politehnica.dp.ua/index.php/it

Статті у виданні перевірені на наявність плагіату за допомогою програмного забезпечення StrikePlagiarism.com від польської компанії Plagiat.pl.

ISSN 2786-507X (Print)

ISSN 2786-5088 (Online)

© Національний технічний університет «Дніпровська політехніка», 2023

УДК 004.056.53 (045)

DOI <https://doi.org/10.32782/IT/2023-2-1>

Анна КОРЧЕНКО

доктор технічних наук, професор, професор кафедри безпеки інформації та телекомунікацій, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, Дніпро, Україна, 49005, annakor@ukr.net,

ORCID: 0000-0003-0016-1966,

Scopus Author ID: 56029291400

Роман КАРПЮК

аспірант кафедри кібербезпеки, Львівський національний університет імені Івана Франка, вул. Університетська, 1, Львів, Україна, 79000, roman.karpiuk@lnu.edu.ua,

ORCID: 0009-0001-0053-1608

Тарас ПАРАЩУК

аспірант кафедри безпеки інформаційних технологій, Національний авіаційний університет, просп. Космонавта Комарова, 1, Київ, Україна, 03058, taras1039@ukr.net,

ORCID: 0000-0002-7014-761X,

Scopus Author ID: 56071538700

Сергій МАЦЮК

кандидат технічних наук, доцент, доцент кафедри безпеки інформації та телекомунікацій, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, Дніпро, Україна, 49005, matsiuk.s.m@nmu.one

ORCID: 0000-0001-6798-5500

Scopus Author ID: 57189702975

Бібліографічний опис статті: Корченко, А, Карпюк, Р, Паращук, Т., Мацюк, С. (2023). Метод формування параметрів та оцінювання загроз в соціотехнічних системах. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 3–11, doi: <https://doi.org/10.32782/IT/2023-2-1>

МЕТОД ФОРМУВАННЯ ПАРАМЕТРІВ ТА ОЦІНЮВАННЯ ЗАГРОЗ У СОЦІОТЕХНІЧНИХ СИСТЕМАХ

З кожним роком значна частина інформаційних технологій в соціотехнічних системах стрімко зростає. А розвиток відповідних систем безпосередньо залежить від відповідних інформаційних технологій, що використовуються в бізнес-процесах підприємств. Це впливає на появу нових загроз, уразливостей та атак, які пов'язані з інформаційною складовою. Рівень реалізації кібератак за останні роки стрімко зростає. Головною ціллю атак на соціотехнічні системи є персональні данні, облікові записи користувачів системи, комерційна таємниця, медична інформація, бази даних, переписка, дані платіжних карт тощо. Також, для попередження атак та виявлення потенційних загроз необхідно комплексно використовувати методи щодо оцінювання уразливостей та загроз, моделювання можливих впливів, а також світові практики та міжнародні стандарти і нормативні документи у галузі управління інформаційною безпекою. Проведений аналіз та розглянуті основні аспекти, які необхідні при формуванні параметрів для оцінювання загроз, створили базис для обґрунтування актуальності теми даної роботи. На основі сформованих аспектів було розроблено метод формування параметрів та оцінювання загроз в соціотехнічних системах. Метод враховує такі параметри, як: пріоритет, що показує ступень необхідності зменшення конкретної загрози; ступінь впливу на систему при реалізації певної загрози; ймовірність реалізації конкретної загрози в системі; чинник людської складової; чинник загрози, що характеризує на загальному рівні причини, які сприяють виникненню загрози. Також, відповідний метод, використовуючи отримані параметри повинен найбільш повно оцінити існуючі загрози в системі та може бути застосований при розробці методології профілювання співробітників.

Ключові слова: соціотехнічні системи, загрози, ризики, бізнес-процеси, чинник загрози, ранжування, експертний метод, ймовірність подій, лінгвістичні параметри, параметри загрози, інформаційна безпека.

Anna KORCHENKO

Doctor of Technical Sciences, Professor, Professor of the Department of Information Security and Telecommunications, National Technical University Dnipro Polytechnic, 19 Dmytra Yavornytskoho Avenue, Dnipro, Ukraine, 49005, annakor@ukr.net

ORCID: 0000-0003-0016-1966

Scopus Author ID: 56029291400

Roman KARPIUK

Postgraduate Student Cyber Security Department, Franko National University of Lviv, Universytetska str., Lviv, Ukraine, 79000, roman.karpiuk@lnu.edu.ua

ORCID: 0009-0001-0053-1608

Taras PARASCHUK

Postgraduate Student of Academic Department of IT-Security, National Aviation University, Kosmonavta Komarova Ave., 1, Kyiv, Ukraine, 03058, taras1039@ukr.net

ORCID: 0000-0002-7014-761X

Scopus Author ID: 56071538700

Sergii MATSIUK

Assistant Professor at the Department of Information Security and Telecommunications, National Technical University Dnipro Polytechnic, 19 Dmytra Yavornytskoho Avenue, Dnipro, Ukraine, 49005, matsiuk.s.m@nmu.one

ORCID: 0000-0001-6798-5500

Scopus Author ID: 57189702975

To cite this article: Korchenko, A, Karpiuk, R, Paraschuk, T., Matsiuk, S. (2023). Metod formuvannia parametriv ta otsiniuvannia zahroz v sotsiotekhnykh systemakh [Method of formation of parameters and assessment of threats in sociotechnical systems]. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 3–11, doi: <https://doi.org/10.32782/IT/2023-2-1>

METHOD OF FORMATION OF PARAMETERS AND ASSESSMENT OF THREATS IN SOCIOTECHNICAL SYSTEMS

Every year a significant part of information technology in sociotechnical systems is growing rapidly. And the development of relevant systems directly depends on the relevant information technologies used in the business processes of enterprises. This affects the emergence of new threats, vulnerabilities and attacks that are associated with the information component. The level of implementation of cyberattacks in recent years is growing rapidly. The main target of attacks on sociotechnical systems is personal data, user accounts of the system, trade secrets, medical information, databases, correspondence, payment card data, etc. Also, in order to prevent attacks and identify potential threats, it is necessary to comprehensively use methods for assessing vulnerabilities and threats, modeling possible impacts, as well as world practices and international standards and regulations in the field of information security management. The analysis carried out and the main aspects that are necessary in the formation of parameters for threat assessment are considered, created the basis for substantiating the relevance of the topic of this work. On the basis of the formed aspects, a method of forming parameters and assessing threats in sociotechnical systems has been developed. The method takes into account such parameters as: priority, showing the degree of need to reduce a specific threat; the degree of influence on the system in the implementation of a certain threat; the likelihood of a specific threat in the system; factor of the human component; threat factor, characterizing at the general level the causes that contribute to the occurrence of the threat. Also, the appropriate method, using the obtained parameters, should most fully assess the existing threats in the system and can be applied in the development of a methodology for profiling employees.

Key words: sociotechnical systems, threats, risks, business processes, threat factor, ranking, expert method, probability of events, linguistic parameters, threat parameters, information security.

Актуальність проблеми та наліз останніх досліджень і публікацій. З кожним роком значна частина інформаційних технологій в соціотехнічних системах стрімко зростає, оскільки використання їх дозволяє автоматизувати інформаційні процеси, заощадити мате-

ріальні та не матеріальні ресурси, зменшити фінансові ризики, збільшити конкурентоспроможність та інше.

Саме тому, подальший розвиток сучасних соціотехнічних систем безпосередньо залежить від інформаційних технологій (IT), що викорис-

товуються в бізнес-процесах підприємств. Але, в свою чергу, це також впливає на появу нових загроз, уразливостей та атак, які пов'язані з інформаційною складовою.

Аналіз останніх аналітичних досліджень (2020 Data Breach Investigations Report; What Does IBM's Cost of a Data Breach...; DDoS Attack Trends for Q3 2021; DDoS Attack Trends for Q4 2021; Cisco Annual Cybersecurity Report 2018; Cisco Annual Cybersecurity Report 2018) від компаній Cisco Systems Inc, DeviceLock, Cloudflare, IBM, Positive Technologies, що займаються інформаційною безпекою або продуктами у цій сфері показують, що рівень кібератак за останні роки стрімко зростає. Головною ціллю атак на соціотехнічні системи є персональні данні, облікові записи користувачів системи, комерційна таємниця, медична інформація, бази даних, переписка, дані платіжних карт тощо.

В роботах (Lasey, 2009; Коцюк; Маслова; Ліпкан, Максименко, Желіховський, 2006; Кравченко, 2018) підтверджуються тенденції зростання кібератак на соціотехнічні системи, та зазначається, що для попередження атак, виявлення потенційних загроз необхідно комплексно використовувати методи щодо оцінювання уразливостей та загроз, моделювання можливих впливів, а також світові практики та міжнародні стандарти і нормативні документи у галузі управління інформаційною безпекою.

Також в роботах (Расторгуєв, Литвиненко, 2014; Дудатьєв, 2014; Методи оцінювання уразливостей та оптимізації інформаційних систем) розглядають особливості функціонування соціотехнічних систем у сучасному довкіллі. Напрямок розвитку кожної із частин, соціальної та технічної складової показує, що зовнішні зміни інформаційного довкілля, які відбуваються з великою інтенсивністю, призводять до того, що швидкість зміни довкілля значно пришвидшується, порівняно зі змінами в самій системі. Як наслідок це призводить до появи різних інформаційних уразливостей у інформаційній та соціальній частині системи, якими можуть скористатися неавторизована сторона (НАС).

Виходячи з цього можна зазначити, що дане питання є актуальним і розгляд відповідного напрямку є обґрунтованим та потребує подальшого дослідження. Тому в даній роботі буде розглянуто формування загальних параметрів, які використовуються для опису конкретної загрози від НАС (порушника) в системі, а також за допомогою математичних і експертних методів та теорії множин буде сформовано метод параметрів та оцінювання можливих загроз для соціотехнічних систем.

Метою статті є розробка методу формування параметрів оцінки загроз в соціотехнічних системах. Перед розробкою або модифікацією систем інформаційного захисту важливо проаналізувати та оцінити можливі загрози в соціотехнічних системах. Оскільки визначення основних понять та аспектів дозволить створити параметри загроз та оцінити їх.

Класифікація загроз базується на основі досліджень атак на соціотехнічні системи, моделей порушника та ресурсів, сфери діяльності системи та її активи, організаційної структури та інше. Саме це дозволяє виділити декілька основних груп:

- загальні особливості діяльності та функціонування системи, її поточний стан (територіальне розміщення, особливості діяльності, поточний стан системи та довкілля в якому функціонує система);
- інформаційно-комунікаційні системи (засоби зв'язку, архітектура мережі, комунікаційні вузли, програмно-апаратні комплекси);
- матеріальні та не матеріальні активи системи (конфіденційна інформація, людські ресурси, фінансові активи);
- організаційна структура та процеси взаємодії між співробітниками (організація робочого процесу, бізнес-процеси, ієрархічні взаємозв'язки працівників).

Виходячи з класифікації можна визначити основні складові, які необхідно враховувати при оцінці загроз в соціотехнічних системах:

- імовірність виникнення конкретної загрози;
- ступінь впливу на систему в результаті реалізації загрози;
- критичні ресурси та активи, що є в системі;
- поточний стан соціотехнічної системи;
- бізнес-процеси, що закладені в основі функціонування системи;
- складність реалізації атаки порушником;
- модель порушника;
- структура соціотехнічної системи, її складові та взаємозв'язки між ними.

Основна частина. На базі проведеного аналізу видів загроз і чинників, які можуть впливати при проведенні оцінювання та визначивши основні аспекти загроз в соціотехнічних системах сформуємо параметри та оцінки.

Базуючись на введеній множині оціночних інтервалів EI для значень параметрів та визначених коефіцієнтах для кожного з інтервалів, які були описані в (Корченко, Мацюк, Чобаль, 2022), сформуємо таблицю можливих значень всіх оціночних інтервалів EI , які експерт буде використовувати при оцінюванні сформованих параметрів загроз (див. табл. 1).

Таблиця 1

Можливі значення оціночного інтервалу

Числові значення	Лінгвістичні значення частоти	Лінгвістичні значення рівня	Коефіцієнт
1–2	Рідко	Дуже низький	0,05
3–4	Іноді	Низький	0,1
5–6	Досить часто	Середній	0,125
7–8	Часто	Вище середнього	0,175
9–10	Дуже часто або завжди	Високий	0,25

Опишемо в загальному виді загрозу T_i^t за допомогою кортежу з параметрами:

$$T_i^t = TN_i, PT_i, DIS_i, FPT_i, HCF_i, TF_i, \quad (1)$$

в якому:

TN_i – назва та опис загрози;

PT_i – пріоритет, що показує ступінь необхідності зменшення конкретної загрози відносно всіх можливих загроз пов'язаних з бізнес-процесами (шкала оцінки від 1 до 10);

DIS_i – ступінь впливу на систему, що характеризує можливість впливу при реалізації конкретної загрози (шкала оцінки від «Дуже низький» до «Високий»);

FPT_i – ймовірність реалізації чи реалізація загрози в системі (шкала оцінки від «Дуже низький» до «Високий»);

HCF_i – чинник людської складової (шкала оцінки від «Дуже низький» до «Високий»);

TF_i – чинник, що сприяють виникненню загрози (значення «Контрольовані» чи «Неконтрольовані»).

Оцінки параметрів визначаються експертом або групою експертів на основі проведення аналізу всіх можливих загроз компанії, що пов'язанні з основними бізнес-процесами. Визначенні експертами оцінки для конкретної загрози повинні відповідати шкалі оцінювання для даної загрози, також процес оцінювання повинен відбуватися максимально об'єктивно.

Пріоритет. Визначення даного параметру відбувається на основі комплексного аналізу всіх загроз, що пов'язанні з бізнес-процесами компанії. Він характеризує ступінь необхідності зменшення конкретної загрози з урахування поточного стану системи, тобто впровадження засобів та методів щодо мінімізації загроз.

Наприклад, при оцінюванні експертом певної загрози, він може аргументувати значення пріоритету:

- впливом зовнішнього довкілля в якому функціонує дана система, тобто економічний, соціальний, екологічний, політичний стан і т. д.;
- внутрішнім станом системи, тобто складом та структурою персоналу їх функціональ-

ними обов'язками, поточними бізнес-процесами, критичними ресурсами системи, поточною інформаційною структурою, технічним чи програмним забезпеченням (ПЗ) системи і т. д.

Ймовірність реалізації загрози НАС. Для визначення точної оцінки для даного параметру необхідно проаналізувати систему захисту, провести аналіз критичності ресурсів соціотехнічної системи, можливі загрози з урахуванням моделі порушника та побудувати між ними математичні залежності. На основі аналізу (Ймовірність усвідомлення потенційним зловмисником реальної загрози безпеці) розглянемо поняття моделі порушника, складності реалізації загрози та її оцінки, що дозволить визначити ймовірність реалізації загрози для конкретної системи.

Складність реалізації загрози характеризується тим на скільки довго зможе знаходитись система в стані, коли загроза безпеці є реальною і є можливість для проведення атаки. Щодо урахування аспекту, який описує мінімально необхідну кількість матеріальних та не матеріальних ресурсів (ПЗ, обладнання, час та ресурси необхідні для проведення атаки, практичні навички, фінанси тощо), що необхідні НАС для реалізації атаки. Його оцінка є складною, оскільки наразі активно розвивається нелегальний ринок послуг та засобів реалізації загроз атак і доволі складно оцінити можливості потенційного порушника.

Рівень складності реалізації загрози може характеризуватись мірою невизначеності для потенційного порушника та наявності загрози системі. Математичною основою для такої міри є поняття інформаційної ентропії визначеної К. Шенноном (Николайчук, Воронич, 2010), в свою чергу даний аспект дозволив автору сформулювати математичну залежність між складністю реалізації і даним поняттям (Ймовірність усвідомлення потенційним зловмисником реальної загрози безпеці):

$$CTR_i = -\log_2(1 - PSR_i), \quad (2)$$

де PSR_i – величина надійності захисту ресурсу системи, або ймовірність відсутності загрози реалізації атаки в соціотехнічній системі.

В свою чергу, значення можна сформулювати та визначити по різному. Для даної роботи пропонується розглянути ймовірність надійності системи захисту ресурсу системи PSR_i , як узагальнюючий параметр на основі:

- цінності ресурсу для соціотехнічної системи з урахуванням моделі порушника та його можливих цілей при доступі в систему;
- надійності системи захисту конкретного ресурсу, тобто ступінь захищеності ресурсу від можливих атак, які були визначанні при побудові політики безпеки соціотехнічної системи.

Позначимо цінність ресурсу як RV_i , відповідно, ступінь захищеності ресурсу RP_i . При оцінці RV та RP експерту слід враховувати шкалу оцінки, наприклад, від 1 до 10 для цих параметрів. Далі встановимо залежності між ними у вигляді формули:

$$PSR = \frac{RV}{RP}. \quad (3)$$

Описана залежність між даними параметрами, при яких ймовірність загрози реалізації атаки в соціотехнічній системі була б мінімальною. Це можна показати лінійно, тобто чим цінніший ресурс, тим відповідний захист повинен бути більшим. Розглянемо всі можливі варіанти залежності:

1. При $PV = RP$, поточний стан захищеності ресурсу знаходиться на границі мінімально можливого захисту від реалізації загрози НАС;
2. При $PV < RP$, поточний стан захищеності ресурсу знаходиться в стані, коли чим більше значення RP відносно PV тим менше вірогідність реалізації загрози;
3. При $PV > RP$, поточний стан захищеності ресурсу знаходиться в стані, коли чим більше значення PV відносно RP тим більша можливість реалізації загрози НАС.

На основі залежностей сформуємо відповідні формули обрахунку ймовірності відсутності загрози реалізації атаки в соціотехнічній системі:

1. При $PV_i \leq RP_i$:

$$PSR_i = 1 - PSR_{const} * \frac{RV_i}{RP_i}, \quad (4)$$

де PSR_{const} – гранична ймовірність загрози реалізації атаки в системі ($PSR_{const} = 0,5$).

2. При $PV_i > RP_i$:

$$PSR_i = PSR_{const} * \frac{RP_i}{RV_i}, \quad (5)$$

де PSR_{const} – гранична ймовірність загрози реалізації атаки в системі ($PSR_{const} = 0,5$).

Тоді, з урахуванням формул (4) та (5) можна подати формулу (2) як:

1. При $PV_i \geq RP_i$:

$$CTR_i = -\log_2 \left(PSR_{const} * \frac{RV_i}{RP_i} \right). \quad (6)$$

2. При $PV_i > RP_i$:

$$CTR_i = -\log_2 \left(1 - PSR_{const} * \frac{RP_i}{RV_i} \right). \quad (7)$$

Базуючись на [16] та з урахуванням описаного в цій роботі, ймовірність реалізації загрози порушником PTR може бути визначена через складність реалізації певної загрози безпеці НАС на основі виразу (2), що можна записати у вигляді:

$$PTR_i = \frac{1}{2^{CTR_i}}. \quad (8)$$

Розглянемо два приклади формування PTR . Нехай оцінки експерта, що описують низку параметрів для першої загрози будуть $RV_1 = 6$ та $RP_1 = 8$, а для другої загрози $RV_2 = 7$ та $RP_2 = 4$.

Оскільки $RV_1 \leq RP_1$, тоді формула визначення CTR_1 буде:

$$\begin{aligned} CTR_1 &= -\log_2 \left(PSR_{const} * \frac{RV_1}{RP_1} \right) = \\ &= -\log_2 \left(\frac{6}{16} \right) = -\log_2 0,375 = 1,415. \end{aligned}$$

В свою чергу $RV_2 > RP_2$, тоді формула визначення CTR_2 буде мати вигляд:

$$\begin{aligned} CTR_2 &= -\log_2 \left(1 - PSR_{const} * \frac{RP_2}{RV_2} \right) = \\ &= -\log_2 \left(1 - \frac{4}{14} \right) = -\log_2 0,714 = 0,486. \end{aligned}$$

З розрахунків видно, що реалізація успішної атаки на першу загрозу в 2,91 рази складніша, чим на другу. Тепер, визначимо значення PTR_1 та PTR_2 :

$$PTR_1 = \frac{1}{2^{CTR_1}} = \frac{1}{2^{1,415}} = 0,375,$$

$$PTR_2 = \frac{1}{2^{CTR_2}} = \frac{1}{2^{0,486}} = 0,714.$$

Дані значення можна також інтерпретувати лінгвістично [18–19] PTR_1 = «Низький», а PTR_2 = «Вище середнього».

Ступінь впливу на систему. Визначення оцінки для певного параметру відбувається на основі можливого впливу на систему після реалізації загрози, що була спричинена порушником.

Сформуємо кортеж параметрів (ймовірність реалізації загрози, складність інфраструктури), які будуть впливати на оцінювання експертом ступеня впливу на систему:

$$DIS_i = CS_i, PSR_i, \quad (9)$$

де CS_i – поточний стан системи, ресурсу або активу відносно якого може бути реалізована загроза в певний момент часу;

PSR_i – величина надійності захисту ресурсу системи.

При оцінці CS_i експерту слід враховувати такі складові, як:

1. Бізнес-процеси, що можуть бути пов'язані або впливати на CS_i , оперуючи поняттями складності процесу, захищеності та доступності;
2. Складність побудови та функціонування соціотехнічної системи й суміжних процесів;
3. Персонал, що обслуговує чи використовує систему, ресурс або актив;
4. Інформаційні ресурси та активи, що можуть стати ціллю НАС;
5. Загальний поточний стан системи.

Мінімальне значення оцінки CS_i на шкалі означає, що поточний стан системи, ресурсу або активу знаходиться в найгіршому положенні за всіма вище описаним складовим і система або її складова потребує максимального вдосконалення. Відповідно до цього, чим вища оцінка тим поточний стан є кращим, шкала оцінки від 1 до 10.

Описавши основні параметри, визначимо ступінь впливу на систему DIS_i і сформуємо в загальному вигляді вираз для певної загрози:

$$DIS_i = CS_i * (1 - PSR_i). \quad (10)$$

Використаємо значення оцінки експерта, що розглядалися при ймовірності реалізації загрози. Для першої загрози $RV_1 = 6$, $RP_1 = 8$ та $CS_1 = 5$, а для другої – $RV_2 = 7$, $RP_2 = 4$, та $CS_2 = 8$, тоді:

$$\begin{aligned} DIS_1 &= CS_1 * \left(PSR_{const} * \frac{RV_1}{RP_1} \right) = \\ &= 5 * \frac{6}{16} = 5 * 0,375 = 1,875 \end{aligned}$$

або «Дуже низький»;

$$\begin{aligned} DIS_2 &= CS_2 * \left(1 - PSR_{const} * \frac{RP_2}{RV_2} \right) = \\ &= 8 * \left(1 - \frac{4}{14} \right) = 5,714 \end{aligned}$$

або «Середній».

З обрахунків видно, що вплив атаки при реалізації другої загрози в 3,05 рази вищий, чим при реалізації першої.

Чинник людської складової. Параметр показує залежність конкретної загрози від людського чинника, яка може бути використана НАС або може впливати на ймовірність реалізації конкретної загрози чи ступінь впливу на систему. Під людською складовою (чинником) розглядається конкретний працівник соціотехнічної системи, що має або може отримати доступ до ресурсів системи, з якими в подальшому буде пов'язана дана загроза.

Значення, які набуває цей параметр є якісними і їх визначає експерт, враховуючи структуру та особливості бізнес-процесів, що можуть бути пов'язаними з відповідною загрозою. Також слід зазначити, що така оцінка є лише загальною й в подальшому спрямована допомогти експерту в ефективному визначенні взаємозалежності між функціональними обов'язками та загрозами, якщо це необхідно.

Чинник загрози. Даний параметр визначає базис, який лежить в основі певної загрози, що в подальшому дозволить більш повно оцінити чинники походження загрози та здатність на них впливати. Значення параметру повинно визначатись на основі попередніх етапів оцінки параметрів і піддаватися контролю за допомогою вживання необхідних заходів по попередженню чи протидії загрозам.

Якщо вживання відповідних заходів не достатньо для зменшення чинників появи загрози або причини її виникнення є випадковими, тоді слід визначати, що чинники загрози є не контрольовані. Також, до випадкових чинників належать ті, для яких необхідні математичні очікування і дисперсії для визначення даних випадкових величин.

В свою чергу чинники, появи загрози, які можна визначити, описати та в подальшому оцінити називаються контрольовані. При створенні і оцінці загроз системи необхідно кожному експерту максимально точно описувати чинники, що лежать в основі загроз, тобто максимальна кількість існуючих загроз повинна бути визначена параметром «Контрольовані».

Саме вище описанні параметри дозволять максимально точно та об'єктивно провести оцінювання складової кортежу T_i^t . Розглянемо множину можливих загроз, яку визначмо у вигляді:

$$T^t = \left\{ \bigcup_{i=1}^n T_i^t \right\} =$$

$$\{T_1^t, T_2^t, T_3^t, \dots, T_n^t\}, \quad (i = \overline{1, n}), \quad (11)$$

де n – визначає кількість всіх можливих загроз для певної соціотехнічної системи.

Відповідно до формули (11) при $m = 11$ розглянемо список загроз для певної компанії:

$$\mathbf{T}^t = \left\{ \bigcup_{i=1}^{11} \mathbf{T}_i^t \right\} = \{ \mathbf{T}_1^t, \mathbf{T}_2^t, \mathbf{T}_3^t, \dots, \mathbf{T}_{11}^t \}.$$

Далі, опишемо кожну загрозу $\mathbf{T}_i^t$ з використанням формули (1) у вигляді кортежу з параметрами $TN, PT, DIS, PTR, HCF, TF$ враховуючи межі шкал кожного із параметрів:

$$\mathbf{T}_1^t = TN_1, PT_1, DIS_1, PTR_1, HCF_1, TF_1 = \langle TN_1, 5, \text{«Вище середнього»}, \text{«Низький»}, \text{«Низький»}, \text{«Контрольовані»} \rangle;$$

$$\mathbf{T}_2^t = TN_2, PT_2, DIS_2, PTR_2, HCF_2, TF_2 = \langle TN_2, 7, \text{«Високий»}, \text{«Середній»}, \text{«Високий»}, \text{«Контрольовані»} \rangle;$$

$$\mathbf{T}_3^t = TN_3, PT_3, DIS_3, PTR_3, HCF_3, TF_3 = \langle TN_3, 7, \text{«Низький»}, \text{«Вище середнього»}, \text{«Середній»}, \text{«Контрольовані»} \rangle;$$

$$\mathbf{T}_4^t = TN_4, PT_4, DIS_4, PTR_4, HCF_4, TF_4 = \langle TN_4, 8, \text{«Високий»}, \text{«Середній»}, \text{«Високий»}, \text{«Не контрольовані»} \rangle;$$

$$\mathbf{T}_5^t = TN_5, PT_5, DIS_5, PTR_5, HCF_5, TF_5 =$$

$$\langle TN_5, 9, \text{«Вище середнього»}, \text{«Середній»}, \text{«Високий»}, \text{«Контрольовані»} \rangle;$$

$$\mathbf{T}_6^t = TN_6, PT_6, DIS_6, PTR_6, HCF_6, TF_6 =$$

$$\mathbf{T}_8^t = TN_8, PT_8, DIS_8, PTR_8, HCF_8, TF_8 = \langle TN_8, 7, \text{«СЕРЕДНІЙ»}, \text{«Середній»}, \text{«Середній»}, \text{«Контрольовані»} \rangle;$$

$$\mathbf{T}_9^t = TN_9, PT_9, DIS_9, PTR_9, HCF_9, TF_9 = \langle TN_9, 7, \text{«Дуже низький»}, \text{«Середній»}, \text{«Високий»}, \text{«Контрольовані»} \rangle;$$

$$\mathbf{T}_{10}^t = TN_{10}, PT_{10}, DIS_{10}, PTR_{10}, HCF_{10}, TF_{10} = \langle TN_{10}, 4, \text{«Вище середнього»}, \text{«Низький»}, \text{«Середній»}, \text{«Контрольовані»} \rangle;$$

$$\mathbf{T}_{11}^t = TN_{11}, PT_{11}, DIS_{11}, PTR_{11}, HCF_{11}, TF_{11} = \langle TN_{11}, 3, \text{«Вище середнього»}, \text{«Дуже низький»}, \text{«Низький»}, \text{«Не контрольовані»} \rangle.$$

Таблиця 2

Приклад можливих загроз

Номер загрози (I)	Назва загрози (TN)	Пріоритет (PT)	Ступінь впливу на систему (DIS)	Ймовірність реалізації загрози НАС (PTR)	Чинник людської складової (NCF)	Чинник загрози (TF)
1	Отримання доступу до внутрішньої поштової системи	5	Вище середнього	Низький	Низький	Контрольовані
2	Отримання доступу до інформації (документів, носіїв, засобів зберігання) пов'язаних з КІ	7	Високий	Середній	Високий	Контрольовані
3	Отримання доступу до робочого телефону	7	Низький	Вище середнього	Середній	Контрольовані
4	Соціальний інжиніринг персоналу	8	Високий	Середній	Високий	Не контрольовані
5	Отримання доступу до знищених документів чи ресурсів	9	Вище середнього	Середній	Високий	Контрольовані
6	Отримання доступу до внутрішньої системи документообігу та суміжних ресурсів	4	Високий	Низький	Середній	Контрольовані
7	Витік інформації про систему через відкриті інформаційні ресурси	6	Низький	Середній	Високий	Контрольовані
8	Отримання інформації про структуру, особливості функціонування системи	7	Середній	Середній	Середній	Контрольовані
9	Викрадення чи отримання дублікату ключа карти	7	Дуже низький	Середній	Високий	Контрольовані
10	Викрадення персональних даних персоналу системи	4	Вище середнього	Низький	Середній	Контрольовані
11	Втрата інформації через стихійні лиха природного чи техногенного характеру	3	Вище середнього	Дуже низький	Низький	Не контрольовані

Описані приклади кортежів загроз для даної компанії визначмо в таблиці додавши назву/описання кожної із них (див. табл. 2).

Описана таблиця дозволяє отримати характеристику певної загрози на базі описаних параметрів, показати відношення між загрозами та встановити певні залежності між параметрами і типом загрози.

Висновки. Оцінка загроз та ризиків в соціотехнічних системах є питанням не новим, але важливим оскільки на даний момент кількість

атак на інформаційні системи, що пов'язані з існуючими загрозами та ризиками зростає. Тому запропонований метод формування параметрів та оцінювання загроз повинен найбільш повно оцінити існуючі загрози системі (встановити чинник загрози, ступінь впливу на соціотехнічну систему, пріоритет, ймовірність реалізації загрози, чинник людської складової). Даний метод може використовуватись при розробці методології профілювання співробітників.

ЛІТЕРАТУРА:

1. 2020 Data Breach Investigations Report. URL: <https://enterprise.verizon.com/resources/reports/2020-data-breach-investigations-report.pdf> (дата звернення: 22.12.2022).
2. What Does IBM's Cost of a Data Breach Report 2021 Mean for Businesses? URL: <https://www.transperere.com/blog/ibm-data-breach-report-2021/> (дата звернення: 22.12.2021).
3. DDoS Attack Trends for Q3 2021. URL: <https://blog.cloudflare.com/ddos-attack-trends-for-2021-q3/> (дата звернення: 26.12.2022).
4. DDoS Attack Trends for Q4 2021. URL: <https://blog.cloudflare.com/ddos-attack-trends-for-2021-q4/> (дата звернення: 26.12.2022).
5. Cisco Annual Cybersecurity Report 2018. URL: https://www.cisco.com/c/dam/m/hu_hu/campaigns/security-hub/pdf/acr-2018.pdf (дата звернення: 28.12.2022).
6. Cisco Annual Cybersecurity Report 2018. URL: https://www.cisco.com/c/dam/m/en_hk/ciscolive/2020-ciso-benchmark-cybersecurity-series.pdf (дата звернення: 28.12.2022).
7. David Lacey (2009). Managing the Human Factor in Information Security, How to win over staff and influence business managers, Chichester, John Wiley & Sons Ltd.
8. Коцюк Ю. А. Роль людського чинника у питаннях захисту інформаційних систем. URL: <https://psj.oa.edu.ua/articles/2012/n20/%D0%9A%D0%BE%D1%86%D1%8E%D0%BA.pdf> (дата звернення: 20.09.2022).
9. Маслова Ю.Ю. Інформаційна безпека і людський фактор. URL: <http://webcache.googleusercontent.com/search?q=cache:XmsbYCGJ78gJ:journals.dut.edu.ua/index.php/dataprotect/article/view/2462/2362+&cd=6&hl=uk&ct=clnk&gl=ua> (дата звернення: 20.09.2022).
10. Ліпкан В. А., Максименко Ю. Є., Желіховський В. М. Інформаційна безпека України в умовах євроінтеграції : навчальний посібник. Київ : КНТ, 2006. 280 с. (Серія: Національна і міжнародна безпека).
11. Кравченко С. І. Безпека соціотехнічних систем. *НБІ технології*. 2018.
12. Расторгуєв В.П., Литвиненко М.В. Інформаційні операції в мережі / В.П. Расторгуєв, М.В. Литвиненко. АНО ЦСОІП, 2014. 128 с.
13. Дудатьєв А.В. Теоретичні аспекти та технології керованого хаосу для реалізації комплексного інформаційного захисту соціотехнічних систем. *Інформаційні технології та комп'ютерна інженерія*. 2014. № 2 (30). С. 28–32.
14. Методи оцінювання уразливостей та оптимізації інформаційних систем в умовах інформаційних впливів. URL: <https://er.nau.edu.ua/bitstream/NAU/14308/1/Диссер.pdf> (дата звернення: 23.09.2022).
15. Метод формування параметрів функціональних обов'язків для оцінки загроз в соціотехнічних системах / А. Корченко, С. Мацюк, О. Чобаль, О. Кручинін, Т. Паращук. *Information Technology: Computer Science, Software Engineering and Cyber Security*. 2022. № 3. С. 19–26.
16. Ймовірність усвідомлення потенційним зловмисником реальної загрози безпеці. URL: https://studme.org/180232/informatika/veroyatnost_realizatsii_potentsialnym_narushitelem_realnoy_ugrozy_bezopasnosti.
17. Николайчук, Я.М., Воронич А. Р. Теоретичні основи мір ентропії їх застосування в інформаційних технологіях формування та опрацювання сигналів. *Оптико-електронні інформаційно-енергетичні технології*. 2010. № 1. С. 50–63.
18. Корченко О.Г. Побудова систем захисту інформації на нечітких множинах: теорія та практичні рішення. Київ : МК-Прес, 2006. 320 с.
19. Анна Корченко, Методи ідентифікації аномальних станів для систем виявлення вторгнень : монографія, Київ : ЦП «Компринт», 2019. 361 с.

REFERENCES:

1. 2020 Data Breach Investigations Report. URL: <https://enterprise.verizon.com/resources/reports/2020-data-breach-investigations-report.pdf> (data zvernennia: 22.12.2022).
2. What Does IBM's Cost of a Data Breach Report 2021 Mean for Businesses? URL: <https://www.transpere.com/blog/ibm-data-breach-report-2021/> (data zvernennia: 22.12.2021).
3. DDoS Attack Trends for Q3 2021. URL: <https://blog.cloudflare.com/ddos-attack-trends-for-2021-q3/> (data zvernennia: 26.12.2022).
4. DDoS Attack Trends for Q4 2021. URL: <https://blog.cloudflare.com/ddos-attack-trends-for-2021-q4/> (data zvernennia: 26.12.2022).
5. Cisco Annual Cybersecurity Report 2018. URL: https://www.cisco.com/c/dam/m/ru_ru/campaigns/security-hub/pdf/acr-2018.pdf (data zvernennia: 28.12.2022).
6. Cisco Annual Cybersecurity Report 2018. URL: https://www.cisco.com/c/dam/m/en_hk/ciscolive/2020-ciso-benchmark-cybersecurity-series.pdf (data zvernennia: 28.12.2022).
7. David Lacey (2009). Managing the Human Factor in Information Security, How to win over staff and influence business managers, Chichester, John Wiley & Sons Ltd.
8. Kotsiuk Yu. A. Rol liudskoho chynnyka u pytanniakh zakhystu informatsiinykh system. URL: <https://psj.ua.edu.ua/articles/2012/n20/%D0%9A%D0%BE%D1%86%D1%8E%D0%BA.pdf> (data zvernennia: 20.09.2022).
9. Maslova Yu. Iu. Informatsiina bezpeka i liudskyi faktor. URL: <http://webcache.googleusercontent.com/search?q=cache:XmsbYCGJ78gJ:journals.dut.edu.ua/index.php/dataprotect/article/view/2462/2362+&cd=6&hl=uk&ct=clnk&gl=ua> (data zvernennia: 20.09.2022).
10. Lipkan V. A., Maksymenko Yu. Ye., Zhelikhovskiy V. M. Informatsiina bezpeka Ukrainy v umovakh yevrointehratsii: Navchalnyi posibnyk. K.: KNT, 2006. 280 s. (Serii: Natsionalna i mizhnarodna bezpeka).
11. Kravchenko S. I. Bezpeka sotsiotekhnykh system. *NBI tekhnologii*. 2018. T. 12. № 2.
12. Rastorhuiev V.P., Lytvynenko M.V. Informatsiini operatsii v merezhi / V.P. Rastorhuiev, M.V. Lytvynenko. *ANO TsSOIP*. 2014. 128 s.
13. Dudatiev A.V. Teoretychni aspekty ta tekhnologii kerovanoho khaosu dlia realizatsii kompleksnoho informatsiinoho zakhystu sotsiotekhnykh system. *Informatsiini tekhnologii ta kompiuterna inzheneriia*. 2014. № 2 (30). S. 28–32.
14. Metody otsiniuvannia urazlyvostei ta optymizatsii informatsiinykh system v umovakh informatsiinykh vplyviv. URL: <https://er.nau.edu.ua/bitstream/NAU/14308/1/Диссеп.pdf> (data zvernennia: 23.09.2022).
15. Korchenko A. Metod formuvannia parametriv funktsionalnykh oboviazkiv dlia otsinky zahroz v sotsiotekhnykh systemakh / A. Korchenko, S. Matsiuk, O. Chobal, O. Kruchinin, T. Parashchuk. *Information Technology: Computer Science. Software Engineering and Cyber Security*. 2022. № 3. S. 19–26.
16. Ymovirnist usvidomlennia potentsiinykh zlovmysnykom realnoi zahrozy bezpetsi. URL: https://studme.org/180232/informatika/veroyatnost_realizatsii_potentsialnym_narushitelem_realnoy_ugrozy_bezopasnosti.
17. Nikolaichuk, Ya.M. Teoretychni osnovy mir entropii yikh zastosuvannia v informatsiinykh tekhnologiiakh formuvannia ta opratsiuvannia sygnaliv. *Optyko-elektronni informatsiino-enerhetychni tekhnologii*. 2010. № 1. S. 50–63.
18. Korchenko O.H. Pobudova system zakhystu informatsii na nechitkykh mnozhynakh: teoriia ta praktychni rishennia. Kyiv : MK-Pres, 2006. 320 s.
19. Anna Korchenko, Metody identyfikatsii anomalnykh staniv dlia system vyivlennia vtornnen : Monohrafiia, Kyiv : TsP "Kompriynt", 2019. 361 s.

УДК 519.8

DOI <https://doi.org/10.32782/IT/2023-2-2>

Ларуса КОРЯШКІНА

кандидат фізико-математичних наук, доцент, доцент кафедри системного аналізу і управління, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, м. Дніпро, Україна, 49005, koriashkina.l.s@nmu.one

ORCID: 0000-0001-6423-092X

Scopus-Author ID: 55844269100

Марина САЗОНОВА

кандидат фізико-математичних наук, доцент, дослідник секції оптимізації та теорії систем кафедри математики, Королівський технологічний інститут (KTH Royal Institute of Technology), офіс 3432, вул. Брінелльвеген, 8, м. Стокгольм, Швеція, SE-100 44, sazonova@kth.se

ORCID: 0000-0002-9513-9659

Микола ОДНОВОЛ

доцент кафедри системного аналізу і управління, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, м. Дніпро, Україна, 49005, odnovol.m.m@nmu.one

ORCID: 0000-0002-2022-7996

Бібліографічний опис статті: Коряшкіна, Л., Сазонова, М., Одновол, М. (2023). Алгоритми територіальної сегментації для мережі сервісних центрів із перекриттям зон обслуговування. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 12–25, doi: <https://doi.org/10.32782/IT/2023-2-2>

АЛГОРИТМИ ТЕРИТОРІАЛЬНОЇ СЕГМЕНТАЦІЇ ДЛЯ МЕРЕЖІ СЕРВІСНИХ ЦЕНТРІВ ІЗ ПЕРЕКРИТТЯМ ЗОН ОБСЛУГОВУВАННЯ

Розглядаються задачі територіальної сегментації заданого регіону на зони обслуговування сервісних центрів. Територія розбивається за критерієм відстані до k найближчих центрів. При моделюванні враховується наступне: самі центри можуть бути функціонуючими або новими, розміщеними у найбільш «вигідних» місцях; їх можливості можуть бути обмеженими. Наведено постановку та чисельний алгоритм розв'язання неперервної лінійної задачі оптимального мультиплексного розбиття обмежених двовимірних множин, яка є математичною моделлю задачі оптимального розміщення сервісних центрів з одночасною територіальною сегментацією на зони їх обслуговування. Продемонстровано роботу цього алгоритму та його окремих випадків: для фіксованого набору центрів; з відшукуванням їх оптимальних координат на заданій множині; з розміщенням додатково кількох центрів та перерозподілом сфер обслуговування для оновленої мережі сервісних центрів.

Ключові слова: територіальне зонування, оптимальне мультиплексне розбиття множин, неперервна модель, розміщення об'єктів, перерозподіл, недиференційовна оптимізація.

Larysa KORIASHKINA

PhD, Associate Professor, Associate Professor at the Department of System Analysis and Control Dnipro University of Technology, Dmytro Yavornytskyi ave., 19, Dnipro, Ukraine, 49005, koriashkina.l.s@nmu.one

ORCID: 0000-0001-6423-092X

Scopus-Author ID: 55844269100

Maryna SAZONOVA

PhD, Associate Professor, Researcher at the Section of Optimization and Systems Theory of the Department of Mathematics, KTH Royal Institute of Technology, office 3432, Brinellvägen str., 8, Stockholm, Sweden, SE-100 44, sazonova@kth.se

ORCID: 0000-0002-9513-9659

Mykola ODNOVOL

Associate Professor at the Department of System Analysis and Control, Dnipro University of Technology, Dmytro Yavornytskyi ave., 19, Dnipro, Ukraine, 49005, koriashkina.l.s@nmu.one

ORCID: 0000-0002-2022-7996

To cite this article: Koriashkina, L., Sazonova, M., Odnovol, M. (2023). Alhorytmy terytorialnoi sehmentatsii dlia merezhi servisnykh tsentriv iz perekryttiam zon obsluhovuvannia [Algorithms of territorial segmentation for a facility network with overlapping service zones]. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 12–25, doi: <https://doi.org/10.32782/IT/2023-2-2>

ALGORITHMS OF TERRITORIAL SEGMENTATION FOR FACILITY NETWORK WITH OVERLAPPING SERVICE ZONES

We consider the problems of territorial segmentation of a given region into service areas for facilities. The optimization criterion is the distance to the k closest centers. When modelling, we consider the following: the centers themselves can be functioning or newly created and located in the most favorable places; they can have limited options. The paper presents a statement and a numerical algorithm for solving a continuous linear problem of optimal multiplex partitioning of bounded sets under constraints. It is a mathematical model of the optimal location of service centers with simultaneous territorial segmentation into their service areas. We demonstrate the work of this algorithm and its cases: for a fixed set of centers, with finding their optimal location on a given set, with the placement of several additional centers and redistribution of service spheres for an updated service network.

Key words: territorial zoning, optimal multiplex partitioning of sets, continuous model, facility location, redistributing, non-differentiable optimization.

Вступ. Розміщення сервісних або виробничих підприємств є класичною проблемою, яка передбачає залучення методів, що оперують з географічною інформацією, включають кількісну оцінку попиту та пропозиції, а також оптимальну організацію комунікаційних мереж. Традиційний аналіз таких задач орієнтований на максимізацію економічної та соціальної вигоди або (та) на мінімізацію транспортних витрат і задоволення попиту на товари та послуги. Це можуть бути довгострокові цілі зростання бізнесу, залучення певних сегментів клієнтів або продаж обмеженої кількості певного продукту. Економічні, демографічні та інфраструктурні умови змінюються у часі, і, отже, місця розташування структурних об'єктів створюваної мережі мають бути адекватно зважені за різними критеріями. Останні, у свою чергу, можуть бути агреговані в єдину цільову функцію (залежно від параметрів, що настроюються), що дозволяє визначити єдиний розв'язок задачі в термінах зважених карт відстаней, прив'язаних до цільової функції. Саме така ідея лежить в основі математичного опису задач оптимальної організації сервісних або виробничих мереж у вигляді неперервних задач оптимального розбиття множин (Kiseleva E., 2015, p. 325–335; Коряшкіна Л.С., 2015, с. 28–36).

Без втрати спільності викладу, як сервісні центри розглядатимемо мережу торгових підприємств. При цьому, за аналогією з роботою (Boots B., 1997, p. 519–536), під моделюванням торгових зон розумітимемо засновані на геометрії процедури створення теоретичних зон з використанням характеристик магазину та припущень про поведінку споживача. Такий підхід є корисним, коли дані про переваги спо-

живачів відсутні, або їх отримання вимагає великої кількості часу або людських ресурсів. Такі моделі варто використовувати для прогнозування торгових зон, для визначення потенційних місць розташування нових об'єктів, оцінювання впливу цих та інших змін на існуючій множині об'єктів.

Будемо вважати, що споживач вибирає k ($k = 1, 2, 3, \dots$) найближчих або найбільш привабливих з N існуючих центрів. Тоді отримані в результаті сегментації торгові зони не будуть взаємовиключними і не визначатимуть просторові монополії.

Метою цієї роботи є розробка алгоритмів територіальної сегментації довільної обмеженої області при різних обмеженнях на основі моделей та методів розв'язання неперервних задач оптимального мультиплексного розбиття множин (OMPM) (Koriashkina L., 2015, p. 15–32). Враховуються такі фактори, як попит та пропозиція на ринку послуг, кількість та розташування існуючих сервісних центрів, неохоплені пропозицією території, місця найбільшої концентрації потенційних клієнтів.

Літературний огляд. Пошук оптимальних місць для розміщення ресурсів є загальною проблемою в різних областях знань (Sahin G., 2007, p. 2310–2331; Mallozzi L., 2018, p. 1–38). Правильне розташування об'єктів, пов'язаних з бізнесом, є важливим для зниження витрат і збільшення прибутку. Кількість можливих комбінацій розподілу ресурсів у реальних задачах зазвичай експоненційно збільшується із входними даними. Іноді знаходження найкращого розв'язку не може бути досягнуто за розумний час. Беручи до уваги NP-повноту задач розміщення, останні роботи використовують

евристику для отримання розумних рішень, як наприклад в роботі (Torkestani S., 2016, 1–22). Lin Y. et al. (Lin Y., 2018, p. 134–144) запропонували алгоритм ройового інтелекту для визначення місцезнаходження об'єкта.

Chen L. et al. [Chen L., 2019.] для оптимізації розміщення каршерингу в одну сторону використовують модель MILP. Представлене в роботі (Lage M.d.O., 2021) дослідження також відноситься до служби каршерингу і спрямоване на виявлення та вибір найкращих місць для встановлення станцій обміну всередині дилерської мережі компанії та обслуговуючої мережі. В роботі залучена географічна інформаційна система (ГІС), яка розраховує просторовий розподіл потенційного попиту на поїздку. Для визначення ідеального розташування станцій запропоновано дві моделі ієрархічного розташування об'єктів. В якості математичного апарату використовувались зважені діаграми Вороного і класична задача Вебера. Критерій оптимізації – максимізація попиту, коли оператор отримує прибуток за рахунок мінімальних перешкод, таких як відстань до станцій.

У статті (Murray A.T., 2020, p. 834–846) обговорюється нещодавно заново відкрита форма класичної проблеми розміщення одного об'єкта Вебера, яка передбачає розміщення виробничого підприємства, де є кілька джерел кожної необхідної сировини (вхідних матеріалів), розподілених по всьому регіону. Це означає, що вибір даного джерела сировини може відрізнятися залежно від місця розташування підприємства. По суті, це робить проблему невиключною навіть при розміщенні лише одного виробництва. Автори розглядають елементи проблеми Вебера, які розглядалися в літературі разом із запропонованими методами вирішення, висвітлюючи елементи, на сьогоднішній день не реалізовані на практиці – вибір розподілу між кількома джерелами даної сировини.

В роботі (Monteiro C.M., 2021) автори моделюють динаміку можливих клієнтів каршерингу, які шукають доступні ресурси та використовують цей сервіс. Клієнти можуть орендувати автомобіль, якщо є вільний автомобіль на станції відправлення клієнта та вільне місце для паркування на станції призначення клієнта. Якщо немає вільного автомобіля або місця для паркування, потенційні клієнти шукають ці ресурси на інших станціях поблизу місця відправлення та призначення. Разом із симуляцією на основі агентів автопарк оптимізується, щоб обслуговувати якомога більше клієнтів і уникнути марної витрати ресурсів. У цій роботі пропонується змішано-цілочисельна модель лінійного про-

грамування для оптимізації розміру автопарку послуги каршерингу для режимів одностороннього та туди-назад, імітуючи взаємодію клієнтів. Різні сценарії аналізуються з використанням реальних параметрів і просторових даних.

В статті (Wang X, 2018, p. 343–357) розроблено модель континуального наближення (КН) для розв'язання задачі динамічного розміщення об'єктів для великомасштабного зростаючого ринку. Тут визначається оптимальне розташування об'єкта та час розгортання, які мінімізують витрати на будівництво об'єкта та обслуговування клієнтів у горизонті планування. Щоб подолати обчислювальні труднощі, модель КН визначає оптимальну густину об'єктів у просторово-часовому континуумі.

Ще одним поширеним методом розподілу ресурсів є модель розміщення-розподілення (Murray A.T., 2019, p. 1106–1130; Itai Feigenbaum, 2017, p. 434–447), яка передбачає вибір найкращих місць для розміщення ключових об'єктів та визначення того, хто який об'єкт обслуговуватиме.

Існують й інші пов'язані моделі, засновані на проблемі покриття множини місцезнаходження, спрямовані на мінімізацію кількості об'єктів для задоволення попиту (Karatas M., 2017, p. 1040–1051; Karatas M., 2020; Rodriguez S.A., 2020).

Класичні моделі розміщення-розподілення зазвичай зводяться до проблеми р-медіани (PMP). Це одна з найбільш вивчених та використовуваних моделей локалізації для вирішення реальних завдань: враховуючи N точок-кандидатів для об'єктів та M точок попиту із заданою вагою (наприклад, кількість людей, які проживають по сусідству), вибираємо рішення P так, щоб сума витрат відстані від кожної точки M до найближчої точки P була якомога найменшою.

Задачі, в яких регіон, що розбивається, є континуальним, у науковій літературі називають неперервними задачами розбиття множин (Kiseleva E., 2015, p. 325–335). У деяких роботах до розв'язування задач розміщення-розбиття на площині з неперервно розподіленим попитом застосовуються методи обчислювальної геометрії, зокрема такі її фундаментальні об'єкти, як діаграми Вороного та їх узагальнення, наприклад (Kiseleva E., 2015, p. 3–12; Gerns A., p. 59–70; Коряшкіна Л.С., 2015, с. 59–76). В роботі (Mallozzi L., 2018, p. 203–220), враховуючи положення деяких об'єктів, вивчається форма оптимальних зон клієнтів у загальній площині попиту, мінімізуючи загальні середні витрати, які залежать від вартості установки

і відстаней переміщення. Також досліджується випадок розмірних об'єктів, що розміщуються.

У роботі (Zheng C., 2020, р. 34407–34417) розроблено метод оптимізації розташування підземної логістичної системи на базі метро (MULS) для транспортування вантажів метро у періоди непікового навантаження. Розроблено покращену р-медіанну модель, яка враховує чотири фактори впливу. Алгоритм найкоротшого шляху використовується для мінімізації транспортних витрат, тоді як витрати на інші фактори розраховуються на основі зібраних у процесі анкетування даних. Аби оптимально розташувати потенційні станції метро та перекроїти обсяг логістичних послуг шляхом додавання зважених умов використовується діаграма Вороного.

В даний час промисловість та академічні кола переживають еволюцію у розробці додатків для дронів наступного покоління, у тому числі автономних, здатних виконувати завдання без допомоги людини-оператора. В роботі (Lynskey J., 2019) автори запропонували алгоритм оптимального розміщення портів, щоб звести до мінімуму середню відстань, яку дрони повинні долати, на основі набору потенційних місць розташування портів для дронів і завдань, створених у цій галузі.

Моделювання розміщення-розподілу завжди вважалося основною частиною просторового аналізу. Та, як зазначено в роботі (Wolf Gert W., 2022, р. 2741–2775), публікації, пов'язані з розміщенням об'єктів, майже зникли з журналів в останній час. Автор припускає, що однією з основних причин цього є обмежені можливості модулів розміщення-розподілу, реалізованих у пакетах ГІС, що не дозволяє використовувати це програмне забезпечення для вирішення більш складних реальних проблем сьогодення. В зазначеній статті продемонстровано, як можна ефективно вирішувати задачі розміщення-розподілу за допомогою двох програмних пакетів CPLEX® та AIMMS®.

Дослідження (Alghanmi Nusaybah., 2022) спрямовано на огляд та аналіз існуючих моделей розміщення та розподілу POD під час надзвичайних ситуацій у галузі суспільної охорони здоров'я на основі повного та часткового розподілу точок попиту. Крім того, автори порівнюють існуючі моделі на основі їх ключових характеристик, сильних сторін та обмежень. Також обговорюються проблеми та майбутні напрямки досліджень моделей розташування-розподілу POD.

У статті (Lebedeva O., 2022, р. 427–433) розглядається проблема оптимізації зон обслуговування транспортних терміналів районів міста

та доставки вантажів з логістичних центрів. Для вирішення цієї проблеми пропонується застосувати адаптовану для транспортної мережі модель Вороного. Модель розроблено за допомогою геометричних інструментів, застосованих для знаходження карти найкоротших відстаней. Розбудова включає простір, що розбивається на кластери. Математичне формулювання моделі кластера може змінюватися в залежності від розміру міста. Чим більше місто, тим більшими є кількість і розміри кластерів. Для вирішення задачі використовувався модуль програмного комплексу MATLAB, зокрема, алгоритм побудови багатокутників Вороного. Можливості програмного забезпечення дозволили провести кластерний аналіз, а на основі порівняльного аналізу розрахувати ряд значень, що відображають просторове розташування.

Математичними моделями територіальної сегментації ринку послуг можуть виступати неперервні задачі оптимального мультиплексного розбиття множин (OMPM), особливо, коли зони можуть перекриватися (Koriashkina L., 2015, р. 15–32). Питанням, пов'язаним з теоретичним обґрунтуванням методів розв'язання таких задач у різних постановках, а також дослідженню властивостей рішень та практичних додатків таких завдань, присвячені роботи (Коряшкіна Л.С., 2017, с. 97–108; Koriashkina L., 2016, р. 232–240).

Матеріали та методи. Наведемо спочатку визначення розбиття k -го порядку обмеженої множини $\Omega \subset E_2$, що задається скінченною множиною точок простору, та можливі його узагальнення.

Нехай у просторі E_2 задані точки τ_i , $i = \overline{1, N}$, які далі називатимемо центрами. Введемо наступні позначення: $\mathbf{N} = \{1, 2, \dots, N\}$ – набір всіх індексів центрів; $M(\mathbf{N}, k)$ – множина всіх k -елементних підмножин множини $\mathbf{N}$, $|M(\mathbf{N}, k)| = C_N^k = L$; $\sigma_l = \{j_1^l, j_2^l, \dots, j_k^l\}$, $l = \overline{1, L}$, – елементи множини $M(\mathbf{N}, k)$; $d(x, y)$ – відстань між точками x і y (може бути L_p -метрикою

$$r_p(x, y) = \left(\sum_{i=1}^2 (x^i - y^i)^p \right)^{\frac{1}{p}} \text{ або } r_\infty(x, y) = \max_{i=1,2} |x^i - y^i|.$$

У подальшому, якщо не зазначено інше, будемо використовувати евклідову L_2 -метрику.

Розбиття k -го порядку обмеженої множини $\Omega = E_2$, що задається скінченною множиною точок $\tau_i, i = \overline{1, N}$ з E_2 – це сукупність підмножин (зон) k -го порядку для кожного можливого набору центрів $\{\tau_{j_1}, \tau_{j_2}, \dots, \tau_{j_k}\}$, $\sigma_l = \{j_1^l, j_2^l, \dots, j_k^l\} \in M(\mathbf{N}, k)$:

$$Z(\tau_1, \tau_2, \dots, \tau_N) = \bigcup_{l=1, L} Z_l(\tau_{j_1^l}, \tau_{j_2^l}, \dots, \tau_{j_k^l}),$$

зону $Z_l(\tau_{j_1^l}, \tau_{j_2^l}, \dots, \tau_{j_k^l})$ формують точки множини Ω , для яких центри $\tau_{j_1^l}, \tau_{j_2^l}, \dots, \tau_{j_k^l}$ є найближчими, тобто

$$Z_l(\tau_{j_1^l}, \tau_{j_2^l}, \dots, \tau_{j_k^l}) = \{x \in \Omega : \max_{j \in \sigma_l} d(x, \tau_j) \leq \min_{j \in N \setminus \sigma_l} d(x, \tau_j)\}, l = \overline{1, L}. \quad (1)$$

При цьому для будь-яких $\sigma_l = \{j_1^l, j_2^l, \dots, j_k^l\}$, $\sigma_m = \{j_1^m, j_2^m, \dots, j_k^m\} \in M(\mathbf{N}, k)$

$$mes(Z_l(\tau_{j_1^l}, \tau_{j_2^l}, \dots, \tau_{j_k^l}) \cap Z_m(\tau_{j_1^m}, \tau_{j_2^m}, \dots, \tau_{j_k^m})) = 0.$$

На практиці часто виникає необхідність задати різні рівні значущості центрів. Якщо кожному центру $\tau_i, i = \overline{1, N}$ поставити у відповідність деяке число $w_i > 0, i = \overline{1, N}$, – його вагу (мультіплікативну), то відстань між точками x і τ_i в (1) визначатиметься за формулою:

$$d(x, \tau_i) = \left(\frac{1}{w_i} \right) r_p(x, \tau_i), i = \overline{1, N}.$$

Якщо кожному центру τ_i відповідає так звана адитивна вага $a_i, i = \overline{1, N}$, то

$$d(x, \tau_i) = r_p(x, \tau_i) + a_i, i = \overline{1, N}.$$

Нехай на множині $\Omega \subset E_2$ задана функція попиту $\rho(x) > 0, x \in \Omega$. Потужністю $Capacity(\tau_i)$ центру τ_i будемо вважати зважену площу об'єднання всіх зон k -го порядку $Z_l(\tau_{j_1^l}, \tau_{j_2^l}, \dots, \tau_{j_k^l})$, для яких $i \in \sigma_l, l = \overline{1, L}$, тобто

$$Capacity(\tau_i) = \sum_{l=1, L} \int_{Z_l(\tau_{j_1^l}, \tau_{j_2^l}, \dots, \tau_{j_k^l})} \gamma_i^l \rho(x) dx, i = \overline{1, N},$$

де коефіцієнти γ_i^l для всіх $i = \overline{1, N}$, $\sigma_l = \{j_1^l, j_2^l, \dots, j_k^l\} \in M(\mathbf{N}, k), l = \overline{1, L}$, задовольняють умови:

$$0 \leq \gamma_i^l \leq 1, \gamma_{j_1^l}^l + \gamma_{j_2^l}^l + \dots + \gamma_{j_k^l}^l = 1. \quad (2)$$

В задачі територіальної сегментації кожен центр τ_i може мати свою обмежену потужність $b_i \geq 0$, і на розбитті множини Ω можуть накладатися обмеження:

$$Capacity(\tau_i) = b_i, i = \overline{1, p}, \\ Capacity(\tau_i) \leq b_i, i = \overline{p+1, N}. \quad (3)$$

Для обмежених множин $\Omega \subset E_2$ виконується рівність:

$$\bigcup_{l=1, L} Z_l(\tau_{j_1^l}, \tau_{j_2^l}, \dots, \tau_{j_k^l}) = \Omega. \quad (4)$$

Нехай $\Omega \subset E_2$ – обмежена замкнена множина; $\rho(x) > 0, x \in \Omega$, – задана функція, $\int_{\Omega} \rho(x) dx = S$; $\tau_i \in \Omega, i = \overline{1, N}$, – задані центри. Для існування розбиття k -го порядку з обмеженнями на потужності центрів (3) достатньо виконання наступних умов:

$$0 \leq b_i \leq S, i = \overline{1, p}; \quad (5)$$

$$\sum_{i=1}^p b_i \leq S \leq \sum_{i=1}^N b_i. \quad (6)$$

Розбиття k -го порядку з обмеженнями на потужності центрів може бути корисними, наприклад, оцінюючи реальні зони обслуговування для сервісних (торгових) підприємств за таких припущеннях:

– співіснування підприємств $\tau_i, i = \overline{1, N}$, що випускають один і той самий товар, причому максимальний обсяг товару, який підприємства можуть забезпечити, обмежений величинами $b_1, b_2, \dots, b_N$;

– константи γ_j^l відображають частину ринку, яку займає підприємство τ_j на території $Z_l(\tau_{j_1^l}, \tau_{j_2^l}, \dots, \tau_{j_k^l})$ серед підприємств $\{\tau_{j_1^l}, \tau_{j_2^l}, \dots, \tau_{j_k^l}\}$, що обслуговують цю територію (тут вважається, що $j \in \sigma_l$). Якщо припустити, що ринок товару розподілений між підприємствами по всій області Ω пропорційно їхнім потужностям, то для всіх $l = \overline{1, L}$ і $j = \overline{1, N}$, таких, що $j \in \sigma_l$, величина γ_j^l визначається так: $\gamma_j^l = b_j / \sum_{q: q \in \sigma_l} b_q$.

Тоді обмеження (3) вказують на те, що можливості центрів $\tau_i, i = \overline{1, p}$, мають бути повністю реалізовані, а для решти центрів $\tau_i, i = \overline{p+1, N}$, вони є обмеженими.

Перейдемо до викладу єдиного підходу щодо оптимально територіальної сегментації обмеженої області $\Omega \subset E_2$, в основі якого лежать моделі та методи розв'язання неперервних задач оптимального мультиплексного розбиття множин при обмеженнях з критерієм розбиття спеціального виду, що забезпечує врахування тієї чи іншої форми задання уподобань клієнтів та транспортних витрат.

Математична модель задачі оптимального розміщення нових сервісних центрів і перерозподілу сфер обслуговування

Нехай Ω – територія деякого регіону, на якій функціонує мережа сервісних підприємств;

$\widehat{\Omega} \subseteq \Omega$ – територія, де можуть бути розміщені нові центри; $\rho(x)$ – функція, що описує попит на послугу в точці x множини Ω ; N – загальна кількість сервісних підприємств в розширеній мережі; m – кількість функціонуючих центрів ($m < N$), решта – нові, які потрібно розмістити; S – загальний попит на послугу на заданій території Ω ; b_i – потужність i -го центру; $c(x, \tau_i)$ – вартість надання послуги клієнту в точці $x \in \Omega$ центром τ_i , яку будемо вважати пропорційною відстані між двома точками; a_i – вартість облаштування нового центру або модернізації існуючого в точці τ_i , або його фіксовані організаційні витрати, розраховані на одну умовну одиницю попиту, $i = \overline{1, N}$.

Для побудови математичної моделі задачі будемо вважати Ω обмеженою, замкнутою множиною з E_2 ; $\tau_i = (\tau_i^{(1)}, \tau_i^{(2)}) \in \Omega$, для всіх $i = \overline{1, N}$, – деякі точки, причому $\tau_1, \tau_2, \dots, \tau_m$ ($m < N$) є заданими, а координати решти центрів $\tau_{m+1}, \tau_{m+2}, \dots, \tau_N$ потрібно визначити, $\tau^{N-m} = (\tau_{m+1}, \tau_{m+2}, \dots, \tau_N) \in \widehat{\Omega} \times \dots \times \widehat{\Omega} = \widehat{\Omega}^{N-m}$.

Нехай $\Sigma_{\Omega}^{N,k}$ – клас всіх можливих розбиттів k -го порядку множини Ω на її підмножини $\Omega_{\sigma_1}, \Omega_{\sigma_2}, \dots, \Omega_{\sigma_L}$:

$$\begin{aligned} \Sigma_{\Omega}^{N,k} &= \left\{ \bar{\omega} = \{ \Omega_{\sigma_1}, \dots, \Omega_{\sigma_L} \} : \right. \\ &: \bigcup_{l=1}^L \Omega_{\sigma_l} = \Omega, \text{mes}(\Omega_{\sigma_i} \cap \Omega_{\sigma_j}) = 0; \\ &\left. \sigma_i, \sigma_j \in M(N, k), i \neq j, i, j = \overline{1, L} \right\}. \end{aligned}$$

Задача А-к. $F(\bar{\omega}, \tau^{N-m}) \rightarrow \min_{\substack{\bar{\omega} \in \Sigma_{\Omega}^{N,k} \\ \tau^{N-m} \in \widehat{\Omega}^{N-m}}}$,

$$F(\bar{\omega}, \tau^{N-m}) = \sum_{l=1}^L \int \sum_{i \in \sigma_l} (c(x, \tau_i) / w_i + a_i) \rho(x) dx,$$

$$\sum_{\substack{l=1 \\ l \in \sigma_l}}^L \int \gamma'_l \rho(x) dx = b_i, \quad i = \overline{1, p}, \quad (7)$$

$$\sum_{\substack{l=1 \\ l \in \sigma_l}}^L \int \gamma'_l \rho(x) dx \leq b_i, \quad i = \overline{p+1, N}. \quad (8)$$

Тут $x = (x^{(1)}, x^{(2)}) \in \Omega$; $c(x, \tau_i)$, $i = \overline{1, N}$ – обмежені, визначені на $\Omega \times \Omega$ функції. Функція $\rho(x)$ – обмежена, невід'ємна на Ω ; $w_i > 0$, $a_i \geq 0$, $b_i \geq 0$, $i = \overline{1, N}$ – задані числа; γ'_i задовольняють умови (3).

Пару $(\bar{\omega}^*, \tau^{N-m*})$, на якій досягається мінімальне значення функціоналу F і виконуються умови (7), (8), будемо називати **оптимальним розв'язком задачі А-к**.

Отже, математичною моделлю задачі оптимального розміщення додаткових (або всіх, коли $m = 0$) центрів мережі сервісних підприємств і перерозподілу сфер обслуговування є неперервна лінійна задача оптимального розбиття k -го порядку множини $\Omega \subset E_n$ на її підмножини $\Omega_{\sigma_1}, \Omega_{\sigma_2}, \dots, \Omega_{\sigma_L}$, що не перетинаються і серед яких можуть бути порожні, при обмеженнях у формі рівностей і нерівностей із розміщенням центрів $\tau_{m+1}, \tau_{m+2}, \dots, \tau_N$.

Метод розв'язання задачі А-к

Задачу **А-к** зведемо до задачі нескінченновимірною математичного програмування з булевими змінними.

Нехай $\bar{\omega} = \{ \Omega_{\sigma_1}, \dots, \Omega_{\sigma_l}, \dots, \Omega_{\sigma_L} \}$ – деяке розбиття k -го порядку множини Ω . Кожній точці $x \in \Omega_{\sigma_l}$, $l = \overline{1, L}$, поставимо у відповідність NL -вимірний вектор $\lambda^l(x) = (\lambda_1^l(x), \dots, \lambda_N^l(x))$, координати якого визначимо наступним чином:

$$\lambda_i^l(x) = \begin{cases} 1, & x \in \Omega_{\sigma_l} \quad i \in \sigma_l, \\ 0 & \text{у протилежному випадку, } i = \overline{1, N}, \end{cases} \quad (9)$$

де $\sigma_l = \{ j_1^l, j_2^l, \dots, j_k^l \}$ – набір індексів центрів $\{ \tau_{j_1^l}, \tau_{j_2^l}, \dots, \tau_{j_k^l} \}$, асоційованих з підмножиною Ω_{σ_l} .

Вектор-функцію $\lambda^l(x) = (\lambda_1^l(x), \dots, \lambda_N^l(x))$, визначену на Ω , з координатами, які задаються формулою (9), будемо називати характеристичною вектор-функцією підмножини Ω_{σ_l} , що є складовою розбиття k -го порядку множини Ω .

Задача В-к. $\min_{(\lambda(\cdot), \tau^{N-m}) \in \Gamma^k \times \widehat{\Omega}^{N-m}} I(\lambda(\cdot), \tau^{N-m})$,

$$I(\lambda(\cdot), \tau^{N-m}) =$$

$$\int \sum_{l=1}^L \left(\sum_{i=1}^N (c(x, \tau_i) / w_i + a_i) \lambda_i^l(x) \right) \rho(x) dx,$$

$$\Gamma^k = \{ \lambda(\cdot) : \lambda(\cdot) \in \Gamma_0^k,$$

$$\int \sum_{l=1}^L \gamma'_l \lambda_i^l(x) \rho(x) dx = b_i, \quad i = \overline{1, p},$$

$$\int \sum_{l=1}^L \gamma'_l \lambda_i^l(x) \rho(x) dx \leq b_i, \quad i = \overline{p+1, N} \};$$

$$\begin{aligned} \Gamma_0^k &= \{ \lambda^l(\cdot) = (\lambda_1^l(\cdot), \dots, \lambda_N^l(\cdot)) : \lambda_i^l(x) = \\ &= 0 \vee 1 \text{ для всіх } x \in \Omega, i = \overline{1, N} \end{aligned}$$

$$I = \overline{1, L}, \sum_{l=1}^L \lambda_i^l(x) = k, I = \overline{1, L}, \text{ м.в. для } x \in \Omega \}.$$

Далі від задачі **В-к** перейдемо до задачі зі значеннями $\lambda_i^l(\cdot)$, $i = \overline{1, N}$, $l = \overline{1, L}$, на відрізок $[0; 1]$ (за аналогією LP-релаксації в дискретній оптимізації).

Задача C-k. $\min_{(\lambda(\cdot), \tau^{N-m}) \in \Gamma_2^k \times \hat{\Omega}^{N-m}} I(\lambda(\cdot), \tau^{N-m}),$

$$\Gamma_2^k = \{\lambda(\cdot) : \lambda(\cdot) \in \Gamma_1^k,$$

$$\int_{\Omega} \gamma_i' \lambda_i'(x) \rho(x) dx = b_i, \quad i = \overline{1, p},$$

$$\int_{\Omega} \gamma_i' \lambda_i'(x) \rho(x) dx \leq b_i, \quad i = \overline{p+1, N}\};$$

$$\Gamma_1^k = \{\lambda'(\cdot) = (\lambda_1'(\cdot), \dots, \lambda_N'(\cdot)) : 0 \leq \lambda_i'(x) \leq 1$$

$$\text{для всіх } x \in \Omega, i = \overline{1, N},$$

$$I = \overline{1, L}, \sum_{i=1}^L \lambda_i'(x) = k, \quad I = \overline{1, L}, \text{ м.в. для } x \in \Omega\}.$$

Очевидно, $\Gamma_0^k \subset \Gamma_1^k$, внаслідок чого виконується включення $\Gamma^k \subset \Gamma_2^k$. Множина Γ_2^k є обмеженою, замкнутою, опуклою у просторі з нор-

мою $\lambda(\cdot) = \left(\int_{\Omega} \sum_{i=1}^L [\lambda_i'(x)]^2 dx \right)^{\frac{1}{2}}$; крім того, Γ_2^k

містить принаймні одну крайню точку.

Для кожного вектору $\tau^{N-m} \in \hat{\Omega}^{N-m}$ функціонал $I(\lambda(\cdot), \tau^{N-m})$ є лінійним, неперервним відносно вектор-функції $\lambda(\cdot)$ на Γ_2^k і, згідно теореми Вейерштрасса, досягає на цій множині своєї нижньої границі. Таким чином, задача **C-k** є розв'язною за $\lambda(\cdot)$. Серед множини точок Γ_2^k , в яких лінійний відносно $\lambda(\cdot)$ функціонал $I(\lambda(\cdot), \tau^{N-m})$ при фіксованому $\tau^{N-m} \in \hat{\Omega}^{N-m}$ досягає на множині Γ_2^k свого мінімального значення, знайдеться принаймні одна крайня точка множини Γ_2^k . Крайні точки Γ_2^k є характеристичними функціями деяких підмножин k -го порядку $\Omega_{\sigma_i}, I = \overline{1, L}$, що складають розбиття k -го порядку множини Ω при кожному фіксованому векторі $\tau^{N-m} \in \hat{\Omega}^{N-m}$. Отже, множина оптимальних розв'язків задачі **C-k** містить оптимальні розв'язки задачі **B-k**, і остання зводиться до задачі **C-k** і відбору з усіх оптимальних її розв'язків тих, які є розв'язками задачі **B-k**.

Для задачі **C-k** складаємо функціонал Лагранжа:

$$W((\lambda(\cdot), \tau^{N-m}), \psi) =$$

$$\int_{\Omega} \sum_{i=1}^L \sum_{j=1}^N [c(x, \tau_i) / w_i + a_i] \lambda_i'(x) \rho(x) dx +$$

$$+ \sum_{p=1}^N \psi_p \left[\int_{\Omega} \sum_{i=1}^L \gamma_i' \lambda_i'(x) \rho(x) dx - b_i \right] =$$

$$= \int_{\Omega} \sum_{i=1}^L \sum_{j=1}^N (c(x, \tau_i) / w_i + a_i + \gamma_i' \psi_i) \lambda_i'(x) \rho(x) dx - \sum_{i=1}^N \psi_i b_i,$$

який визначений на декартовому добутку

$$(\Gamma_2^k \times \hat{\Omega}^{N-m}) \times \Psi,$$

$$\text{де } \Psi = \{\psi \in R^N, \psi_i \geq 0, i = \overline{p+1, N}\}.$$

Пару $((\hat{\lambda}(\cdot), \hat{\tau}^{N-m}), \hat{\psi})$ будемо називати сідловою точкою функціонала Лагранжа $W((\lambda(\cdot), \tau^{N-m}), \psi)$ на множині $(\Gamma_2^k \times \hat{\Omega}^{N-m}) \times \Psi$,

якщо $\forall (\lambda(\cdot), \tau^{N-m}) \in \Gamma_2^k \times \hat{\Omega}^{N-m}$ і $\forall \psi \in \Psi$ виконується наступна нерівність:

$$W((\hat{\lambda}(\cdot), \hat{\tau}^{N-m}), \hat{\psi}) \leq W((\lambda(\cdot), \tau^{N-m}), \hat{\psi}) \leq W((\lambda(\cdot), \tau^{N-m}), \hat{\psi}).$$

Задача, двоїста до задачі **C-k**, має такий вигляд:

$$H(\psi) = \min_{(\lambda(\cdot), \tau^{N-m}) \in \Gamma_2^k \times \hat{\Omega}^{N-m}} W((\lambda(\cdot), \tau^{N-m}), \psi) \rightarrow \max_{\psi \in \Psi}. \quad (10)$$

Задача (10) є скінченновимірною, на відміну від **C-k**. Аби задачі **C-k** і (7) були пов'язані співвідношенням двоїстості $I_* = W^*$, і верхня грань в двоїстій задачі досягалась, необхідно і достатньо існування сідлової точки функціонала $W((\lambda(\cdot), \tau^{N-m}), \psi)$ в сенсі введеного визначення.

Відтак, розв'язання пари двоїстих задач **C-k** і (10) еквівалентно знаходженню сідлової точки функціоналу Лагранжа $W((\lambda(\cdot), \tau^{N-m}), \psi)$ на множині $(\Gamma_2^k \times \hat{\Omega}^{N-m}) \times \Psi$.

Зафіксуємо довільний вектор $\bar{\psi} \in \Psi$ і розглянемо задачу

$$W((\lambda(\cdot), \tau^{N-m}), \bar{\psi}) \rightarrow \min_{(\lambda(\cdot), \tau^{N-m}) \in \Gamma_2^k \times \hat{\Omega}^{N-m}}, \quad (11)$$

яка є задачею оптимального мультиплексного розбиття множини без обмежень, записаною відносно характеристичних вектор-функцій підмножин $\Omega_{\sigma_i}, I = \overline{1, L}$, що складають розбиття k -го порядку множини Ω . Метод розв'язання цієї задачі докладно описаний у роботі (Коряшкіна Л.С., 2015). Зробивши теоретичні викладки, аналогічні наведеним у зазначеній роботі, можна отримати оптимальний розв'язок задачі (11) у такому вигляді:

для $i = \overline{1, N}, I = \overline{1, L}$ і майже всіх $x \in \Omega$:

$$\lambda_i'(x) = \begin{cases} 1, & \text{якщо } c(x, \tau_i) / w_i + a_i + \gamma_i' \bar{\psi}_i \leq c(x, \tau_j) / w_j + a_j + \gamma_j' \bar{\psi}_j, \\ & \text{одночасно } \forall i \in \sigma_i, j \in N \setminus \sigma_i; \\ 0 & \text{у протилежному випадку, } i = \overline{1, N}; \end{cases} \quad (12)$$

де $\tau_{*m+1}, \dots, \tau_{*N}$ є оптимальним розв'язком задачі

$$G(\tau^{N-m}, \bar{\psi}) \rightarrow \min_{\tau^{N-m} \in \hat{\Omega}^{N-m}}, \quad (13)$$

$$G(\tau^{N-m}, \bar{\psi}) = \int_{\Omega} \min_{\sigma_i \in M(N,k)} \sum_{i=1, L} [c(x, \tau_i) / w_i + a_i + \gamma_i' \bar{\psi}_i] \rho(x) dx. \quad (14)$$

Відтак, аби для фіксованого набору центрів $\tau^{N-m} \in \hat{\Omega}^{N-m}$ можливе розбиття $(\Omega_{\sigma_1}, \dots, \Omega_{\sigma_i}, \dots, \Omega_{\sigma_N}) \in \sum_{\Omega}^N$ множини Ω було оптимальним для задачі **A-k**, необхідно і достатньо існування таких констант $\psi_1, \dots, \psi_p, \psi_{p+1}, \dots, \psi_N$ ($\psi_{p+1}, \dots, \psi_N$ – невід’ємні), що м. в. для $x \in \Omega$

$$c(x, \tau_i) / w_i + a_i + \gamma_i' \psi_i \leq c(x, \tau_j) / w_j + a_j + \gamma_j' \psi_j, \quad i \in \sigma_i, j \in N \setminus \sigma_i. \quad (15)$$

В точках x на оптимальній границі між підмножинами Ω_{σ_i} і Ω_{σ_j} в нерівності (12) досягається знак рівності.

Враховуючи (13) – (14), функціонал задачі (10), двоїстої до задачі **C-k**, в наступному уточненому вигляді:

$$H(\psi) = \min_{\tau^{N-m} \in \hat{\Omega}^{N-m}} \int_{\Omega} \min_{\sigma_i \in M(N,k)} \sum_{i=1, L} [c(x, \tau_i) / w_i + a_i + \gamma_i' \psi_i] \rho(x) dx$$

Сідлова точка функціонала Лагранжа задачі **C-k** на множині $(\Gamma_2^k \times \hat{\Omega}^{N-m}) \times \Psi$ визначається для всіх $i = \overline{1, N}, l = \overline{1, L}$ і майже всіх $x \in \Omega$ так:

$$\hat{\lambda}_i^l(x) = \begin{cases} 1, \text{ якщо } c(x, \hat{\tau}_i) / w_i + a_i + \gamma_i' \hat{\psi}_i \leq c(x, \hat{\tau}_j) / w_j + a_j + \gamma_j' \hat{\psi}_j, \\ \quad \forall i \in \sigma_i, j \in N \setminus \sigma_i, \\ 0 \text{ у протилежному випадку,} \end{cases}$$

де $\hat{\tau}_{m+1}, \dots, \hat{\tau}_N, \hat{\psi}_1, \dots, \hat{\psi}_N$ складають оптимальний розв’язок наступної задачі скінченновимірної умовної оптимізації:

$$G(\psi) = \min_{\tau^{N-m} \in \hat{\Omega}^{N-m}} G_1(\tau, \psi) \rightarrow \max, \quad (16)$$

за умов

$$\psi_i \geq 0, \quad i = \overline{p+1, N}, \quad (17)$$

$$G_1(\tau^{N-m}, \psi) = \int_{\Omega} \min_{\sigma_i \in M(N,k)} \sum_{i=1, L} \left[\frac{c(x, \tau_i)}{w_i} + a_i + \gamma_i' \psi_i \right] \rho(x) dx - \sum_{i=1}^N \psi_i b_i. \quad (18)$$

З усіх розв’язків задачі **C-k** вибираємо ті, що є і розв’язками задачі **B-k**, і приходимо до наступного твердження.

Твердження 1. Оптимальний розв’язок задачі **B-k** має наступний вигляд: для $i = \overline{1, N}, l = \overline{1, L}$ і майже всіх $x \in \Omega$:

$$\lambda_i^l(x) = \begin{cases} 1, \text{ якщо } c(x, \tau_i) / w_i + a_i + \gamma_i' \psi_i \leq c(x, \tau_j) / w_j + a_j + \gamma_j' \psi_j, \\ \quad \forall i \in \sigma_i, j \in N \setminus \sigma_i, \\ 0 \text{ у протилежному випадку,} \end{cases}$$

за $\tau_{m+1}, \dots, \tau_N, \psi_1^*, \dots, \psi_N^*$ вибирається оптимальний розв’язок задачі (16), (17) з негладкою ціловою функцією (18).

Чисельний алгоритм. Від задачі (16) перейдемо до задачі:

$$\max_{\psi \in E^N} \min_{\tau^{N-m} \in \hat{\Omega}^{N-m}} U(\tau^N, \psi),$$

$$\text{де } U(\tau^{N-m}, \psi) = G_1(\tau^{N-m}, \psi) - Q \sum_{i=p+1}^N \max(0, -\psi_i),$$

Q – додатне число (більше за множники Лагранжа). Далі вектор τ^{N-m} будемо позначати τ . Нехай

$$g_U(\tau, \psi) = (g_U^{\tau}(\tau, \psi), -g_U^{\psi}(\tau, \psi)) = (g_U^{\tau_{m+1}}(\tau, \psi), \dots, g_U^{\tau_N}(\tau, \psi), -g_U^{\psi_1}(\tau, \psi), \dots, -g_U^{\psi_N}(\tau, \psi))$$

– вектор узагальненого псевдо-градієнта функції $U(\tau^{N-m}, \psi)$ в точці $(\tau, \psi) = (\tau_{m+1}, \dots, \tau_N, \psi_1, \dots, \psi_N)$. Його компоненти $g_U^{\tau_j}(\tau, \psi), j = m+1, \dots, N$, обчислюємо за допомогою скінченних різниць, оскільки зміна цих координат викликає перерозподіл зон відповідальності не тільки центрів, що розміщуються, але й існуючих. Компоненти $-g_U^{\psi_i}(\tau, \psi), i = 1, \dots, N$, визначаються за формулою:

$$g_U^{\psi_i}(\tau, \psi) =$$

$$= \begin{cases} \int_{\Omega} \sum_{l=1}^L \gamma_l' \lambda_i^l(x) \rho(x) dx - b_i, & i = \overline{1, m}, \\ \int_{\Omega} \sum_{l=1}^L \gamma_l' \lambda_i^l(x) \rho(x) dx - b_i + Q \max(0, \text{sign}(-\psi_i)), & i = \overline{m+1, N}, \end{cases} \quad (19)$$

де $\lambda_i^l(x), i = \overline{1, N}, l = \overline{1, L}$, обчислюються так:

$$\lambda_i^l(x) = \begin{cases} 1, \text{ якщо } c(x, \tau_i) / w_i + a_i + \gamma_i' \psi_i \leq c(x, \tau_j) / w_j + a_j + \gamma_j' \psi_j, \\ \quad \forall i \in \sigma_i, j \in N \setminus \sigma_i, \\ 0 \text{ у протилежному випадку.} \end{cases} \quad (20)$$

Ключовою частиною алгоритму розв’язання задачі **B-k** є $r(\alpha)$ -алгоритм з константним коефіцієнтом розтягу простору α й адаптивним способом регулювання крокового множника (Шор Н. З., 1999, с. 33–47).

Ініціалізація. Область Ω заключаємо в прямокутник Π , $\rho(x) = 0$ для $x \in \Pi \setminus \Omega$. Π покриваємо прямокутною сіткою, задаємо початкове наближення $(\tau, \psi) = (\tau^{(0)}, \psi^{(0)})$, параметри $\alpha, q_1, q_2, n_h, \varepsilon$ модифікації $r(\alpha)$ – алгоритму.

Крок 1.

Обчислюємо $\lambda^{(0)}(x) = (\lambda_1^{(0)}(x), \dots, \lambda_N^{(0)}(x))$, $l = \overline{1, L}$, $G_1(\tau^{(0)}, \psi^{(0)})$ і вектор $g_U(\tau^{(0)}, \psi^{(0)})$ у вузлах сітки при $\tau = \tau^{(0)}, \psi = \psi^{(0)}$.

Вибираємо початковий крок $h_0 > 0$, задаємо $B_0^T = I_{2(N-m)}$, $B_0^V = I_N$, – квадратні матриці $2(N-m) \times 2(N-m)$, $N \times N$ відповідно, і знаходимо:

$$\tau^{(1)} = P_{\Pi} \left(\tau^{(0)} - h_0 g_U^T(\tau^{(0)}, \psi^{(0)}) \right),$$

$$\psi^{(1)} = \psi^{(0)} - h_0 g_U^V(\tau^{(0)}, \psi^{(0)}).$$

Крок 2. Нехай в результаті обчислень $k, k = 1, 2, \dots$ кроків алгоритму отримано $\tau^{(k)}, \psi^{(k)}$, $\lambda^{(k-1)}(x)$, $l = \overline{1, L}$, у вузлах сітки, матриці B_k^T , B_k^V .

Опишемо $(k+1)$ -й крок, який включає наступні етапи.

1. Обчислюємо $\lambda^{(k)}(x)$, $l = \overline{1, L}$, у вузлах сітки за формулою (17) при $\tau = \tau^{(k)}, \psi = \psi^{(k)}$.

2. Обчислюємо $g_U(\tau^{(k)}, \psi^{(k)})$ у вузлах сітки при $\tau = \tau^{(k)}, \psi = \psi^{(k)}$, $\lambda^l(x) = \lambda^{(k)}(x)$, $l = \overline{1, L}$.

3. Проводимо чергову ітерацію $r(\alpha)$ -алгоритму, обчислювальна формула якого має вигляд:

$$\tau^{(k+1)} = P_{\Pi} \left(\tau^{(k)} - h_k B_{k+1}^T g_U^T(\tau^{(k)}, \psi^{(k)}) \right),$$

$$\psi^{(k+1)} = \psi^{(k)} - h_k B_{k+1}^V g_U^V(\tau^{(k)}, \psi^{(k)}).$$

Тут B_{k+1}^T, B_{k+1}^V – оператори відображення перетвореного простору в основний простір з коефіцієнтом розтягу α , які перераховуються

за формулою $B_{k+1}^w = B_k^w \left(I + \left(\frac{1}{\alpha} - 1 \right) \theta_k^w (\theta_k^w)^T \right)$, де

w – змінна τ або ψ , I – одинична матриця відповідного розміру, θ_k^w – нормований вектор різниці двох послідовних псевдоградієнтів в перетвореному просторі, тобто:

$$\theta_k^w = \frac{(B_{k+1}^w)^T (g_U^w(\tau^{(k)}, \psi^{(k)}) - g_U^w(\tau^{(k-1)}, \psi^{(k-1)}))}{(B_{k+1}^w)^T (g_U^w(\tau^{(k)}, \psi^{(k)}) - g_U^w(\tau^{(k-1)}, \psi^{(k-1)}))},$$

за умови $(B_{k+1}^w)^T (g_U^w(\tau^{(k)}, \psi^{(k)}) - g_U^w(\tau^{(k-1)}, \psi^{(k-1)})) \geq \varepsilon_0$, і $\theta_k^w = 0$ в решті випадків, ε_0 – точність представлення машинного нуля. Множник h_k регулюється адаптивним способом з параметрами h_0, q_1, q_2, n_h , сообразно якому множ-

ник обирається з умови мінімуму різниці $[G_1(\tau^{(k-1)}, \psi^{(k)}) - G_1(\tau^{(k)}, \psi^{(k-1)})]$ за напрямом узагальненого антипсевдоградієнта $-g_U(\tau, \psi)$ в перетвореному просторі.

4. Якщо умова $(\tau^{(k+1)}, \psi^{(k+1)}) - (\tau^{(k)}, \psi^{(k)}) \leq \varepsilon, \varepsilon > 0$, не виконується, переходимо до $(k+2)$ -го кроку алгоритму з новими значеннями величин $\tau^{(k+1)}, \psi^{(k+1)}$, $\lambda^{(k)}(x)$, $l = \overline{1, L}$, у вузлах сітки, інакше – переходимо до п. 5.

5. Покладемо $\tau_* = \tau^{(l)}, \psi_* = \psi^{(l)}$, $\lambda^l(x) = \lambda^{(s)}(x)$, $l = \overline{1, L}$, де S – номер ітерації, на якій умова закінчення ітераційного процесу виконалася.

6. Обчислюємо оптимальне значення функціоналу вихідної задачі **B-k** при $\tau = \tau_*$, $\psi = \psi_*$, $\lambda^l(\cdot) = \lambda_*^l(\cdot)$, $l = \overline{1, L}$ і для контролю правильності розрахунків – значення цільової функції задачі (16) з такими самими параметрами.

Кінець алгоритму.

Для негладких функцій параметри доцільно вибирати наступним чином: $\alpha = 2 \div 3$, $h_0 = 1$, $q_1 = 1$, $q_2 = 1.1 \div 1.2$, $n_h = 2 \div 3$, де q_1 коефіцієнт зменшення кроку, якщо умова релаксації ітераційного процесу за поточним напрямом спуску виконується за один крок, q_2 – коефіцієнт збільшення кроку, натуральне число n_h ($n_h > 1$) задає кількість кроків одновимірного спуску, після яких крок збільшуватиметься в q_3 разів.

Експерименти. Наведемо результат розв'язання модельних задач оптимального розміщення нових об'єктів деякої мережі сервісних підприємств з перерозподілом сфер обслуговування існуючих і нових підприємств. Для всіх прикладів: $\Omega = \{x \in R^2 : 0 \leq x_i \leq 10, i = 1, 2\}$; $\rho(x) = 1 \quad \forall x \in \Omega$; $w_i = 1, \forall i = \overline{1, N}$; функція відстані – метрика Мінковського з параметром p : $c(x, \tau_i) = \sqrt[p]{(x_1 - \tau_1^i)^p + (x_2 - \tau_2^i)^p}$.

Задача 1. На рис. 1 представлено зонування території до та після розміщення двох нових центрів (7-го та 8-го) за таких початкових даних:

$$N = 8, m = 6, p = 1, k = 1,$$

$$a = (1; 1; 1.2; 1.04; 0.94; 1; 1; 1.18).$$

Значення функціоналу знизилось на 8,4 % з 331,97 до 304,251 одиниць.

З табл. 1 видно, наскільки зменшилося і майже рівномірно розподілилося навантаження на центри після розміщення нових і перерозподілу зон обслуговування.

Далі розглянемо випадки, коли зони обслуговування можуть перекриватися двократно.

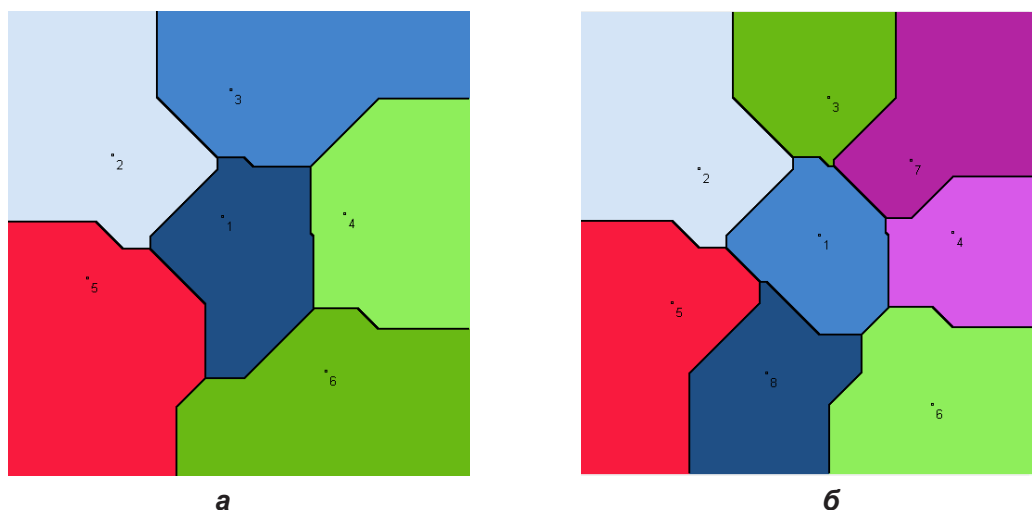


Рис. 1. Монопольні зони обслуговування в задачі 1 для центрів:
а – шести фіксованих, б – шести фіксованих і двох, оптимально розміщених

Таблиця 1
Потужності центрів в задачі 1

№ центра	Потужність центра	
	до перерозподілу	після перерозподілу
1	11,189	9,015
2	17,362	17,362
3	17,259	9,434
4	15,218	9,618
5	20,632	15,457
6	18,334	14,406
7	0	14,043
8	0	10,659

Задача 2. Початкові дані: $N = 9, m = 6, p = 2, k = 2$, параметри $a_i, \forall i = 1, N$, наведено в табл. 3.

На рис. 2, а подано оптимальне розбиття другого порядку, отримане у результаті розв'язання задачі ОМРМ без врахування інтегральних обмежень для фіксованих шести центрів. Результат розв'язання задачі ОМРМ, у якій додатково розміщені три центри – 7й, 8й і 9й, наведено на рис. 2, б.

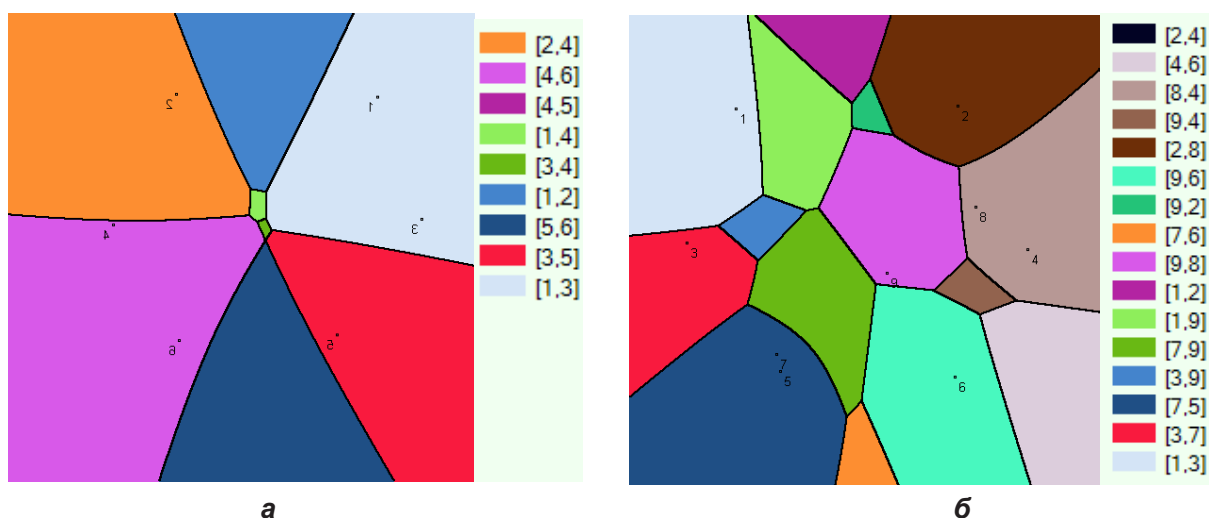


Рис. 2. Дуплексне розбиття множини:
а – для шести фіксованих центрів, б – для шести фіксованих і трьох оптимально розміщених центрів

Таблиця 2
Навантаження на центри в задачі 3

№ центра		Потужність центра	
		до перерозподілу	після перерозподілу
1	2,139	13,3941	10,2494
2	2,346	13,9087	7,5423
3	2,135	17,0553	9,8256
4	1,335	22,2905	9,7908
5	2,369	14,4236	6,7787
6	1,618	18,9269	10,2223
7	1,63	0	13,8029
8	1,53	0	14,995
9	0,906	0	16,792

За рахунок розміщення трьох додаткових центрів були перерозподілені зони обслуговування старих центрів, перераховано їх потужності. Табл. 2 містить інформацію про те, як змінилося навантаження на старі центри, і яку потужність повинні мати нові центри. У табл. 3 наведені ці зони для кожного центру після розміщення нових. Значення цільового функціоналу зменшилося з 873,586 до 755,162.

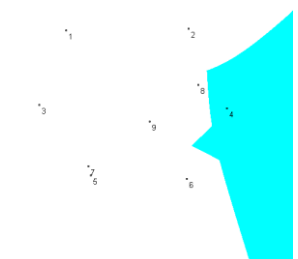
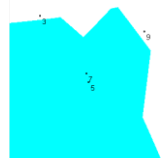
З результатів розв'язання тестових задач видно, що розміщення нових центрів має істот-

ний вплив на потужності найближчих до них центрів. Потужності більш віддалених центрів можуть або зменшуватися незначно, або ж залишатися незмінними.

Висновки. У даній роботі описаний підхід щодо територіальної сегментації обмеженої області, заснований на використанні моделей і методів розв'язання неперервних задач оптимального мультиплексного розбиття множин. Зонування території здійснюється у такий спосіб, що певна зона виявляється не просторовою монополією одного центру, а може обслуговуватися будь-яким з k найближчих або найбільш привабливих центрів. На відміну від досліджених раніше задач ОМРМ тут розглянутий випадок розширення функціонуючої мережі сервісних центрів, коли частина з них вже існує, а декілька нових центрів потрібно розмістити на певній території, в зв'язку з чим одночасно перерозподіляються зони обслуговування для усіх центрів. За допомогою розробленого програмного забезпечення для зонування територій продажу можна швидко оновлювати карти для обліку нових торгових представників та інших незначних змін у бізнесі з мінімальним втручанням.

Таблиця 3

Зони обслуговування після розміщення нових центрів

№	Зона	№	Зона	№	Зона
1		2		3	
4		5		6	
7		8		9	

Зазвичай, планування розміщення нових сервісних центрів передбачає врахування таких факторів, як наявність комунікацій у місцях передбачуваного розміщення центрів, можливі додаткові витрати на будівництво власних комунікацій, податок на землю тощо. При моделюванні реальних задач сегментації виника-

ють також питання, пов'язані з описом кордонів територій, що сегментуються, обчисленням фактичних відстаней між сервісними центрами та їх клієнтами. Тоді поряд з представленим математичним і алгоритмічним забезпеченням мають залучатися сучасні геоінформаційні технології.

ЛІТЕРАТУРА:

1. Коряшкіна Л.С. Розширення одного класу нескінченновимірних оптимізаційних задач. *Вісн. Черкаського ун-ту. Сер. Прикл. матем. Інф.* 2015. № 18(351). С. 28–36.
2. Kiseleva E., Koriashkina L. Theory of Continuous Optimal Set Partitioning Problems as a Universal Mathematical Formalism for Constructing Voronoi Diagrams and Their Generalizations. I. Theoretical Foundations. *Cybernetics and Systems Analysis*. 2015. Vol. 51. Issue 3. P. 325–335.
3. Boots B., South R. Modeling retail trade areas using higher-order, multiplicatively weighted Voronoi diagrams. *Journal of Retailing* 73. 1997. P. 519–536.
4. Koriashkina L.S., Cherevatenko A.P. Continuous problems of optimal multiplex-partitioning of sets without constraints and solving methods. *Journal of Computational & Applied Mathematics*. 2015. № 2 (119). P. 15–32.
5. Sahin G., Sural H. A review of hierarchical facility location models. *Comput. Oper. Res.* 2007. Vol. 34. P. 2310–2331.
6. Mallozzi L., Puerto J., Rodríguez-Madrena M. On Location-Allocation Problems for Dimensional Facilities. *Journal of Optimization Theory and Applications*. 2018. P. 1–38.
7. Torkestani S.S., Seyedhosseini S.M., Makui A. Hierarchical Facility Location and Hub Network Problems: A literature review. *J. Ind. Syst. Eng.* 2016. Vol. 9. P. 1–22.
8. Lin Y., Jia H., Yang Y., Tian G., Tao F., Ling L. An improved artificial bee colony for facility location allocation problem of end-of-life vehicles recovery network. *J. Clean. Prod.* 2018. Vol. 205. P. 134–144.
9. Chen L., Zhang W., Ma J., Wang L. Newly-added Station Location-allocation in One-way Station-based Carsharing Systems. In *Proceedings of the 2019 Chinese Control Conference (CCC), Guangzhou, China, 27–30 July 2019*. Vol. 2019.
10. Wang X., Lim M., Ouyang Y. A Continuum Approximation Approach to the Dynamic Facility Location Problem in a Growing Market. *Transportation Science*. 2018. 51:1. P. 343–357.
11. Lage M.d.O., Machado C.A.S., Monteiro C.M., Davis C.A. Jr., Yamamura C.L.K., Berssaneti F.T., Quintanilha J.A. Using Hierarchical Facility Location, Single Facility Approach, and GIS in Carsharing Services. *Sustainability*. 2021. 13, 12704. <https://doi.org/10.3390/su132212704>
12. Murray A.T., Church R.L., Feng X. Single facility siting involving allocation decisions. *Eur. J. Oper. Res.* 2020. Vol. 284. P. 834–846.
13. Monteiro C.M., Machado C.A.S., Lage M.O., Berssaneti F.T., Davis C.A., Quintanilha J.A. Optimization of carsharing fleet size to maximize the number of clients served. *Comput. Environ. Urban Syst.* 2021. Vol. 87. 101623.
14. Murray A.T., Xu J., Wang Z., Church R.L. Commercial GIS location analytics: Capabilities and performance. *Int. J. Geogr. Inf. Sci.* 2019. Vol. 33. P. 1106–1130.
15. Itai Feigenbaum, Jay Sethuraman, Chun Ye. Approximately optimal mechanisms for strategyproof facility location: Minimizing Lp norm of costs. *Mathematics of Operations Research*. 2017. Vol. 42 (2). P. 434–447.
16. Karatas M. A multi-objective facility location problem in the presence of variable gradual coverage performance and cooperative cover. *Eur. J. Oper. Res.* 2017. V. 262. P. 1040–1051.
17. Karatas M., Dasci A. A two-level facility location and sizing problem for maximal coverage. *Comput. Ind. Eng.* 2020. Vol. 139. 106204.
18. Rodriguez S.A., De la Fuente R.A., Aguayo M.M. A facility location and equipment emplacement technique model with expected coverage for the location of fire stations in the Concepción province, Chile. *Comput. Ind. Eng.* 2020. Vol. 147, 106522.
19. Zheng C., Zhao X., Shen J. Research on Location Optimization of Metro-Based Underground Logistics System with Voronoi diagram. *IEEE Access* 2020. Vol. 8. P. 34407–34417.
20. Kiseleva E., Koriashkina L. The Theory of Continuous Optimal Set Partitioning Problems as a Universal Mathematical Formalism for Constructing the Voronoi Diagram and its Generalizations. II. Algorithms for constructing Voronoi Diagrams based on the theory of optimal partitioning of sets. *Cybernetics and Systems Analysis*. 2015. Vol. 51, Issue 4. P. 3–12.

21. Gamsa A., Lee D., Liu C.H., Wagner D. Higher Order City Voronoi Diagrams. *Proc. 13th Scandinavian Symposium and Workshops on Algorithm Theory (SWAT'12). Helsinki, 2012. (Lecture Notes in Computer Science; Vol. 7357).* P. 59–70.
22. Коряшкіна Л.С., Череватенко А.П. Оптимальные диаграммы Вороного высших порядков ограниченных множеств и алгоритмы их построения. *Вісник Черкаського університету. Серія Прикладна математика. Інформатика.* 2016. №. 1-2. С. 59–76.
23. Mallozzi L., Justo P. The geometry of optimal partitions in location problems. *Optimization Letters* 12. 2018. P. 203–220.
24. Lynskey J., Thar K., Oo T.Z., Hong C.S. Facility Location Problem Approach for Distributed Drones. *Symmetry*. 2019. 11. 118. <https://doi.org/10.3390/sym11010118>
25. Wolf Gert W. Solving location-allocation problems with professional optimization software. *Transactions in GIS* 26. 2022. P. 2741–2775.
26. Alghanmi Nusaybah, Alotaibi Reem, Alshammari Sultanah, Alhothali Areej, Bamasag Omaimah, Faisal Kamil. A Survey of Location-Allocation of Points of Dispensing During Public Health Emergencies *Front. Public Health, 10 March 2022 Sec. Disaster and Emergency Medicine.* 2022. Vol. 10. <https://doi.org/10.3389/fpubh.2022.811858>
27. Lebedeva O., Kripak M., Gozbenko V. Increasing effectiveness of the transportation network by using the automation of a Voronoi diagram. *Transp. Res. Procedia*. 2018. Vol. 36. P. 427–433.
28. Коряшкіна Л.С., Череватенко А.П., Коряшкіна Е.О. Інтеграція ГІС-технологій і методів розв'язання неперервних задач оптимального мультиплексного розбиття множин. *Системні дослідження та інформаційні технології.* 2017. № 4. С. 97–108.
29. Koriashkina L., Cherevatenko A., Mykhalova O. The continuous problems of the optimal multiplex partitioning an application of sets. *Power Engineering and Information Technologies in Technical Objects Control. Annual Proceedings.* 2016. P. 233–240.
30. Шор Н. З., Бардадым Т. А., Журбенко Н. Г., Лиховид А. П., Стецюк П. И. Использование методов негладкой оптимизации в задачах стохастического программирования. *Кибернетика и системный анализ.* 1999. № 5. С. 33–47.

REFERENCES:

1. Koriashkina L.S. (2015). Rozshyrennia odnogo klasu neskinchenovymirnykh optymizatsiynykh zadach. *Visn. Cherkaskoho un-tu. Ser. Prykl. matem. Inf.* № 18(351). S. 28–36 [in Ukrainian].
2. Kiseleva E., Koriashkina L. Theory of Continuous Optimal Set Partitioning Problems as a Universal Mathematical Formalism for Constructing Voronoi Diagrams and Their Generalizations. I. Theoretical Foundations. *Cybernetics and Systems Analysis.* 2015. Vol. 51, Issue 3. P.325–335.
3. Boots B., South R. Modeling retail trade areas using higher-order, multiplicatively weighted Voronoi diagrams. *Journal of Retailing* 73. 1997. P. 519–536.
4. Koriashkina L.S., Cherevatenko A.P. Continuous problems of optimal multiplex-partitioning of sets without constraints and solving methods. *Journal of Computational & Applied Mathematics.* 2015. № 2 (119). P. 15–32
5. Sahin G., Sural H. A review of hierarchical facility location models. *Comput. Oper. Res.* 2007. Vol. 34. P. 2310–2331.
6. Mallozzi L., Puerto J., Rodríguez-Madrena M. On Location-Allocation Problems for Dimensional Facilities. *Journal of Optimization Theory and Applications.* 2018. P. 1–38.
7. Torkestani S.S., Seyedhosseini S.M., Makui A. Hierarchical Facility Location and Hub Network Problems: A literature review. *J. Ind. Syst. Eng.* 2016. Vol. 9, P. 1–22.
8. Lin Y., Jia H., Yang Y., Tian G., Tao F., Ling L. An improved artificial bee colony for facility location allocation problem of end-of-life vehicles recovery network. *J. Clean. Prod.* 2018. Vol. 205. P. 134–144.
9. Chen L., Zhang W., Ma J., Wang L. Newly-added Station Location-allocation in One-way Station-based Carsharing Systems. *In Proceedings of the 2019 Chinese Control Conference (CCC), Guangzhou, China, 27–30 July 2019.* Vol. 2019.
10. Wang X, Lim M, Ouyang Y. A Continuum Approximation Approach to the Dynamic Facility Location Problem in a Growing Market. *Transportation Science.* 2018. 51:1. P. 343–357.
11. Lage M.d.O., Machado C.A.S., Monteiro C.M., Davis C.A. Jr., Yamamura C.L.K., Berssaneti F.T., Quintanilha J.A. Using Hierarchical Facility Location, Single Facility Approach, and GIS in Carsharing Services. *Sustainability.* 2021. 13, 12704. <https://doi.org/10.3390/su132212704>
12. Murray A.T., Church R.L., Feng X. Single facility siting involving allocation decisions. *Eur. J. Oper. Res.* 2020. Vol. 284. P. 834–846.

13. Monteiro C.M., Machado C.A.S., Lage M.O., Berssaneti F.T., Davis C.A., Quintanilha J.A. Optimization of carsharing fleet size to maximize the number of clients served. *Comput. Environ. Urban Syst.* 2021. Vol. 87. 101623.
14. Murray A.T., Xu J., Wang Z., Church R.L. Commercial GIS location analytics: Capabilities and performance. *Int. J. Geogr. Inf. Sci.* 2019. Vol. 33. P. 1106–1130.
15. Itai Feigenbaum, Jay Sethuraman, Chun Ye. Approximately optimal mechanisms for strategyproof facility location: Minimizing Lp norm of costs. *Mathematics of Operations Research.* 2017. Vol. 42 (2). P. 434–447,
16. Karatas M. A multi-objective facility location problem in the presence of variable gradual coverage performance and cooperative cover. *Eur. J. Oper. Res.* 2017. V. 262. P. 1040–1051.
17. Karatas M., Dasci A. A two-level facility location and sizing problem for maximal coverage. *Comput. Ind. Eng.* 2020. Vol. 139. 106204.
18. Rodriguez S.A., De la Fuente R.A., Aguayo M.M. A facility location and equipment emplacement technique model with expected coverage for the location of fire stations in the Concepción province, Chile. *Comput. Ind. Eng.* 2020. Vol. 147, 106522.
19. Zheng C., Zhao X., Shen J. Research on Location Optimization of Metro-Based Underground Logistics System with Voronoi diagram. *IEEE Access* 2020. Vol. 8. P. 34407–34417.
20. Kiseleva E., Koriashkina L. The Theory of Continuous Optimal Set Partitioning Problems as a Universal Mathematical Formalism for Constructing the Voronoi Diagram and its Generalizations. II. Algorithms for constructing Voronoi Diagrams based on the theory of optimal partitioning of sets. *Cybernetics and Systems Analysis.* 2015. Vol. 51, Issue 4. P. 3–12.
21. Gemsa A., Lee D., Liu C.H., Wagner D. Higher Order City Voronoi Diagrams. *Proc. 13th Scandinavian Symposium and Workshops on Algorithm Theory (SWAT'12). Helsinki, 2012. (Lecture Notes in Computer Science; Vol. 7357).* P. 59–70.
22. Korjashkina L.S., Cherevatenko A.P. (2016). Optimal'nye diagrammy Voronogo vysshih porjadkov ograničennyh mnozhestv i algoritmy ih postroenija. *Visnik Cherkas'kogo universitetu. Serija Prikladna matematika. Informatika.* №. 1-2. S. 59–76 [in Russian].
23. Mallozzi L., Justo P. The geometry of optimal partitions in location problems. *Optimization Letters* 12. 2018. P. 203–220.
24. Lynskey J., Thar K., Oo T.Z., Hong C.S. Facility Location Problem Approach for Distributed Drones. *Symmetry.* 2019. 11. 118. <https://doi.org/10.3390/sym11010118>
25. Wolf Gert W. Solving location-allocation problems with professional optimization software. *Transactions in GIS* 26. 2022. P. 2741–2775.
26. Alghanmi Nusaybah, Alotaibi Reem, Alshammari Sultanah, Alhothali Areej, Bamasag Omaimah, Faisal Kamil. A Survey of Location-Allocation of Points of Dispensing During Public Health Emergencies *Front. Public Health, 10 March 2022 Sec. Disaster and Emergency Medicine.* 2022. Vol. 10. <https://doi.org/10.3389/fpubh.2022.811858>
27. Lebedeva O., Kripak M., Gozbenko V. Increasing effectiveness of the transportation network by using the automation of a Voronoi diagram. *Transp. Res. Procedia.* 2018. Vol. 36. P. 427–433.
28. Koriashkina L.S., Cherevatenko A.P., Koriashkina E.O. Intehratsiia HIS-tekhnohohii i metodiv rozviazannia neperervnykh zadach optymalnoho multypleksnoho rozbytta mnozhyn. *Systemni doslidzhennia ta informatsiini tekhnolohii.* 2017. № 4. S. 97–108 [in Ukrainian].
29. Koriashkina L., Cherevatenko A., Mykhalova O. The continuous problems of the optimal multiplex partitioning an application of sets. *Power Engineering and Information Technologies in Technical Objects Control. Annual Proceedings.* 2016. P. 233–240.
30. Shor N. Z., Bardadym T. A., Zhurbenko N. G., Lihovid A. P., Stecjuk P. I. Ispol'zovanie metodov negladkoj optimizacii v zadachah stohasticheskogo programmirovaniia. *Kibernetika i sistemnyj analiz.* 1999. № 5. S. 33–47 [in Russian].

УДК 04.5

DOI <https://doi.org/10.32782/IT/2023-2-3>

Володимир КУВАЄВ

доктор технічних наук, професор кафедри програмного забезпечення комп'ютерних систем, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, м. Дніпро, Україна, 49005, kuvaiev.v.m@ntu.one

ORCID: 0000-0001-6329-071X

Scopus Author ID: 6602411915

Леонід МЕЩЕРЯКОВ

доктор технічних наук, професор кафедри програмного забезпечення комп'ютерних систем, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, м. Дніпро, Україна, 49005, meshcheriakov.l.i@ntu.one

ORCID: 0000-0002-9579-1970

Scopus Author ID: 57205282540

Альона ХАРЬ

асистент кафедри програмного забезпечення комп'ютерних систем, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, м. Дніпро, Україна, 49005, khar.a.t@ntu.one

ORCID: 0000-0003-3176-7792

Артем ПОЛІТОВ

студент факультету фізики, електроніки та комп'ютерних систем, Дніпровський національний університет імені Олеся Гончара, просп. Гагаріна, 72, м. Дніпро, Україна, 49000, politov3601@gmail.com

ORCID: 0000-0003-2822-8026

Бібліографічний опис статті: Куваєв, В., Мещеряков, Л., Харь, А., Політов, А. (2023). Інтерфейс технічного супроводження складних інформаційно-керуючих систем автоматизації. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 26–33, doi: <https://doi.org/10.32782/IT/2023-2-3>

ІНТЕРФЕЙС ТЕХНІЧНОГО СУПРОВОДЖЕННЯ СКЛАДНИХ ІНФОРМАЦІЙНО-КЕРУЮЧИХ СИСТЕМ АВТОМАТИЗАЦІЇ

Розвиток комп'ютерних технологій у сфері автоматизації технологічних процесів сформував достатньо типові підходи до структури технічних засобів і програмного забезпечення комп'ютерних систем керування. Найбільш поширеним є підхід, що передбачає строгу функціональну ієрархію і відповідну до неї технічних засобів та спеціалізованих інструментальних засобів для розробки прикладного програмного забезпечення для кожного рівня ієрархії. В усіх сучасних складних автоматизованих системах керування передбачається реалізація людино-машинного інтерфейсу (HMI – Human-Machine Interface), який, за визначенням, під «машиною» розуміє систему з усіма технічними засобами, що беруть участь у процесі вимірювання, контролю, сигналізації та керуванні, а під «людиною» – оператора-технолога, який бере безпосередню участь у процесі керування. Власне кажучи, безпосередньо HMI програмується засобами SCADA-системи (Supervisory Control And Data Acquisition – диспетчерське управління та збір даних). Ринок пропонує велику кількість SCADA-систем різних розробників з схожим функціоналом підсистем для оператора-технолога. У той же час, вони залишають по за уваги питання технічного супроводження систем автоматизації побудованих за допомогою засобів SCADA, що ускладнює пошуки порушень у зовнішніх ланцюгах систем автоматизації, уповільнює налаштування системи після ремонтів діючих складних систем автоматизації, які вже введені в експлуатацію і працюють. Питання функціоналу та досвід створення інтерфейсу технічного супроводження складних систем автоматизації, що аналізуються на прикладі системи керування швидкісним режимом прокатки безперервного прокатного стану ПАО «АрселорМіттал Кривий Ріг», можуть бути використані для подальшого розвитку SCADA-систем.

Метою роботи є обґрунтування рішень що до структури і функціоналу інтерфейсу технічного супроводження складних систем автоматизації. Реалізація поставленої мети передбачає обґрунтування структури інтерфейсу з урахуванням функціональних аспектів і форм подання інформації щодо різних параметрів, які дозволяють контролювати зовнішні ланцюги підключення системи автоматизації та внутрішні параметри, які настроюються.

Методологія вирішення поставленого завдання полягає у аналізі і структуруванні масиву інформації, що використовує персонал, який супроводжує експлуатацію системи, та подальшому її групуванню за функціональними ознаками.

Наукова новизна полягає у обґрунтуванні доцільності і необхідності доповнення традиційного людино-машинного інтерфейсу SCADA-систем функціоналом, що забезпечує технічне супроводження систем автоматизації.

Ключові слова: SCADA-система, людино-машинний інтерфейс, технічне супроводження систем автоматизації.

Volodymyr KUVAIEV

Doctor of Technical Sciences, Professor at the Department of Software Engineering, Dnipro University of Technology, Dmytro Yavornytskyi ave., 19, Dnipro, Ukraine, 49005, kuvaiev.v.m@nmu.one

ORCID: 0000-0001-6329-071X

Scopus Author ID: 6602411915

Leonid MESHCHERIAKOV

Doctor of Technical Sciences, Professor at the Department of Software Engineering, Dnipro University of Technology, Dmytro Yavornytskyi ave., 19, Dnipro, Ukraine, 49005, meshcheriakov.l.i@nmu.one

ORCID: 0000-0002-9579-1970

Scopus Author ID: 57205282540

Alona KHAR

Assistant Lecturer at the Department of Software Engineering, Dnipro University of Technology, Dmytro Yavornytskyi ave., 19, Dnipro, Ukraine, 49005, khar.a.t@nmu.one

ORCID: 0000-0003-3176-7792

Artem POLITOV

Student at the Faculty of Physics, Electronics and Computer Systems, Oles Honchar Dnipro National University, Gagarina ave., 72, Dnipro, Ukraine, 49000, politov3601@gmail.com

ORCID: 0000-0003-2822-8026

To cite this article: Kuvaev, V., Meshcheriakov, L., Khar, A., Politov, A. (2023). Interfeis tekhnichnoho suprovodzhennia skladnykh informatsiino-keruiuchykh system avtomatyzatsii [Interface for technical support of complex automation information control systems]. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 26–33, doi: <https://doi.org/10.32782/IT/2023-2-3>

INTERFACE FOR TECHNICAL SUPPORT OF COMPLEX AUTOMATION INFORMATION CONTROL SYSTEMS

The development of computer technologies in the field of technological processes automation has formed typical approaches to the structure of technical means and software of computer control systems. The most common is the approach that involves a strict functional hierarchy and corresponding technical means and specialized tools for developing application software for each level of the hierarchy. In all modern complex automated control systems, the implementation of a human-machine interface (HMI – Human-Machine Interface) is assumed, that understands the “machine” as a system with all technical means involved in the process of measurement, control, signaling and management, and the “person” as the operator-technologist, which is directly involved in the management process. Strictly speaking, the NMI is directly programmed by means of the SCADA system (Supervisory Control and Data Acquisition). The market offers many SCADA systems from different developers with similar functionality of subsystems for the operator-technologist. At the same time, they neglect the issue of technical support of automation systems built with the help of SCADA tools, which complicates the search for violations in the external circuits of automation systems, slows down the system adjustment after repairs of existing complex automation systems that have already been put into operation and are working. Functional issues and the experience of creating an interface for technical support of complex automation systems, which are analyzed using the example of control systems for the high-speed rolling mode of a continuous rolling mill of ArcelorMittal Kryvyi Rih PJSC, can be used for the further development of SCADA systems.

The purpose of the work is to justify decisions regarding the structure and functionality of the technical support interface of complex automation systems. The realization of the set goal involves substantiating the structure of the interface, considering functional aspects and forms of presenting information about various parameters that allow controlling the external connection chains of the automation system and the internal parameters that can be adjusted.

The methodology for solving the given task consists in the analysis and structuring of the information mass used by the personnel accompanying the operation of the system, and its subsequent grouping according to functional characteristics.

The scientific novelty consists in substantiating the expediency and necessity of supplementing the traditional human-machine interface of SCADA systems with functionality that provides technical support for automation systems.

Key words: SCADA system, human-machine interface, technical support of automation systems.

Актуальність проблеми. Розвиток комп'ютерних технологій у сфері автоматизації технологічних процесів сформував достатньо типові підходи до структури технічних засобів і програмного забезпечення комп'ютерних систем керування. Найбільш поширеним є підхід, що передбачає строгу функціональну ієрархію і відповідну до неї технічних засобів та спеціалізованих інструментальних засобів для розробки прикладного програмного забезпечення для кожного рівня ієрархії. В усіх сучасних складних автоматизованих системах керування передбачається реалізація людино-машинного інтерфейсу (HMI – Human-Machine Interface), який, за визначенням, під «машиною» розуміє систему з усіма технічними засобами, що беруть участь у процесі вимірювання, контролю, сигналізації та керуванні, а під «людиною» – оператора-технолог, який бере безпосередню участь у процесі керування. Власне кажучи, безпосередньо HMI програмується засобами SCADA-системи (Supervisory Control And Data Acquisition – диспетчерське управління та збір даних). Ринок пропонує велику кількість SCADA-систем різних розробників з схожим функціоналом підсистем для оператора-технолога. У той же час, вони залишають поза уваги питання технічного супроводження систем автоматизації побудованих за допомогою засобів SCADA, що ускладнює пошуки порушень у зовнішніх ланцюгах систем автоматизації, уповільнює налаштування системи після ремонтів систем автоматизації, які вже введені в експлуатацію і працюють. Таким чином питання функціоналу та досвід створення інтерфейсу технічного супроводження складних систем автоматизації є актуальними.

Аналіз останніх досліджень і публікацій. На поточний час найбільш широке застосування при розробці людино-машинних інтерфейсів систем технологічної автоматизації знайшли програмні засоби SCADA (Пупена, 2020; Bailey, 2005; activolt.com.ua; cora-data.com.ua). Таки системи надають широкий спектр інструментарію для програмування інтерфейсу операторів технологічних процесів, який спрямований на створення достатньо типових підсистем автоматизації.

Більшість програм SCADA/HMI має типовий набір функціональних можливостей для реалізації завдань керування технологічних процесів, доступ до яких оператори-технологи та інший виробничий персонал має доступ через АРМи (автоматизоване робоче місце):

– збирання інформації про контрольовані технологічні параметри (даних реального часу) з контролерів та засобів віддаленого введення/виведення;

– графічне представлення стану технологічного процесу і устаткування в зручній для сприйняття формі у вигляді мнемосхем;

– вторинне оброблення інформації (масштабування, обмеження вводу, перевірка коректності тощо);

– приймання команд оператора і передача їх на контролер, або засіб віддаленого виведення;

– збереження даних реального часу в архівах даних і графічне представлення історичної інформації в зручній для сприйняття формі у вигляді графіків, гістограм тощо;

– сповіщення експлуатаційного і обслуговуючого персоналу про виявлені аварійні події в технологічному процесі і програмно-апаратних засобах;

– фіксація в електронних журналах виникнення аварійних подій у контрольованому технологічному процесі та дій експлуатаційного персоналу;

– формування звітів на основі архівної інформації, тривог та даних реального часу;

– обмін інформацією з автоматизованими системами керування виробництвом та підприємством у складі інтегрованих систем керування;

– виконання певної логіки оброблення даних з використанням вбудованих мов, наприклад, VBA;

– захист від несанкціонованого доступу до компонентів і файлів;

– формування, записування, читання та збереження даних у вигляді рецептів;

– створення оперативних календарних планів керування устаткуванням.

Цей перелік і особливості реалізації наведеного функціоналу у SCADA різних виробників наведено в (Пупена, 2020). Як можна побачити, сучасні SCADA не передбачають реалізацію функціоналу щодо технічного супроводження

складних систем автоматизації та, відповідно, АРМів для персоналу, який обслуговує такі системи.

Слід зазначити ще один аспект з яким будуть стискатися при організації АРМів для персоналу, який обслуговує системи автоматизації – перевірку ланцюгів підключення датчиків, виконавчих механізмів тощо частіше зручно проводити з точок розміщення контролерів, які входять до системи автоматизації, а от дії з налаштування параметрів системи та окремих її елементів – з місць операторів-технологів або місць де є візуальний контроль за датчиками чи механізмами, які налаштовуються.

Більш гнучкі можливості для реалізації задач різного функціонального призначення, зокрема у організації комп'ютерних вузлів змішаного типу – які одночасно вирішують завдання, що притаманні контролерам, і забезпечують інтерфейси різних типів, надає використання операційних системи реального часу (Зайцев, 2019).

Так операційна система реального часу QNX була використана для реалізації програмних рішень по забезпеченню надійного функціонування складних інформаційно-керуючих систем критичних до режиму реального часу як на обчислювальному вузлі який виконував функції контролера, так і вузлі через який організовано інтерфейс оператора та вирішувалися задачі критичні до реального часу (Куваєв, 2021). Сумісна робота задач, критичних до режиму реально часу, і задач «м'якого» реального часу, була досягнута завдяки використанню для програмування інтерфейсу графічної оболонки Photon, що працює під керівництвом операційної системи реального часу QNX (rts.ua).

Слід зазначити, що якщо засоби SCADA надають інструментарій орієнтований на реалізацію HMI оператора-технолога, тоді як Photon надають інструментарій до реалізації інтерфейсу оператора без прив'язки до його функціоналу, а рішення до його структури, інформаційної насиченості, форми надання інформації тощо лежать на розробнику. Подібні інструментарії, що використовуються для створення людино-машинного інтерфейсу описані у (Доценко, 2022), але рекомендації що до його компонування відсутні. В роботі (Куваєв, 2023) описаний HMI оператора-технолога, що було виконано з використанням графічної оболонки Photon, у той же час питання розробки інтерфейсу технічного супроводження складних систем автоматизації в неї не розглядалися.

Метою статті є обґрунтуванні програмних рішень що до структури і функціоналу інтерфейсу технічного і програмного супроводження

складних систем автоматизації. Реалізація поставленої мети передбачає обґрунтування структури інтерфейсу з урахуванням функціональних аспектів і форм подання інформації щодо різних параметрів, які дозволяють контролювати зовнішні ланцюги підключення системи автоматизації, внутрішні параметри, які налаштовуються.

Виклад основного матеріалу. Використання комп'ютерних технологій у системах, що призначені для керування технологічними процесами, дозволяють інтегрувати у них цілу низку функцій, що раніше були притаманні різним системам автоматизації, та отримати синергетичний ефект від такого об'єднання. У той же час суттєво зросла складність таких систем і, відповідно, вимоги до організації у них людино-машинного інтерфейсу. Перш за все, це торкнулося операторів технологічних процесів, що є достатньо логічним, оскільки від якості налаштування і контролю технологічного процесу залежить його продуктивність, якість продукції і своєчасність виявлення порушень як режимів роботи обладнання, так і його стану. Тем не менш, як питання скорочення часу локалізації і усунення таких порушень, так і питання налаштування параметрів елементів апаратно-програмного забезпечення таких систем, наприклад, після їх заміни у ремонті, залишилися поза увагою розробників програмного інструментарію HMI SCADA-систем, які і користувачів цього продукту.

Проблема створення зручного людино-машинного інтерфейсу для налагодження і супроводження технічного забезпечення складної комп'ютерної системи, критичної до режиму реального часу, вирішувалася при впровадженні і супроводженні комп'ютерної системи керування швидкісним режимом прокатки сортової лінії дрібно-сортно-дротового стану ДСДС 250/150 ПАТ «АрселорМіттал Кривий Ріг».

З функціональною структурою і основними програмними рішеннями цієї системи можна ознайомитися в (Куваєв, 2021). Людино-машинний інтерфейс для супроводження технічного забезпечення комп'ютерної системи керування швидкісним режимом прокатки створювався як окрема підсистема, доступ до якої здійснювався з HMI оператора технолога, опис якої наведено в (Куваєв, 2023).

Загальний вигляд людино-машинного інтерфейсу для супроводження технічного забезпечення комп'ютерної системи керування швидкісним режимом прокатки сортової лінії клітей стану ДСДС 250/150 ПАТ «АрселорМіттал Кривий Ріг» наведено на рис. 1, а чистових груп клітей дрібно-сортного стану ДСС 250-2 того ж підприємства – на рис. 2.

Структури меню людино-машинних інтерфейсів для супроводження технічного забезпечення цих систем ідентичні. При їх розробці були використані підходи, що аналогічні описаним у (Куваєв, 2023) – вони виконані у вигляді дворівневого блокнути у якому права сторінка є ведучою і вміщає перелік модулів з прив'язкою до обчислювальних вузлів комп'ютерної системи, в яких вони встановлені, та перелік параметрів, що доступні до налаштування. Ліва сторінка є підлеглою вибраному пункту меню на правій сторінці, відображає інформацію і вміщає елементи що потрібні для встановлення відповідних параметрів системи або перевірки її зовнішніх ланцюгів підключення датчиків, органів керування, виконуючих механізмів тощо.

Так, рис.1 ілюструє елементи меню, що використовуються для налаштування параметрів

встановлення датчиків прогину прокату, який використовується для контролю відхилення прокату від осі прокатки і стабілізації цього відхилення в положення, що задано оператором.

Для коректної роботи системи у кожному міжклітьовому проміжку, де прокатка ведеться з автоматичною стабілізацією прогину прокату, треба знати положення осі прокатки, тобто положення прокату що відповідає «нульовому» прогину, і положення дна передавального жолоба у перетині контролю положення прокату датчиком прогину – максимально можливого прогину прокату у міжклітьовому проміжку. Датчики встановлюються безпосередньо поблизу руху прокату, працюють у важких умовах і захищені спеціальними кожухами. Вони періодично знімаються для профілактики при проведенні штатних ремонтних робіт на прокатному стані.

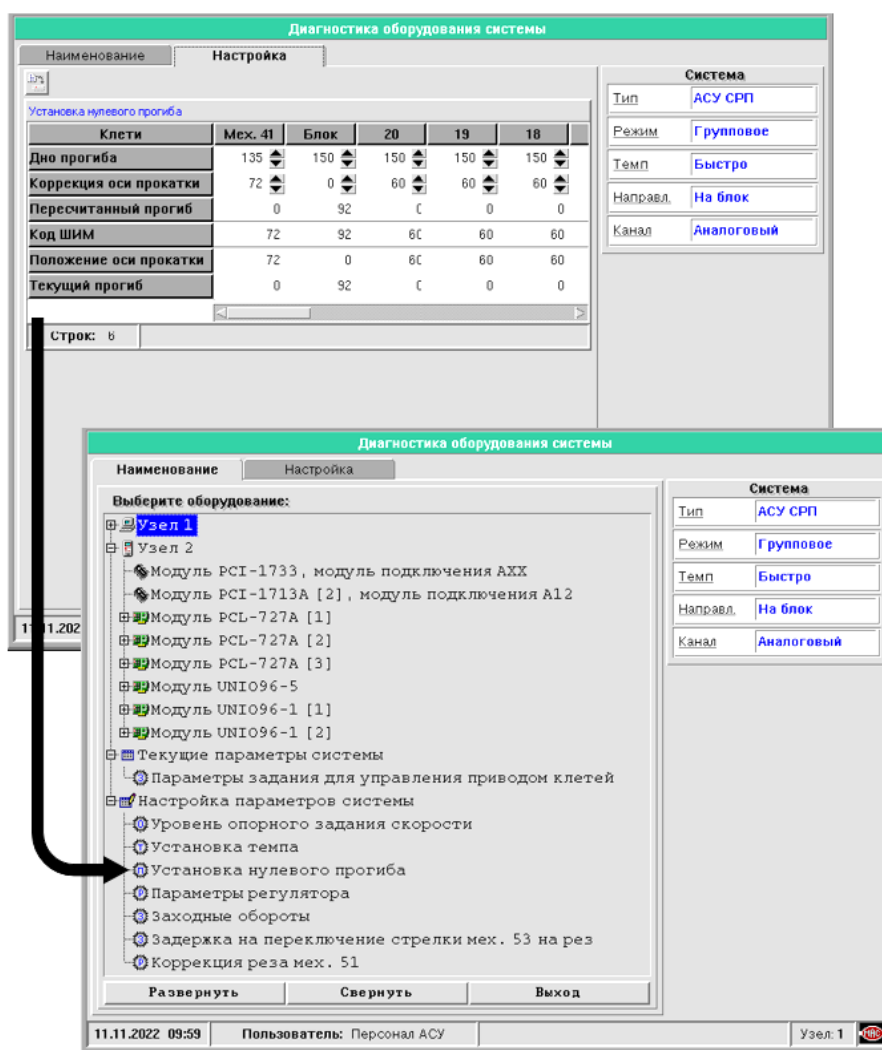


Рис. 1. Людино-машинний інтерфейс для супроводження технічного забезпечення комп'ютерної системи керування швидкісним режимом прокатки сортової лінії дрібносоротно-дротового стану ДСДС 250/150 ПАТ «АрселорМіттал Кривий Ріг». Налаштування параметрів встановлення датчиків прогину прокату

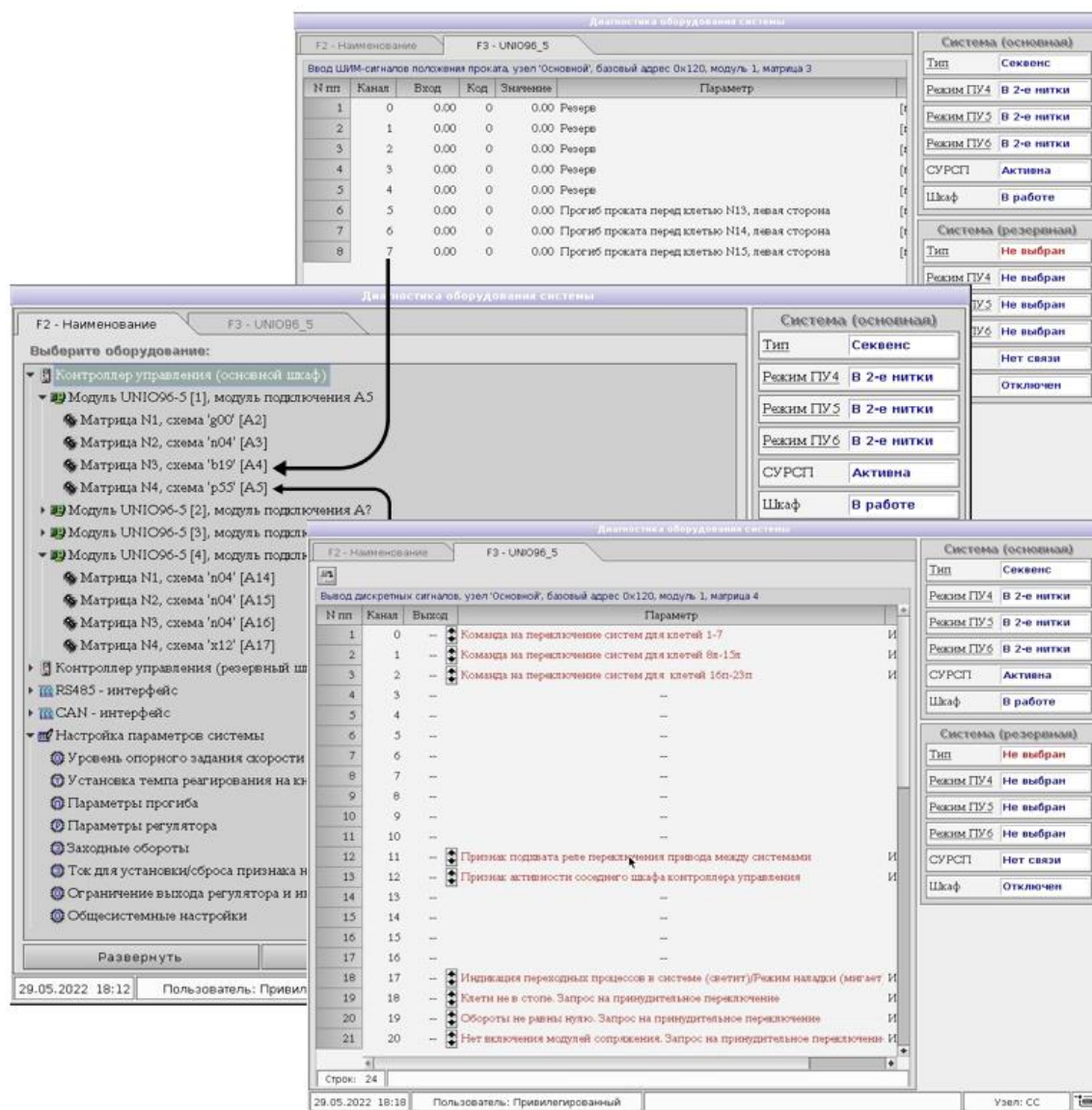


Рис. 2. Людино-машинный интерфейс для супроводження технічного забезпечення комп'ютерної системи керування швидкісним режимом прокатки чистових груп клітей дрібносортового стана ДСС 250-2 ПАТ «АрселорМіттал Кривий Ріг». Контроль вхідних ШІМ сигналів від датчиків прогину прокату і імітація вихідних дискретних сигналів

Наявність сервісу по встановленню «нульового прогину» у людино-машинному інтерфейсу технічного супроводження системи спрощує процедуру налаштування положення датчика прогину при його встановленні на штатне місце. Так відкривши пункт меню «Встановлення нульового прогину» і перейшовши на сторінку «Налаштування» можна безпосередньо контролювати код широтно-імпульсному модульованому (ШІМ) сигналу, який надходить від датчику прогину прокату і який пропорційний величині прогину прокату. Встановивши на дно передавального жолоба імітатор гарячого прокату, корегують

кут нахилу датчика у вертикальній площині таким чином, щоб код ШІМ, що надходить з датчика, був близьким до максимально можливого значення. Після завершення регулювання кута нахилу датчика, він фіксується, а поточний код ШІМ заносять у стовбець таблиці, який відповідає датчику, що налаштовується, у рядок «Дно прогину». Потім встановлюють імітатор у положення осі прокатки і код ШІМ датчика прогину заносять у той ж стовбець, рядок «Корекція осі прокатки». Таким чином, наявність цього сервісу значно спрощує і скорочує час налаштування параметрів встановлення датчиків прогину.

За допомогою інтерфейсу технічного супроводження системи автоматизації можуть бути безпосередньо проконтрольовані вхідні сигнали, що надходять у систему. Так, наприклад, сигнали ШІМ з датчиків прогину прокату можна також проконтролювати через меню технічного супроводження системи до входів модулів вводу відповідних сигналів. На рис. 2 наведені відеограми входів матриці № 3 модуля UNIO95-5, що запрограмована на введення ШІМ-модульованих сигналів. Доступ до цієї інформації здійснюється вибором на сторінці «Найменування» пункту «Пристрій-матриця» відповідного елемента з наступним переходом на підлеглу сторінку.

Інша, матриця № 4 модуля UNIO95-5, запрограмована на вивід дискретних сигналів. Оскільки операція ручного виводу сигналів може бути небезпечною, то вона є доступною (активною) тільки в режимі «теплого» резерву роботи пристрою при перевірці вихідних ланцюгів, а назви параметрів надані червоним кольором.

Запропонований комплекс рішень по людино-машинному інтерфейсу технічного супроводження складних інформаційно-керуючих систем автоматизації дозволив скоротити час і спростити обслуговування системи автоматизації як під час вводу в експлуатацію, так і в процесі її подальшого супроводження.

Висновки. При розробці складних інформаційно-керуючих систем автоматизації слід перед-

бачати крім традиційного людино-машинного інтерфейсу оператора-технолога ще й людино-машинний інтерфейс технічного супроводження цієї системи, що скорочує витрати часу на її налагодження і супроводження. Такий інтерфейс повинен забезпечувати швидкий і зручний доступ до вхідних/вихідних сигналів з можливостями їх безпосереднього зчитування/видачі через модулі введення/виведення з/на зовнішні пристрої, а також налаштування параметрів елементів апаратно-програмного комплексу систем після їх заміни, після ремонтів, при внесенні змін у технологічний процес тощо.

Організація людино-машинного інтерфейсу технічного супроводження у вигляді дворовневого блокнота, в якому на першій сторінці відображається меню, за яким відбувається вибір модуля введення/виведення для перевірки зовнішніх ланцюгів підключення системи чи параметрів, що підлягають налаштуванню, а на другій, підлеглий, відображається у табличному вигляді інструментарій для такого налаштування, забезпечує прозору та зручну реалізацію такого інтерфейсу.

Людино-машинний інтерфейс технічного супроводження складних інформаційно-керуючих систем автоматизації скорочує час на пошук порушень у зовнішніх ланцюгах підключення системи і налаштування її параметрів, що в кінцевому рахунку запобігає непродуктивним простоям технологічного процесу і спрощення вимог до технічного персоналу, що її обслуговує.

ЛІТЕРАТУРА:

1. Пупена О.М. Розроблення людино-машинних інтерфейсів та систем збирання даних з використанням програмних засобів SCADA/HMI : навч. посіб. Київ : Видавництво Ліра-К, 2020. 594 с.
2. Bailey D., Wright E. Practical SCADA for Industry. Newnes, 2005. 304 p.
3. Автоматизована система диспетчерського керування WindEx. URL: https://activolt.com.ua/wp-content/uploads/2021/04/windex-2_2021.pdf (дата звернення 24.04.2023).
4. Навчальні матеріали SCADA zenon. Базовий курс. URL: http://www.copa-data.com.ua/files/pdf/zenon_basic_training_UKR.pdf (дата звернення 24.04.2023).
5. Зайцев В.Г., Цибаєв Є.І. Комп'ютерні системи реального часу : навчальний посібник. Київ : КПІ ім. Ігоря Сікорського, 2019. 162 с.
6. Куваєв В., Іщук П., Політов А., Буряк В. Програмні рішення по забезпеченню надійного функціонування складних інформаційно-керуючих систем критичних до режиму реального часу. *Information Technology: Computer Science, Software Engineering and Cyber Security*. 2021. № 1. С. 16–24. DOI: <https://doi.org/10.32782/IT/2021-1-3>
7. Графічна оболонка Photon – революція у світі інтерфейсів. URL: <https://www.rts.ua/rus/articles/360/0/32/> (дата звернення 24.04.2023).
8. Доценко С. І. Людино-машинний інтерфейс : навч. посібник. Харків : УкрДУЗТ, 2022. 135 с.
9. Куваєв В., Мещеряков Л., Харь А., Політов А. Розробка інтерфейсу оператора складних інформаційно-керуючих систем критичних до режиму реального часу. *Information Technology: Computer Science, Software Engineering and Cyber Security*. 2023. № 1. С. 41–48.

REFERENCES:

1. Pupena, O.M. (2020). Rozroblennia liudyno-mashynnykh interfeisiv ta system zbyrannia danykh z vykorystanniam prohramnykh zasobiv SCADA/HMI: Navch. posib. [Development of human-machine interfaces and data collection systems using SCADA/HMI software]. Kyiv : Lira-K [in Ukrainian].
2. Bailey, D., Wright, E. (2005). Practical SCADA for Industry. Newnes.
3. Avtomatyzovana systema dyspetcherskoho keruvannia WindEx [WindEx automated dispatch control system]. *activolt.com.ua*. Retrieved from https://activolt.com.ua/wp-content/uploads/2021/04/windex-2_2021.pdf [in Ukrainian].
4. Navchalni materialy SCADA zenon. Bazovyi kurs [Educational materials SCADA zenon. Basic course]. *copa-data.com.ua*. Retrieved from http://www.copa-data.com.ua/files/pdf/zenon_basic_training_UKR.pdf [in Ukrainian].
5. Zaitsev, V.G., Tsybaev, E.I. (2019). Kompiuterni systemy realnoho chasu. Navchalnyi posibnyk [Real-time computer systems. Tutorial]. Kyiv: KPI im. Ihoria Sikorskoho [in Ukrainian].
6. Kuvaev, V., Ishchuk, P., Politov, A., Buriak, V. (2021). Prohramni rishennia po zabezpechenniu nadiinoho funktsionuvannia skladnykh informatsiino-keruiuchykh system krytychnykh do rezhymu realnoho chasu [Software solutions to ensure the reliable operation of complex information and control systems critical to real-time]. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 1, 16–24. [in Ukrainian]. doi: <https://doi.org/10.32782/IT/2021-1-3>
7. Hrafichna obolonka Photon – revoliutsiia u sviti interfeisiv [The Photon graphic shell is a revolution in the world of interfaces]. *rts.ua*. Retrieved from <https://www.rts.ua/rus/articles/360/0/32/> [in Ukrainian].
8. Dotsenko, S.I. (2022). Liudyno-mashynnyi interfeis: navch. posibnyk [Human-machine interface: tutorial]. Kharkiv: UkrDUZT [in Ukrainian].
9. Kuvaev, V., Meshcheriakov L., Khar A., Politov A. (2023). Rozrobka interfeisu operatora skladnykh informatsiino-keruiuchykh system krytychnykh do rezhymu realnoho chasu [Development of the operator interface of complex information and control systems critical to the real-time mode]. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 1, 41–48 [in Ukrainian].

УДК 004

DOI <https://doi.org/10.32782/IT/2023-2-4>

Борис МОРОЗ

доктор технічних наук, професор, професор кафедри програмного забезпечення комп'ютерних систем, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, м. Дніпро, Україна, 49005, moroz.boris.1948@gmail.com

ORCID: 0000-0002-5625-0864

Scopus-Author ID: 57218242332

Леонід КАБАК

кандидат технічних наук, доцент, доцент кафедри програмного забезпечення комп'ютерних систем, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, м. Дніпро, Україна, 49005, kabak.leo@gmail.com

ORCID: 0000-0001-6267-1772

Scopus-Author ID: 5720222205

Нонна ВАРЕХ

кандидат наук із соціальних комунікацій, доцент кафедри мовних та гуманітарних дисциплін, Технічний університет «Метінвест політехніка», вул. Південне шосе, 80, м. Запоріжжя, Україна, 69008, nonna.varekh@gmail.com

ORCID: 0000-0002-2779-9225

WoS ResearcherID: ABG-3294-2021

Дмитро МОРОЗ

викладач кафедри програмного забезпечення комп'ютерних систем, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, м. Дніпро, Україна, 49005, dmitriy@moroz.cc

ORCID: 0000-0003-2577-3352

Scopus-Author ID: 57369936300

Бібліографічний опис статті: Мороз, Б., Кабак, Л., Варех, Н., Мороз, Д. (2023). Система класифікації текстових документів із використанням технологій Big Data. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 34–40, doi: <https://doi.org/10.32782/IT/2023-2-4>

СИСТЕМА КЛАСИФІКАЦІЇ ТЕКСТОВИХ ДОКУМЕНТІВ ІЗ ВИКОРИСТАННЯМ ТЕХНОЛОГІЙ BIG DATA

У роботі було розглянуто модель системи класифікації документів з використанням технології Big Data. При використанні технології Big Data на сервері накопичується великий масив документів, які потрібно попередньо обробити та завантажити у базу даних. В документах потрібно визначити ключові слова за допомогою яких їх потрібно віднести до однієї або декількох тематичних розділів. Крім того розроблена система повинна працювати швидко та передбачати автоматичне навчання.

Отже розробка моделей та методів класифікації текстових документів на дійсний час є актуальним завданням. Дуже інтенсивний розвиток цих методів спостерігається в останній час при стрімкому розвитку обчислювальної техніки, та при переході багатьох організацій на електронний документообіг. В результаті дослідження було розроблено метод та модель системи; запропоновано комбінацію підходів для навчання моделі; визначено найбільш продуктивну модель для навчання системи.

Метою роботи є проведення аналізу існуючих методів класифікації текстових документів та розробити модель та метод класифікації текстових документів з використанням технології MapRaduce.

Методологія вирішення поставленого завдання полягає в проведенні порівняльного аналізу показників продуктивності різних конфігурацій системи, які запроваджені з урахуванням попередніх досліджень моделей систем класифікації документів, які використовують технологію Big Data.

Наукова новизна. У роботі запропоноване нове рішення для виконання точної байєсовської класифікації на основі Spark. Цей класифікатор використовує велику кількість операцій в пам'яті сервера, щоб класифікувати велику кількість текстових документів на основі великого навчального набору даних з використанням MapReduce. Фаза карти обчислює кількість входжень ключових слів у різних розподілах даних навчання. Після цього кілька редукторів обчислюють вірогідність віднесення документу до певних класів,

на підставі обчислень отриманих на етапі карти. Ключовий момент цієї пропозиції полягає в управлінні набором текстових документів, зберігаючи їх в пам'яті, коли це можливо.

Висновки. Результати даної роботи можуть бути використані для реалізації ефективної системи класифікації текстової документації, яка використовує точний байєсовський класифікатор, з використання мови програмування Python в поєднанні з сервісом Hadoop Big Data.

Ключові слова: Big Data, Hadoop, Map Reduce, Apache Spark, алгоритми машинного навчання, системи класифікації, байєсовський класифікатор.

Borys MOROZ

Doctor of Technical Sciences, Professor, Professor at the Department of Software Engineering, Dnipro University of Technology, Dmytro Yavornytskyi ave., 19, Dnipro, Ukraine, 49005, moroz.boris.1948@gmail.com

ORCID: 0000-0002-5625-0864

Scopus-Author ID: 57202222055

Leonid KABAK

Candidate of Technical Sciences, Associate Professor, Associate Professor at the Department of Software Engineering, Dnipro University of Technology, Dmytro Yavornytskyi ave., 19, Dnipro, Ukraine, 49005; Associate Professor at the Department of Digital Technologies and Project-Analytical Solutions, Technical University "Metinvest Polytechnic", Pivdenne shose str., 80, Zaporizhzhia, Ukraine, 69008, kabak.leo@gmail.com

ORCID: 0000-0001-6267-1772

Nonna VAREKH

Candidate of Sciences in Social Communications, Associate Professor at the Department of Linguistic and Humanitarian Disciplines, Technical University "Metinvest Polytechnic", Pivdenne shose str., 80, Zaporizhzhia, Ukraine, 69008, nonna.varekh@gmail.com

ORCID: 0000-0002-2779-9225

WoS ResearcherID: ABG-3294-2021

Dmitriy MOROZ

Lecturer at the Department of Software Engineering, Dnipro University of Technology, Dmytro Yavornytskyi ave., 19, Dnipro, Ukraine, 49005, dmitriy@moroz.cc

ORCID: 0000-0003-2577-3352

Scopus-Author ID: 57369936300

To cite this article: Moroz, B., Kabak, L., Varekh, N., Moroz, D. (2023). Systema klasyfikatsii tekstovoykh dokumentiv iz vykorystanniam tekhnolohii Big Data [Text document classification system with Big Data technologies usage]. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 34–40, doi: <https://doi.org/10.32782/IT/2023-2-4>

TEXT DOCUMENT CLASSIFICATION SYSTEM WITH BIG DATA TECHNOLOGIES USAGE

The aim. The paper considered a model of the document classification system using Big Data technology. When using Big Data technology, a large array of documents accumulates on the server which must be pre-processed and uploaded to the database. In the documents you need to define keywords with a help of which you need to assign them to one or more thematic sections. In addition, the developed system should operate fast and provide automatic learning.

Therefore, the development of models and methods of classification of text documents for real time is an urgent task. A very intensive development of these methods has been observed recently with the rapid development of computer technology and with the transition of many organizations into electronic document management. As a result of the study, a method and a system model were developed; a combination of approaches for model training is proposed; the most productive model for system training is determined.

Scientific novelty. The paper proposes a new solution for performing accurate Bayesian classification based on Spark. This classifier uses a large number of in-memory server operations to classify a large number of text documents based on a large training dataset using MapReduce. The map phase calculates the number of occurrences of keywords in different distributions of the training data. After that, several reducers calculate the probability of assigning of a document to certain classes based on the calculations obtained at the map stage. The key point of this proposal is to manage a set of text documents keeping them in memory whenever possible.

Conclusions. The results of this work could be used for implementation of an effective classification system for text documents that uses an accurate Bayesian classifier developed with the Python programming language in combination with the Hadoop Big Data service.

Key words: Big Data, Hadoop, Map Reduce, Apache Spark, Machine Learning Algorithm, systems of classification, Bayes Classifier.

Актуальність проблеми. При використанні в автоматизованих інформаційних системах технологій Big Data виникає необхідність в обробці великих масивів документів. Однією з актуальних задач в процесі обробки текстових документів є їх попередня класифікація, яка передбачає віднесення документа до одного або декількох тематичних розділів. В дійсний час розроблено багато методів що дозволяють вирішити такі завдання. Однак більшість методів розраховані на пошук документів в мережі Internet за допомогою ключових слів. Пошукова система видає користувачу безліч документів з яких користувач вибирає необхідні йому документи.

При використанні технології Big Data на сервері накопичується великий масив документів, які потрібно попередньо обробити та завантажити у базу даних. В документах потрібно визначити ключові слова за допомогою яких їх потрібно віднести до однієї або декількох тематичних розділів. Крім того розроблена система повинна працювати швидко та передбачати автоматичне навчання.

Отже розробка моделей та методів класифікації текстових документів на дійсний час є актуальним завданням. Дуже інтенсивний розвиток цих методів спостерігається в останній час при стрімкому розвитку обчислювальної техніки, та при переході багатьох організацій на електронний документообіг.

Отже, основна мета цієї роботи є розробка інформаційної системи в якій передбачається обробка та класифікація документів, які потрапляють до сервера, з використанням існуючих та з використанням нових методів класифікації.

Аналіз останніх досліджень і публікацій. Проблемою розпізнавання образів та автоматизованої класифікації текстових документів вчені займаються починаючи з 60-х років класичними роботами в цьому напрямку є роботи (Gonzalez, Thomason, 1974; Gonzalez, Tou, 1968; Salton, 1986) вчених Gonzalez R.C., Thomason M G. в цих роботах було досліджені різні типи класифікаторів. Класичним методом класифікації в асоціативно-статистичному підході є метод і алгоритм, що його реалізує, запропонований Дж. Солтоном в 1975 р (Salton, 1986). У 1979 р. алгоритм був уточнений і доопрацьований у Дж. Солтоном і названий TFxIDF. Цей алгоритм є на сьогоднішній день найбільш ефективним, поширеним і використовується в сучасних інформаційно-пошукових системах.

Метод, розроблений Дж. Солтоном, ґрунтується на так званій «векторній моделі тексту». У різних джерелах можна зустріти різну назву

цієї моделі: «векторна», «лінійна» або «алгебраїчна».

В дійсний час існує багато методів класифікації текстових документів які використовують нейронні мережі, різноманітні методи кластеризації, методи опорних векторів. У статті пропонується розробка автоматизованої системи класифікації тестових документів з використанням платформи Apache Spark яка вбудована в систему Big Data Hadoop. Класичними роботами у цьому напрямку є роботи авторів П. Семберескі та Г. Мацієвський (Semberecki, Maciejewski, 2016). У цій статті було показано, як ці послідовні кроки можна реалізувати на платформі Apache Spark, призначеній для розподіленої обробки великих даних. Авторами статті було проілюстровано запропонований метод, який є зразком класифікатора, призначеного для прогнозування тематичної категорії документа в англомовній Вікіпедії.

У роботі (Pintye, Kail, Kacsuk, Lovas) було розглянуто різноманітність технологій і протоколів. У документі зосереджено увагу на широко поширеному кластері Apache Spark із Jupyter як особливо адресованою структурою, а також на інструменті Ossorus, що не залежить від хмари, для автоматизації етапів його розгортання та обслуговування. У цій статті був представлений підхід та було продемонстровано та перевірено за допомогою нової багатообіцяючої програми класифікації тексту в інфраструктурі академічних досліджень Угорщини, MTA Cloud на основі OpenStack. У статті пояснюється концепція, застосовані компоненти та ілюструється їх використання за допомогою реальних вимірювань у випадку використання.

У статті (Chaudharil, Patil, Ghorpade, 2020) розглянуті різні дисципліни мають які мають справу з великими даними, які включають велику кількість функцій. Зібрані дані можуть аналізуватися для виявлення знань і використовуються для прийняття рішень. Для того щоб отримати необхідні знання використовують різні алгоритми машинного навчання. В роботі доказано цей факт що для отримання чудових результатів класифікації потрібно об'єднати методи навчання та інструментальні засоби класифікації. В роботі досліджені існуючі методи інтелектуального аналізу даних, які найбільш широко використовуються для класифікації та кластеризації даних, такі як k-Means, Support Vector Machine, Naive Bayes і k-Nearest Neighbor разом із Map Reduce, Apache Spark. Інструменти аналізу даних Map Reduce, Apache Spark, застосований до класифікації або кластеризації даних, забезпечує кращу продуктивність. У цій статті представлено екосистему Hadoop та

дослідження різні методи класифікації та кластеризації за допомогою Map Reduce і Apache Spark.

У роботі (Gopalani, Arora, 2015) обговорюється два варіанти порівняння – Hadoop Map Reduce і нещодавно представлений Apache Spark – обидва з яких забезпечують модель обробки для аналізу великих даних. Незважаючи на те, що обидва ці варіанти базуються на концепції великих даних, їх продуктивність суттєво відрізняється залежно від варіанту використання, який реалізується. Ось що робить ці два варіанти гідними аналізу з огляду на їхню мінливість і різноманітність у динамічному полі великих даних. У цій статті було досліджено ці дві системи разом із аналізом продуктивності за допомогою стандартного алгоритму машинного навчання для кластеризації (K-Means).

У статті (Maillo, Ramírez, Triguero, Herrera, 2017) досліджується k-Nearest Neighbors класифікатор, який є простим, але ефективним широко відомим методом інтелектуального аналізу даних. Фактичне застосування цієї моделі в області великих даних є неможливим через обмеження часу та пам'яті. У роботі запропоновано кілька розподілених альтернатив на основі MapReduce, щоб цей метод міг обробляти великомасштабні дані. Однак їх продуктивність можна ще більше покращити за допомогою нових конструкцій, які відповідають новим технологіям. В цій статті пропонується метод який на основі MapReduce використовувати k-Nearest Neighbors класифікатор найбільш ефективно.

Мета статті – провести аналіз існуючих методів класифікації текстових документів та розробити модель та метод класифікації текстових документів з використанням технології MapReduce.

Виклад основного матеріалу. В класичному вигляді система класифікації працює наступним чином на сервер потрапляє безліч документів. Кожен документ є членом колекції документів D_i є членом можливих категорій S_j , таких як $\{s_1, s_2, \dots, s_j\}$, тоді класифікація тексту є операцією зіставлення кожного документа з одним або декількома класами. У роботі досліджена класифікація яка виконана на максимум п'ять категорій. Категорії освіта, спорт, культура, політика та світ. Типові кроки для застосування алгоритмів машинного навчання на основі текстових даних показано на рис. 1. Ті самі кроки використовувались і в програмі класифікації.

Система класифікації буде відносити документ який поступає на сервер до певного класу S_j . Для того щоб віднести документ до певного класу система повинна мати інформацію про документ а саме інформацію про ключові слова що входять в документ. Виходячи з цього під системою класифікації розуміється система яка відносить D_i до певних класів S_j . При прийнятті рішень про віднесення документа до певного класу система повинна робити на підставі якоїсь інформації. Ця інформація у теорії розпізнавання образів і називається вектором ознак. В якості таких ознак саме і розглядається саме факт входження певного набору ключових термінів у документ $V = \{v_1, v_2, \dots, v_m\}$. Таким чином вектором ознак документа D_i буде $x' = \{x_1, x_2, \dots, x_n\}$ де x приймає значення 0 якщо x_n термін не входить до документа та 1 якщо входить. Таким чином завданням системи класифікації на підставі класифікації попер D_{i-1} документів віднести D_i до одного чи декількох класів. Розглянемо роботу системи класифікації. Система

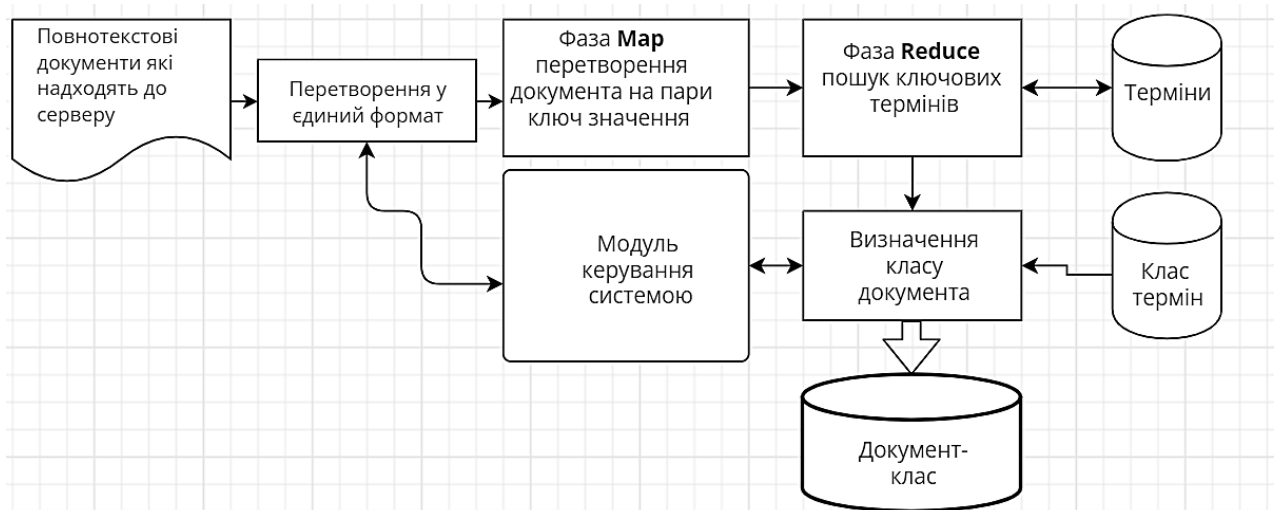


Рис. 1. Модель системи класифікації

класифікації має віднести до певного класу. Якщо система вірно віднесла документ до певного класу це вважається успіхом в зворотному випадку програшом системи, чи помилкою системи. У статичній теорії розпізнавання образів найбільш поширеним є байєсовський класифікатор (Gonzalez, Thomason, 1974) який засновано на мінімізації середніх програшів системи класифікації. Під програшом системи будемо вважати невірну класифікацію документу. Будемо вважати що система віднесе документ до класу S_j з вірогідністю $P(S_j)$. Стратегією системи класифікації буде вказівка класу S_j до якого вона відносить даний документ. Зазначена гра характеризується матрицею втрат L_{ij} . Припустимо, що на вхід системи класифікації потрапляють незалежні документи D_i , керуючись ймовірностями $P(S_j)$. Тоді математичне очікування втрат, пов'язаних із віднесенням системою образу x^i до класу S_i , визначиться формулою

$$r_i(x^i) = \sum_{j=1}^n L_{ij} P(S_j | x^i). \quad (1)$$

В цій формулі $P(S_j | x^i)$ – умовна ймовірність віднесення документа до класу S_j , коли на вхід системи було подано образ x^i .

Байєсовським називається класифікаційне правило, яке мінімізує $r_i(x^i)$ (Gonzalez, Thomason, 1974).

Найчастіше природно припустити, що за правильної класифікації рівні нулю, а інакше рівні одиниці, тобто.

$$L_{ij} = 1 - \delta_{ij} \quad (2)$$

де δ_{ij} – символи Кронекера.

Для матриці втрат (2) байєсовське класифікаційне правило визначатиметься за допомогою функції вигляду:

$$d_i(x^i) = P(x^i | S_i) P(S_i). \quad (3)$$

Якщо компоненти випадкового вектора $x^i = \{x_1, x_2, \dots, x_m\}$ незалежні, умовна ймовірність визначається виразом

$$P(x^{(i)} | S_i) = \prod_{k=1}^m \gamma_{ik}, \quad (4)$$

де

$$\gamma_{ik} = \begin{cases} p_{ik}, & \text{если } x_k = 1 \\ 1 - p_{ik}, & \text{если } x_k = 0. \end{cases} \quad (5)$$

Тут p_{ik} – ймовірність появи k -го терміна зі словника у документі, що відноситься до класу S_i .

Далі розглянемо алгоритм роботи системи.

Крок 1. На вхід системи потрапляє документ, система на фазі Мар робить лексичний розбір документу.

Для цього користуємось додатком Zeppelin. Для прикладу у папку tmp серверу Hadoop розміщуємо документ Doc1.txt, який потребує класифікації.

Далі читаємо документ. Перетворюємо документ на пари ключ значення, та підраховуємо входження термінів.

```
%spark2.pyspark
lines = spark.sparkContext.textFile("/tmp/doc1.txt")
words = lines.flatMap(lambda line: line.split(" "))
# Видаляємо порожні слова.
wordsFiltered = words.filter(lambda w: len(w) > 0)
# підраховуємо кількість слів в документі.
wordc = wordsFiltered.count()
```

Після текстового розбору документа потрібно відшукати значущі слова. Для оцінки значущості на практиці використовують методи, які враховують частоту народження термінів і характеристики, що відображають деякі структурні властивості тексту, наприклад, спільну частоту народження (асоційованість) і щільність розподілу термінів у тексті (надфразові властивості). Використовуючи різні математичні моделі, всі ці методи мають загальне обґрунтування в рамках нейропсихологічної моделі «грубою» обробки інформації у правій півкулі мозку. Перевагою такого підходу є алгоритмічна простота, яка не вимагає точного лінгвістичного аналізу.

Статистичний аналіз тексту використовується на вирішення завдання виділення ключових слів довільного документа. В усіх документах можна назвати статистичні закономірності. Внутрішня структура тексту описується законами Дж. Зіпфа (Zipf, 1949).

Перший закон Зіпфа «ранг-частота» говорить про те, що якщо виміряти кількість входжень кожного слова в текст і взяти тільки одне значення з кожної групи, що має однакову частоту, після чого розташувати частоти в міру їх зменшення та пронумерувати (порядковий номер частоти називається рангом частоти), то слова, що найчастіше зустрічаються, матимуть ранг 1, наступні за ними – 2 і т. д. Ймовірність появи довільно вибраного слова P у тексті дорівнюватиме відношенню кількості входжень цього слова K до загального числа слів у тексті N .

$$P = \frac{K}{C}. \quad (6)$$

Дж. Зіпф виявив таку закономірність: добуток ймовірності P виявлення слова на ранг частоти R дорівнює константі C .

$$C = P \times R. \quad (7)$$

За першим законом Зіпфа, якщо найпоширеніше слово зустрічається у мові K разів,

то наступне за частотою слово зустрічається з частотою ймовірності раз. Значення константи C у різних мовах по-різному, але всередині однієї мовної групи залишається незмінним (Zipf, 1949).

Другий закон Зіпфа «кількість-частота» говорить про те, що різні слова можуть входити до тексту з однаковою частотою, причому частота та кількість слів, що входять до тексту із цією частотою, теж пов'язані між собою. Відносини частоти входження слова до рангу слів представлено рис. 2 (Zipf, 1949).

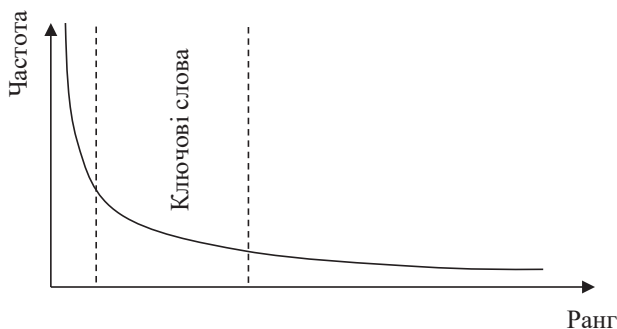


Рис. 2. Відношення частоти входження слова до рангу слів

З наведеного на рис. 2 графіка ми бачимо, що значні слова лежать у середній частині діаграми. Слова, які зустрічаються як занадто часто в тексті (в основному вони виявляються прийменниками та займенниками), так і дуже рідко, в більшості випадків не мають вирішального значення (Zipf, 1949).

Крок 2. На фазі Reduce підраховуємо кількість входження кожного слова в документ. Та за формулою (7) підраховуємо ранг слова.

Для цього знаходимо найпопулярнішими, виконавши підрахунок слів за допомогою перетворень `map()` і `reduceByKey()` для створення кортежів типу `(word, count)`.

```
wordCounts = wordsFiltered.map(lambda word: (word, 1)).
reduceByKey(lambda a,b: a+b)
```

Далі перетворюємо пари ключ значення на дата фрейм. Та знаходимо значимі слова.

```
wordsCounts = (filteredWordCounts.map(lambda (w, c):
    Row(word=w, rang=c)) .toDF())
```

Ми можемо для прикладу переглянути значущі слова які виявила система.

```
wordCounts.show()
```

rang	word
1	colleges
1	particularly
1	institutional
1	professors
1	multitasking
1	through
1	cost
1	multitask
1	speed
1	relationships
1	group
1	work

Рис. 3. Ранг слів що входять до документу

Далі на підставі вибраних з документу значущих термінів система за допомогою байєсовського правила формула (3) приймає рішення саме до якого класу віднести цей документ. На рисунку 3 наведено входження ключових слів у документ. Як ми бачимо всі ключові слова мають ранг один так як документ взятий для прикладу був не великим.

Висновки із цього дослідження і перспективи подальших розвідок у цьому напрямку.

У роботі запропоноване нове рішення для виконання точної байєсовської класифікації на основі Spark. Цей класифікатор використовує велику кількість операції в пам'яті сервера, щоб класифікувати велику кількість випадків текстових документів на основі великого навчального набору даних з використанням MapReduce. Фаза карти обчислює кількість входжень ключових слів у різних розподілах даних навчання. Після цього кілька редукторів обчислюють вірогідність віднесення документу до певних класів, на підставі обчислень отриманих на етапі карти. Ключовий момент цієї пропозиції полягає в управлінні набором текстових документів, зберігаючи їх в пам'яті, коли це можливо. В іншому випадку він розбивається на мінімальну кількість частин, застосовуючи MapReduce для кожної частини, використовуючи навички кешування Spark для повторного використання попереднього розділеного навчального набору.

ЛІТЕРАТУРА:

1. Gonzalez R.C., Thomason M G. Tree Grammars and Their Application to Pattern Recognition. Tech. Rep. TR-EE/CS-74-10, Electrical Engineering Dept., Univ. of Tennessee, Knoxville. 1974. P. 364.
2. Gonzalez R C., Thomason M.G. Inference of Tree Grammars for Syntactic Pattern Recognition. Tech. Rept. TR-EE/CS-74-20, Electrical Engineering Dept., University of Tennessee, Knoxville. 1974. P. 160.
3. Gonzalez R.C., Tou J.T. Some Results in Minimum-Entropy Feature Extraction. IEEE Convention Record. Region III. 1968.
4. Salton G. Another look at automatic text-retrieval systems. *Commun. ACM*. 1986. № 7. P. 648–656. 2000. ISBN 951-22-5145-0
5. Semberecki P., Maciejewski H. Distributed Classification of Text Documents on Apache Spark Platform. International Conference on Artificial Intelligence and Soft Computing. June 2016. P. 621–629. DOI:10.1007/978-3-319-39378-0_53 [Scopus].
6. I. Pintye, E. Kail, P. Kacsuk, R. Lovas. Big data and machine learning framework for clouds and its usage for text classification. Volume 33. Issue 19. Special Issue: Human oriented solutions for intelligent analysis, multimedia and communication systems (Human Oriented Solutions 2020). Science Gateways Special Issue (Science Gateways 2020) 10 October 2021. <https://doi.org/10.1002/cpe.6164>.
7. Ratna S. Chaudhari¹, Seema S. Patil, Smita J. Ghorpade. Classification and clustering methods along with Map Reduce, Apache Spark: a study. *IJRAR*. November 2020. Volume 7. Issue 4.
8. Gopalani S., Arora R. Comparing Apache Spark and Map Reduce with Performance Analysis using K-Means. *March 2015 International Journal of Computer Applications*. 113(1). P. 8–11. DOI:10.5120/19788-0531
9. Maillio J., Ramírez S., Triguero I., Herrera F. kNN-IS: An Iterative Spark-based design of the k-Nearest Neighbors classifier for big data. June 2016. *Knowledge-Based Systems*. 1 February 2017. Volume 117. P. 3–15. DOI:10.1016/j.knosys.2016.06.012
10. Zipf G.K. Human Behavior and the Principle of Least Effort. Cambridge, 1949. P. ix, 3, 5–8.

REFERENCES:

1. Gonzalez R.C., Thomason M G. Tree Grammars and Their Application to Pattern Recognition. Tech. Rep. TR-EE/CS-74-10, Electrical Engineering Dept., Univ. of Tennessee, Knoxville. 1974. P. 364.
2. Gonzalez R C., Thomason M.G. Inference of Tree Grammars for Syntactic Pattern Recognition. Tech. Rept. TR-EE/CS-74-20, Electrical Engineering Dept., University of Tennessee, Knoxville. 1974. P. 160.
3. Gonzalez R.C., Tou J.T. Some Results in Minimum-Entropy Feature Extraction. IEEE Convention Record. Region III. 1968.
4. Salton G. Another look at automatic text-retrieval systems. *Commun. ACM*. 1986. № 7. P. 648–656. 2000. ISBN 951-22-5145-0
5. P. Semberecki, H. Maciejewski. Distributed Classification of Text Documents on Apache Spark Platform. International Conference on Artificial Intelligence and Soft Computing. June 2016. P. 621–629. DOI:10.1007/978-3-319-39378-0_53 [Scopus].
6. Pintye I., Kail E., Kacsuk P., Lovas R. Big data and machine learning framework for clouds and its usage for text classification. Volume 33. Issue 19. Special Issue: Human oriented solutions for intelligent analysis, multimedia and communication systems (Human Oriented Solutions 2020). Science Gateways Special Issue (Science Gateways 2020) 10 October 2021. <https://doi.org/10.1002/cpe.6164>.
7. Ratna S. Chaudhari¹, Seema S. Patil, Smita J. Ghorpade. Classification and clustering methods along with Map Reduce, Apache Spark: a study // *IJRAR* November 2020, Volume 7, Issue 4.
8. Gopalani S., Arora R. Comparing Apache Spark and Map Reduce with Performance Analysis using K-Means. *March 2015 International Journal of Computer Applications* 113(1). P. 8–11. DOI:10.5120/19788-0531
9. Maillio J., Ramírez S., Triguero I., Herrera F. kNN-IS: An Iterative Spark-based design of the k-Nearest Neighbors classifier for big data. June 2016. *Knowledge-Based Systems* Volume 117, 1 February 2017, Pages 3–15. DOI:10.1016/j.knosys.2016.06.012
10. Zipf G.K. Human Behavior and the Principle of Least Effort. Cambridge, 1949. P. ix, 3, 5–8.

УДК 004.8

DOI <https://doi.org/10.32782/IT/2023-2-5>

Костянтин ОНИЩЕНКО

старший викладач кафедри програмної інженерії, Харківський національний університет радіоелектроніки, просп. Науки, 14, м. Харків, Україна, 61166, kostiantyn.onyshchenko@nure.ua

ORCID: 0000-0002-7746-4570

Scopus Author ID: 57209338164

Анатолій АФАНАСЬЄВ

студент кафедри інформатики, Харківський національний університет радіоелектроніки, просп. Науки, 14, м. Харків, Україна, 61166, anatolii.afanasiev@nure.ua

ORCID: 0009-0005-0707-981X

Бібліографічний опис статті: Онищенко, К., Афанасьєв, А. (2023). Оцінка ефективності використання згорткових та рекурентних нейронних мереж у завданні обробки текстових даних. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 41–49, doi: <https://doi.org/10.32782/IT/2023-2-5>

ОЦІНКА ЕФЕКТИВНОСТІ ВИКОРИСТАННЯ ЗГОРТКОВИХ ТА РЕКУРЕНТНИХ НЕЙРОННИХ МЕРЕЖ У ЗАВДАННІ ОБРОБКИ ТЕКСТОВИХ ДАНИХ

Дана робота присвячена оцінці ефективності використання згорткових та рекурентних нейромереж у завданні обробки текстових даних на прикладі виявлення фейкових новин. Нині зусилля світової спільноти спрямовані на боротьбу з подібною інформацією в цілому, що зумовлює актуальність порушеного питання. Проблематика виявлення фейкових новин полягає у достовірності визначення тої чи іншої інформації як фейкової чи правдивої. Метою роботи є порівняння точності визначення фейкових новин для архітектур згорткових та рекурентних нейромереж, в яких закладена модель синтаксичного аналізу текстів статті шляхом формування міток новин з використанням TF-IDF та Word Embedding. Для досягнення поставленої мети було проведено аналіз області застосування та визначено ключові особливості цього типу інформації. Було розглянуто теоретичну основу обраних архітектур та встановлено їх конфігурації відповідно до поставленої задачі. Для практичної реалізації вибраних типів нейромереж було створено експериментальне середовище. Була виявлена відносна ефективність використання рекурентних нейромереж у порівнянні зі згортковими, а також визначено можливі сценарії, в яких отримані результати можуть змінюватися. У результаті аналізу було встановлено, що на доступних даних згорткова нейромережа має більшу швидкість порівняно з рекурентною нейромережею, але при цьому надає менш точні результати класифікації. З урахуванням запропонованого правила порівняння ефективності, ймовірності помилок різних видів та можливості подолання розбіжностей між алгоритмами, отриманий приріст у продуктивності можна вважати незначним. Цей висновок відповідає світовій науковій практиці, яка рекомендує використовувати одну з запропонованих моделей або їх комбінацію під час аналізу текстової інформації, зокрема, у випадку наявності двох класів (фейкові та нефейкові дані), або при перевірці достовірності зображень.

Ключові слова: класифікація, CNN, RNN, ефективність, нейронні мережі.

Kostiantyn ONYSHCHENKO

Senior Lecturer at the Department of Software Engineering, Kharkiv National University of Radio Electronics, Nauky ave., 14, Kharkiv, Ukraine, 61166, kostiantyn.onyshchenko@nure.ua

ORCID: 0000-0002-7746-4570

Scopus Author ID: 57209338164

Anatolii AFANASIEV

Student at the Department of Informatics, Kharkiv National University of Radio Electronics, Nauky ave., 14, Kharkiv, Ukraine, 61166, anatolii.afanasiev@nure.ua

ORCID: 0009-0005-0707-981X

To cite this article: Onyshchenko, K., Afanasiev, A. (2023). Otsinka efektyvnosti vykorystannia zghortkovykh ta rekurentnykh neironnykh merezh u zavdanni obrobky tekstovykh danykh [Evaluation of the effectiveness of using convolutional and recurrent neural networks in the task of text data processing].

EVALUATION OF THE EFFECTIVENESS OF USING CONVOLUTIONAL AND RECURRENT NEURAL NETWORKS IN THE TASK OF TEXT DATA PROCESSING

This work is dedicated to evaluating the effectiveness of using convolutional and recurrent neural networks in the task of processing textual data, specifically in detecting fake news. Currently, the efforts of the global community are focused on combating such misinformation as a whole, which underscores the relevance of the addressed issue. The problem of detecting fake news lies in accurately determining the authenticity of the information. The objective of this work is to compare the accuracy of identifying fake news between convolutional and recurrent neural network architectures, which incorporate a syntactic analysis model of article texts by forming news labels using TF-IDF and Word Embedding. To achieve this goal, an analysis of the application domain was conducted, and the key characteristics of this type of information were identified. The theoretical foundation of the selected architectures was examined, and their configurations were established in accordance with the defined task. An experimental environment was created for the practical implementation of the chosen types of neural networks. The relative effectiveness of using recurrent neural networks compared to convolutional ones was revealed, and potential scenarios were identified where the obtained results may vary. As a result of the analysis, it was determined that the convolutional neural network performs computations faster than the recurrent neural network on the available data, but it provides less accurate classification results. Taking into account the proposed rule for performance comparison, the marginal increase in productivity can be considered insignificant, considering the probabilities of different types of errors and the potential for resolving discrepancies between algorithms. This conclusion aligns with the global scientific practice, which recommends using one of the proposed models or their combination for analyzing textual information, particularly when dealing with two classes (fake and non-fake data) or verifying the authenticity of images.

Key words: classification, CNN, RNN, efficiency, neural networks.

Вступ. Відповідно до Кембриджського словника фейкові новини – це історії, які виглядають як новини, поширюються в інтернеті або за допомогою інших засобів масової інформації, та, зазвичай, створені для впливу на політичні погляди або як жарт (Cambridge Dictionary, 2023). Насправді, їх історія доволі розлога, однак зі зростанням популярності соціальних мереж, особливо анонімних, ця проблема гостро постала перед світовим суспільством. У якості її каталізатора виступає ще і розвиток технологій спрямованих на коригування відео- та аудіо-інформації, і технологій створення ботів.

Станом на сьогодні кількість фейків зростає як в абсолютному, так і відносному значенні. Деякі з них можна віднести до неправильного суб'єктивного сприйняття реальної інформації, однак суттєва кількість є результатом пропаганди (Wesolowski, 2022). Найбільшими каналами такої інформації є Twitter (BBC News, 2022) та Telegram (Ukraine Crisis Media Center, 2022).

Загалом для боротьби з подібними новинами було створено багато ініціатив, це зокрема французький закон проти маніпуляції інформацією (Blossman, 2019), який був ухвалений задля боротьби з дискредитацією іммігрантів та Європейського Союзу після оголошення результатів Брекзиту (Danzig, 2017). У цьому законі зазначено, що платформи, які перевищують певну кількість відвідувань на

день, повинні мати законного представника у Франції та публікувати свої алгоритми, при цьому необхідно повідомляти про будь-який спонсорований контент, публікуючи ім'я автора та сплачену суму. Закон також вимагає від суддів кваліфікувати фейкові новини за такими трьома критеріями: явність, свідоме поширення в масових масштабах, призведення до порушень громадського порядку або скомпрометування результатів виборів. При цьому початкове рішення приймає спеціально створений комітет з етики.

Створення правової бази для регуляції фейкової інформації спостерігається і в Україні, зокрема це стаття 259 Кримінального кодексу, що регламентує відповідальність за завідомо неправдиве повідомлення про загрозу безпеці громадян (Закон України № 1292-IX, 2021).

Світова тенденція спрямована на боротьбу з подібною інформацією в цілому є позитивною, однак проблема полягає у визначенні того чи є інформацією фейковою чи ні. Якщо рішення базуватиметься не на фактах, а на експертній оцінці, то таку ситуацію можна розглядати як маніпулятивну і спрямовану на авторитаризацію країни.

Для демократичних країн процес визначення фейкових новин можна автоматизувати з використанням штучного інтелекту. Початково у якості базового методу застосовували наївний Байес (Jain, 2022) чи SVM

(Reddy, 2019, с. 945–946), однак з розвитком нейромереж найбільшої популярності здобули згорткові та рекурентні нейромережі.

Варто зазначити, що з огляду на численні дослідження, ці моделі можуть давати різний результат в залежності від предметної області (Pai, 2020). Тому було прийнято рішення розглянути ефективність використання цих моделей для визначення фейкових новин, з метою оптимізації їх вибору при побудови відповідної програмної системи.

Аналіз предметної галузі та наявних рішень. У першу чергу завдання визначення фейкових новин стосується процесу обробки природної мови (Natural Language Processing, NLP). Розглянемо існуючі рішення для вирішення задач NLP.

TF-IDF характеристика. Програмні системи не здатні напряму обробляти текстову інформацію, тому першочергово кожен текст має набути кількісного вигляду. Існує декілька основних методик подібної конвертації, однак наразі базовою є TF-IDF (Stecanella, 2019). TF – скорочення від term frequency – це частота кожного вживаного слова (точніше набору слів, які більш доцільно називати термінами). IDF – скорочення від inverse document frequency – це зворотне число термінів на документ. В цілому показник TF-IDF вказує на те, наскільки рідкісним є певний термін. Наприклад, вставні слова, сполучники чи вигукі будуть найпоширенішими, а, відповідно, матимуть низький TF-IDF.

Саме по собі отримання TF-IDF характеристики дозволяє лише ранжувати ті чи інші терміни, однак не перетворити текст у числове представлення. Для виконання цієї дії існує технологія Word Embedding, яка здійснює відображення слів чи фраз у вектори дійсних чисел. Формальна зазначена технологія є набором різноманітних методів, одним з яких є GloVe (Global Vectors). GloVe – це алгоритм для перетворення немаркованих даних (у нашому випадку термінів), у безперервні вектори, для зменшення розмірності (Pennington, 2022). Вектори GloVe попередньо навчені на даних з Вікіпедії та Gigaword 5, тож доволі добре вловлюють семантику речень. Варто зауважити, що цей алгоритм спрямований саме на тексти англійською мовою, а не українською.

Розглянемо як саме відрізняється фейкова новина від справжньої (без прив'язки до архітектурних моделей нейромереж). Задля демонстрації скористаємось набором даних щодо президентських виборів у 2016 році в Сполучених Штатах (Risdal, 2016). Він містить 20 015 новин, з яких 11 941 фейкові,

а 8074 справжні. Справжніми новинами у цьому набору виступають дані з відомих авторитетних веб-сайтів новин, наприклад, New York Times, Washington Post тощо. Вони містять багато різноманітної інформації, зокрема назва статті, текст, зображення, автор, веб-сайт та багато чого іншого. У подальшому будемо використовувати лише назву та текст. Розглянемо TF-IDF характеристику для заголовків, аби продемонструвати різницю між справжніми та фейковими новинами (рис. 1). Якщо заголовку немає, то встановлюємо його значення рівне «notitle».

Заголовки фейкових новин часто містять такі слова, як «notitle», «IN», «THE», «CLINTON» і багато незв'язних цифр, що представляють спеціальні символи. З цього можна зробити декілька цікавих висновків. По-перше, більшість фейкових новин не мають заголовків. Ці фейкові новини широко поширюються у вигляді твітів із кількома ключовими словами та гіперпосиланнями на новини в соціальних мережах. По-друге, у фейках більшість символів у великому регістрі. Мета – привернути увагу читачів, а справжні новини містять менше великих літер та загалом пишуться у стандартному форматі. По-третє, справжні новини містять більш детальні описи. Наприклад, імена (Джеб Буш, Мітч МакКоннелл тощо) і дієслова (left, claim, debate, poll тощо). Тепер для більш глибокого розуміння розглянемо проблему пошуку фейків з точки зору обчислювальної лінгвістики, психологічного позиціонування, лексичного різноманіття та аналізу настроїв.

Обчислювальна лінгвістика. Перше, що варто зауважити, фейкові новини в середньому містять менше слів, ніж справжні новини: 4360 слів проти 3943 (84 речення проти 69). Крім того діапазон слів у несправжніх новинах значно більший за реальні дані. Незважаючи на меншу кількість символів та речень, самі речення у фейкових новинах коротші. Це пов'язано з тим, що редактори та журналісти користуються певними газетними нормами, що включають довжину та вибір слів, відсутність граматичних помилок тощо. Натомість фейкові новини на цих правилах, зазвичай, не ґрунтуються. Відповідно до отриманих статистичних даних, справжні новини мають менше знаків питання, ніж фейкові. Причини можуть полягати в тому, що у фейкових новинах багато риторичних питань, які використовуються, щоб свідомо підкреслити ідеї та підсилити настрої. Окрім цього, можна зауважити, що обидва типи новин мають дуже мало вигуків, а, отже, і знаків оклику, хоча їхня кількість все ж більша саме у фейкових новинах.

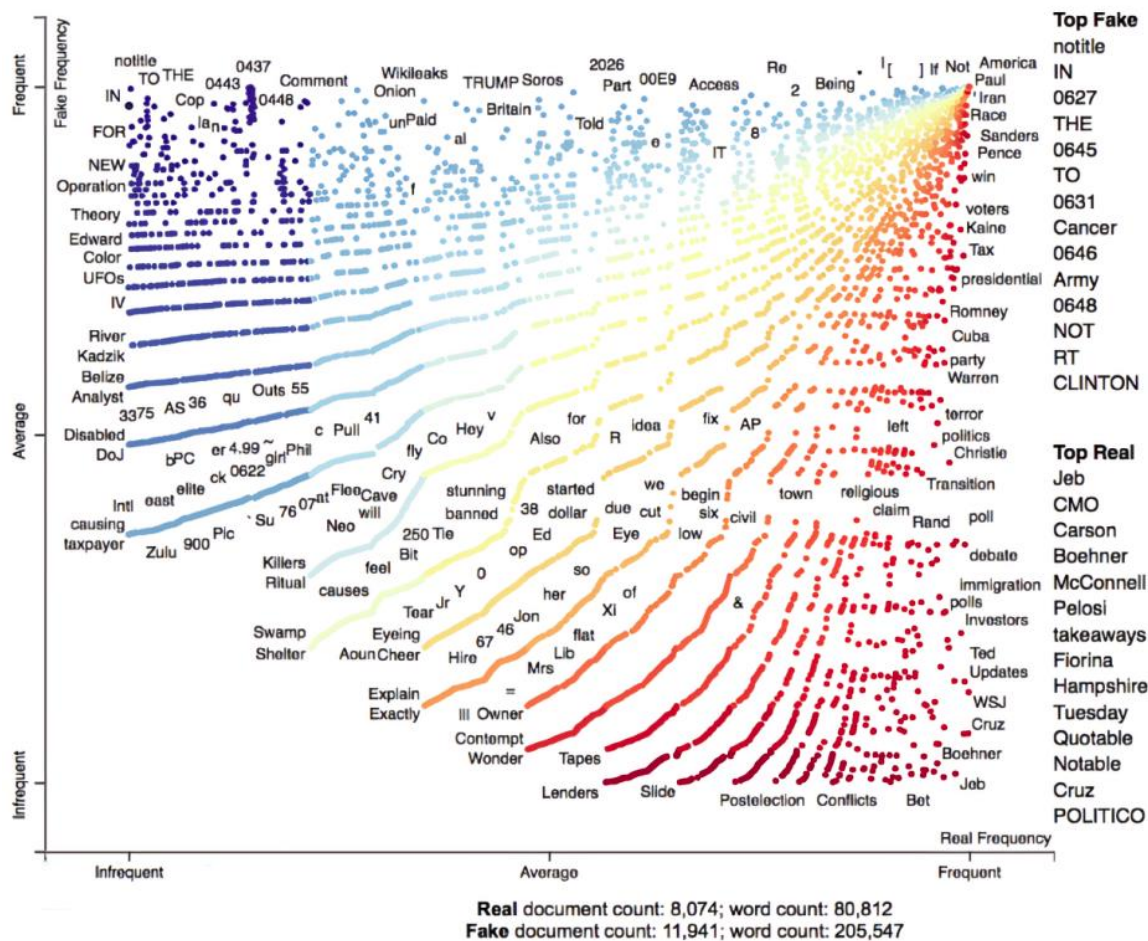


Рис. 1. TF-IDF характеристика для слів із заголовків

Постановка завдання. З'ясувавши основні нюанси, які можуть вплинути на перебіг дослідження формалізуємо мету роботи.

Дано набір з m текстових новин, який можна зобразити наступним чином:

$$\mathring{A} = \{A_i^T\}_i^m.$$

У задачі виявлення фейкових новин необхідно передбачити, чи є статті з $\mathring{A}$ фейковими чи ні. При цьому набір міток, що вказуватимуть на правдивість інформації можна представити у наступному вигляді:

$$y = \{1, 0\}_i^m,$$

де 1 – позначає справжню новину, а 0 – фейкову.

Набір функцій X_i^T з $A_i^T \in \mathring{A}$ має отримуватися шляхом синтаксичного аналізу текстів статті, зокрема з використанням TF-IDF та Word Embedding.

Метою чинної роботи є порівняння точності визначення фейкових новин для архітектур

CNN та RNN, в яких закладена наступна модель формування міток новин з $\mathring{A}$:

$$f : \{X_i : T_i\}_i^m \in \mathbb{X} \rightarrow y.$$

Дослідження ефективності. Перед тим як безпосередньо визначити експериментальні умови, наведемо опис алгоритму відповідно до якого здійснюватиметься підготовка даних до аналізу. Спочатку з тексту видаляється всі не буквені вирази (наприклад, числа, коми, крапки та інші розділові знаки) за допомогою бібліотеки *re*, яка дає доступ до регулярний виразів. Після очищення текст потрапляє до функції обробки (рис. 2). Основними в цій функції є методи, що надаються бібліотекою *nlTK*, яка дозволяє обробляти природну мову, виходячи із вбудованих корпусів слів. Першочергово необхідно видалити з набору ті слова, що не несуть інформаційного навантаження (наприклад, «and», «or» тощо). Ці слова заважатимуть коректному проведенню аналізу.

Наступним кроком необхідно нівелювати лінгвістичну варіативність, що породжується

використанням морфем. Для цього застосовується операція стемінгу – процес скорочення слова до його основи. Наприклад, слова «eating» та «eaten» будуть замінені на слово «eat». Після цієї обробки масив слів скорочується до набору основ.

Вбудований в бібліотеку *nltk* клас, що здійснює зазначений процес, не завжди працює коректно, тому для покращення результату було вирішено поєднати стемінг з іншою операцією – лематизацією. Вона дозволяє привести словоформу до лема (нормальної словникової форми). Наприклад, при обробці слова «good» та «better» матимуть різні основи, однак лема у цих слів одна – «good». Після проведення зазначених дій набір слів потребує ще однієї перевірки на неінформативний вміст.

Після обробки тексту, утворюється словник необхідний для коректного визначення TF-IDF характеристики. Задля врахування емоційного забарвлення слів наступним кроком є знаходження для кожного слова частотно-полярної характеристики, шляхом перемноження показників TF-IDF та полярності. Останній можна отримати за допомогою використання класу *SentimentIntensityAnalyzer*, що знаходиться в модулі *vader* бібліотеки *nltk*.

Ефективність певної моделі (E) будемо визначати наступним чином:

$$E = f(\vartheta, RMSE, v),$$

де ϑ – швидкість тренування, $RMSE$ – корінь середньоквадратичної похибки, що описує точність класифікації, v – обсяг даних. Показ-

ник швидкодії можна визначити за допомогою сторонніх модулів чи програмного забезпечення після імплементації моделей. Наприклад, модуль *time* чи програмні системи *Postman*, *JMeter*.

Для визначення точності прогнозу скористаємось спеціальними вибірками з даними щодо російського вторгнення в Україну (DS-Pr1nce, 2022) та виборів президента Сполучених Штатів у 2020 році (Dedhia, 2022). Вказані дані будуть поділені за принципом Парето (Алесу, 2010, с. 4) на дві групи у співвідношенні 80 до 20. 80 % використовуватимуться у якості вхідних даних для здійснення прогнозу. 20 % натомість слугуватимуть реальним значенням. Після класифікації для знаходження кореня середньоквадратичної похибки достатньо скористатися наступною формулою:

$$RMSE = \sqrt{\frac{1}{N} \sum_{i=1}^N (y_i - \hat{y}_i)^2},$$

де N – кількість класифікованих значень; y_i – реальне значення показника; $\hat{y}_i$ – отримане значення показника. В якості обсягу даних розглядатимемо час, який витрачається користувачем для того, щоб підготувати модель до використання. При цьому, аби коректно порівняти дві моделі розглядатиметься обернений показник часу.

Правило порівняння ефективності. Задля можливості порівняння ефективності двох моделей введемо змінну C , що визначається наступною формулою:

```

1  def process_text(text, language):
2      stemmer = PorterStemmer()
3      lemmatizer = WordNetLemmatizer()
4      stop_words =
5          set(stopwords.words(
6              get_language_name(language)
7          ))
8      sentences = nltk.sent_tokenize(text)
9      sentences_processed = []
10     for sentence in sentences:
11         words_processed = []
12         words = nltk.word_tokenize(sentence)
13         words = [word for word in words if word not in stop_words]
14         for word in words:
15             word = stemmer.stem(word)
16             word = lemmatizer.lemmatize(word)
17             words_processed.append(word)
18         words =
19             [word for word in words_processed if word not in stop_words]
20         sentences_processed.append(" ".join(words))
21     text_processed = " ".join(sentences_processed)
22     return text_processed

```

Рис. 2. Код функції обробки

$$C_{AB} = \sum \text{normalize} \left(\frac{m_A}{m_B} \right),$$

де m_A – значення метрики для моделі A , m_B – значення метрики для моделі B .

Отримавши результати для кожної моделі, можна, використовуючи нечіткий параметр C , встановити яка модель є більш ефективною:

- $C_{AB} \geq 3.5$: модель A є більш ефективною за модель B ;
- $C_{AB} \leq 2.5$: модель A є менш ефективною за модель B ;
- $2.5 \leq C_{AB} \leq 3.5$: значення ефективності обох моделей є близьким один до одного.

Числові дані повинні бути зібрані для всіх метрик кожної моделі, тому використовується метод контрольованого експерименту. У наслідок цього було обрано стабільне та постійне середовище виконання – фізичний пристрій на базі *Ubuntu*. Для визначення швидкодії кожна модель імплементується за допомогою *Python 3*. Аби уникнути проблем при мануальному тестуванні та спростити процес налагодження прийнято рішення використати програмне середовище *Jupyter* та модуль *time* для оцінки швидкодії.

Як зазначалося раніше для визначення точності використовуються дві тестові вибірки. Ці дані розподіляються на дві частини у співвідношенні 80/20. Результат точності отриманий як співвідношення отриманого класу та реального для двох вибірок агрегується. Заключним показником для аналізу є швидкість підготовки даних. У цьому випадку показник залежить виключно від людини, що здійснює тестування.

Для цього експерименту можна визначити наступні помилки та невизначеності:

- при перевірці швидкості роботи варто очікувати похибки та помилки пов'язані з модулем оцінки часу виконання *time* та програмним середовищем *Jupyter*;
- при перевірці точності роботи моделей можуть виникнути проблеми із даними, що використовуються в якості тестової вибірки, оскільки вони прямо впливають на отриманий результат;
- при перевірці швидкості підготовки даних для аналізу можна виокремити дві основні проблеми: людський фактор та похибка інструменту вимірювання часу.

Задля зменшення впливу вказаних проблем було вирішено виконати 10 замірів швидкодії для звичайної та модифікованої моделей, по одному заміру для визначення точності результатів для кожної тестової вибірки та по 5 замірів для визначення часу підготовки до

аналізу. Для кожної метрики після проведення відповідних дій передбачається агрегація результатів

Аналіз результатів. Перед тим, як безпосередньо порівняти ефективність роботи двох алгоритмів, наведемо результати їх тренування – графіки, що зображають порівняння втрат (*loss*) з кількістю епох. Для CNN маємо наступні результати (рис. 3). Як видно з вказаного рисунку найкращою кількістю епох є трохи більше 1-єї.

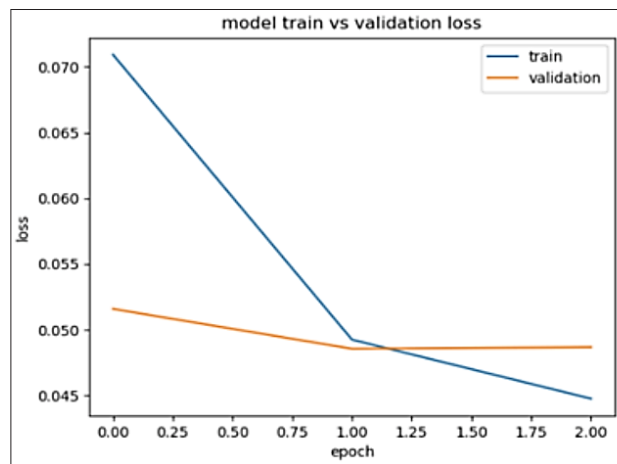


Рис. 3. Порівняння *loss* з кількістю епох для CNN

Перейдемо до наступної моделі – RNN (точніше LSTM). Для неї отримуємо наступний результат (рис. 4). Як бачимо у цьому разі оптимальна кількість епох рівна 1.

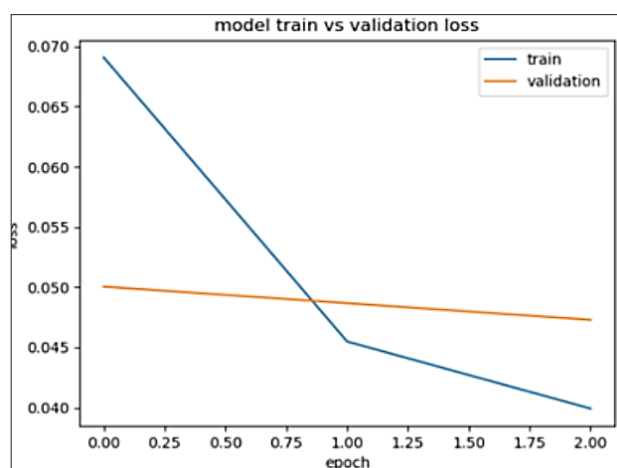


Рис. 4. Порівняння *loss* з кількістю епох для RNN

Перейдемо до визначення параметрів ефективності та почнемо із замірів швидкості тренування за допомогою бібліотеки *time*.

Результати наведено на рисунку 5. Знайдемо середнє значення для кожного з алгоритмів. Маємо для CNN значення часу близько 301 с, для RNN – 319 с.

Перейдемо до наступного показника – точність класифікації. Як зазначалося вище точність прогнозу вимірюється за допомогою показника *RMSE*. Маємо наступну точність отриманих результатів для обох вибірок: CNN (війна) – 93 %, CNN (вибори) – 94 %, RNN (війна) – 94 %, RNN (вибори) – 96 %.

Заключним показником є час підготовки даних для роботи архітектурних моделей. Результати замірів для даного показника наведені на рисунку 6.

Знайдемо середнє значення показника для кожного з алгоритмів. Маємо для базового алгоритму – 391 с, для модифікованого алгоритму – 391 с.

Знайшовши відповідні метрики для порівняння ефективності двох алгоритмів перейдемо до знаходження показника C_{AB} . У якості моделі *A* розглядатимемо CNN, а моделі *B* – RNN. Маємо $C_{AB} \approx 2.83$. Це дозволяє констатувати вищу ефективність RNN проти CNN, хоча ця відмінність неістотна. Даючи вигравш у швидкості тренування CNN програє у точності прогнозування. У випадку імплементації алгоритмів в рамках хмарної платформи, яка дозволяє значно пришвидшити швидкість роботи, отримана різниця в швидкодії між моделями буде неістотною. Однак варто зауважити, що точність може змінюватися в залежності від даних та їх обсягів.

Висновки. Дана робота мала на меті дослідження ефективності використання нейромережових архітектур типу CNN та RNN для виявлення фейкових новин. Задля цього була проаналізована предметна галузь пов'язана з неправдивою інформацією та виявлені ключові паттерни, які характеризують подібні дані. Окрім цього було проаналізовано теоретичне підґрунтя обох типів моделей. У ході чого виявлено, що стандартні моделі не можуть ефективно застосовуватися для текстової класифікації, тому було прийнято рішення для подальшого розгляду мати на увазі наступні положення:

- під CNN розуміти згорткову нейромережу з одношаровою згорткою і відповідним налаштуванням усіх її параметрів;
- під RNN розуміти налаштовану для аналізу тексту мережу LSTM, що має як довготривала, так і короткочасну пам'ять.

Наступним кроком стало визначення експериментального середовища, що включало себе:

- визначення функції ефективності;
- визначення правила порівняння ефективності двох моделей;
- план експерименту.

У якості тестових вибірок вирішено обрати дані, що стосується російського вторгнення на територію України та президентських виборів у Сполучених Штатах в 2020 році.

У результаті було виявлено, що на представлених даних модель CNN працює швидше за RNN, однак дає менш точний результат класифікації. З огляду на запропоноване правило порівняння ефективності, зазначені

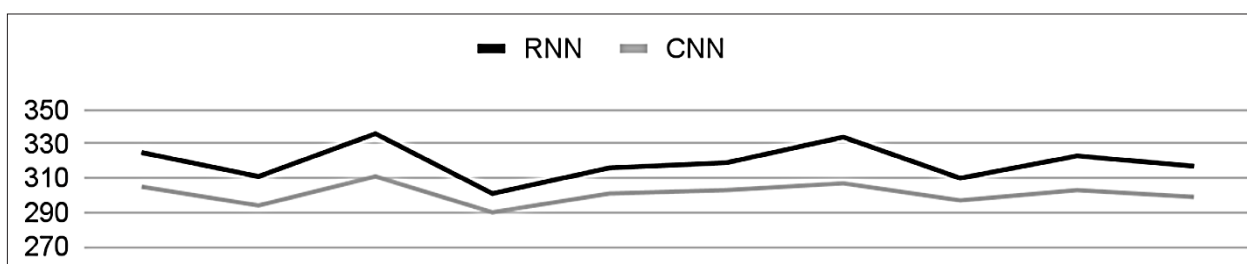


Рис. 5. Швидкість тренування нейромереж

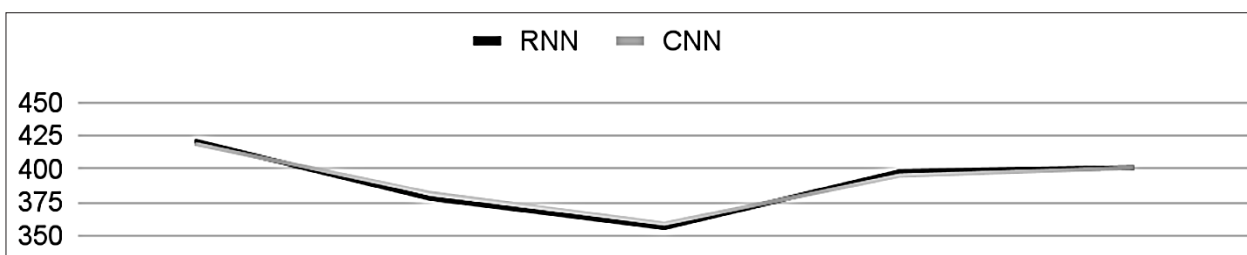


Рис. 6. Швидкість підготовки вхідних даних

імовірності похибок різного роду та шляхи подолання диспропорції між алгоритмами, отриманий вигравш ефективності можна вважати неістотним.

Зазначений висновок відповідає світовій науковій практиці, яка рекомендує під час

аналізу текстової інформації, особливо, за наявності двох класів (фейкові та не фейкові дані у чинному випадку), використовувати будь-яку із запропонованих моделей, або їх конгрегації у разі наявності перевірки зображень на достовірність (Saran, 2021).

ЛІТЕРАТУРА:

1. Fake News. *Cambridge Dictionary*. URL: <https://dictionary.cambridge.org/dictionary/english/fake-news> (дата звернення: 05.04.2023).
2. Wesolowski K. Fake news further fogs Russia's war on Ukraine. *Deutsche Welle*. URL: <https://www.dw.com/en/fact-check-fake-news-thrives-amid-russia-ukraine-war/a-61477502> (дата звернення: 08.04.2023).
3. Twitter steps up Ukraine misinformation fight. *BBC News*. URL: <https://www.bbc.com/news/business-61512261> (дата звернення: 08.04.2023).
4. Hybrid Warfare Analytical Group. How russian disinformation works in Telegram. *Ukraine Crisis Media Center*. URL: <https://uacrisis.org/en/ros-dezinfo-telegram> (дата звернення: 08.04.2023).
5. Blocman A. Laws to combat manipulation of information finally adopted. *IRIS Merlin*. URL: <https://merlin.obs.coe.int/article/8446> (дата звернення: 09.04.2023).
6. Danzig J. How fake news caused Brexit. *UACES*. URL: <https://europe.ideason europe.eu/2017/11/14/fake-news-caused-brexit/> (дата звернення: 09.04.2023).
7. Про внесення змін до статті 259 Кримінального кодексу України щодо посилення відповідальності за завідомо неправдиве повідомлення про загрозу безпеці громадян : Закон України від 02.03.2021 р. № 1292-IX. URL: <https://zakon.rada.gov.ua/laws/show/1292-20> (дата звернення: 09.04.2023).
8. Classifying Fake News Detection Using SVM, Naive Bayes and LSTM / P. Jain та ін. 2022 *12th International Conference on Cloud Computing, Data Science & Engineering*, м. Noida, 27–28 січ. 2022 р. URL: <https://doi.org/10.1109/Confluence52989.2022.9734129> (дата звернення: 20.04.2023).
9. A Study on Fake News Detection Using Naive Bayes, SVM, Neural Networks and LSTM / P. S. Reddy та ін. *Journal of Advanced Research in Dynamical and Control Systems*. 2019. Т. 11, № 06. С. 942–947.
10. Pai A. CNN vs. RNN vs. ANN – Analyzing 3 Types of Neural Networks in Deep Learning. *Analytics Vidhya*. URL: <https://www.analyticsvidhya.com/blog/2020/02/cnn-vs-rnn-vs-mlp-analyzing-3-types-of-neural-networks-in-deep-learning/> (дата звернення: 02.05.2023).
11. Stecanella B. Understanding TF-IDF: A Simple Introduction. *Monkey Learn*. URL: <https://monkeylearn.com/blog/what-is-tf-idf/> (дата звернення: 08.09.2023).
12. Pennington J., Socher R., Manning C. D. GloVe: Global Vectors for Word Representation. *Stanford Edu*. URL: <https://nlp.stanford.edu/projects/glove/> (дата звернення: 15.05.2023).
13. Risdal M. Getting Real about Fake News. *kaggle*. URL: <https://www.kaggle.com/datasets/mrisdal/fake-news> (дата звернення: 16.05.2023).
14. Lying words: Predicting deception from linguistic styles. / M. L. Newman та ін. *Personality and social psychology bulletin*. 2003. Т. 29, № 5. С. 665–675.
15. IBM Cloud Education. Natural Language Processing (NLP). *IBM Cloud Learn Hub*. URL: <https://www.ibm.com/cloud/learn/natural-language-processing> (дата звернення: 20.05.2023).
16. DS-Pr1nce. War in Ukraine: Russian social network discussions. *kaggle*. URL: <https://www.kaggle.com/datasets/ustyk5/war-in-ukraine-russian-social-network-discussions> (дата звернення: 25.05.2023).
17. Dedhia R. Fake News. *kaggle*. URL: <https://www.kaggle.com/datasets/ronikdedhia/fake-news> (дата звернення: 02.06.2023).
18. Alecu F. The Pareto Principle in the Modern Economy. *Economics of Knowledge*. 2010. Т. 2, № 3. С. 2–5.
19. Saran A. Text Classification – CNN with LSTM. *Medium*. URL: <https://anandsarank.medium.com/cnn-with-lstm-for-text-classification-53d18e5f7f5c> (дата звернення: 11.06.2023).

REFERENCES:

1. Fake News. (n.d.). Cambridge Dictionary. <https://dictionary.cambridge.org/dictionary/english/fake-news>
2. Wesolowski, K. (2022, April 17). Fake news further fogs Russia's war on Ukraine. dw.com. <https://www.dw.com/en/fact-check-fake-news-thrives-amid-russia-ukraine-war/a-61477502>
3. BBC News. (2022, May 19). Twitter steps up Ukraine misinformation fight. <https://www.bbc.com/news/business-61512261>

4. *How russian disinformation works in Telegram*. (n.d.). Uacrisis.org. <https://uacrisis.org/en/ros-dezinfo-telegram>
5. Blocman, A. (n.d.). Laws to combat manipulation of information finally adopted. IRIS Merlin. <https://merlin.obs.coe.int/article/8446>
6. Danzig, J. (2017, November 14). How fake news caused Brexit. UACES. <https://europe.ideasoneurope.eu/2017/11/14/fake-news-caused-brexite>
7. Pro vnesennia zmin do stat' 259 Kryminalnoho kodeksu Ukrainy shchodo posylennia vidpovidalnosti za zavidomo nepravdyve povidomlennia pro zahrozu bezpetsi hromadian, Zakon Ukrainy №. 1292-IX [On making changes to Article 259 of the Criminal Code of Ukraine regarding increased liability for knowingly false notification of a threat to the safety of citizens: Law of Ukraine No. 1292-IX] (2021). <https://zakon.rada.gov.ua/laws/show/1292-20#Text> [in Ukrainian].
8. Jain, P., Sharma, S., Monica, & Aggarwal, P. K. (2022). Classifying Fake News Detection Using SVM, Naive Bayes and LSTM. In *2022 12th International Conference on Cloud Computing, Data Science & Engineering (Confluence)*. IEEE. <https://doi.org/10.1109/confluence52989.2022.9734129>
9. Reddy, P. S. (2019). A Study on Fake News Detection Using Naive Bayes, SVM, Neural Networks and LSTM. *Journal of Advanced Research in Dynamical and Control Systems*, 11(6), 942–947.
10. Pai, A. (2020, February). *CNN vs. RNN vs. ANN – Analyzing 3 Types of Neural Networks in Deep Learning*. Analytics Vidhya. <https://www.analyticsvidhya.com/blog/2020/02/cnn-vs-rnn-vs-mlp-analyzing-3-types-of-neural-networks-in-deep-learning>
11. Stecanella, V. (n.d.). *Understanding TF-IDF: A Simple Introduction*. MonkeyLearn Blog. <https://monkeylearn.com/blog/what-is-tf-idf/>
12. Pennington, J., Socher, R., & Manning, C. D. (n.d.). *GloVe: Global Vectors for Word Representation*. Stanford Edu. <https://nlp.stanford.edu/projects/glove/>
13. Risdal, M. (n.d.). *Getting Real about Fake News*. kaggle. <https://www.kaggle.com/datasets/mrisdal/fake-news>
14. Newman, M. L. (2003). Lying words: Predicting deception from linguistic styles. *Personality and social psychology bulletin*, 29(5), 665–675.
15. IBM Cloud Education. (n.d.). *Natural Language Processing (NLP)*. IBM Cloud Learn Hub. <https://www.ibm.com/cloud/learn/natural-language-processing>
16. DS-Pr1nce. (n.d.). *War in Ukraine: Russian social network discussions*. kaggle. <https://www.kaggle.com/datasets/ustyk5/war-in-ukraine-russian-social-network-discussions>
17. Dedhia, R. (n.d.). *Fake News*. kaggle. <https://www.kaggle.com/datasets/ronikdedhia/fake-news>
18. Alecu, F. (2010). The Pareto Principle in the Modern Economy. *Economics of Knowledge*, 2(3), 2–5.
19. Saran, A. (n.d.). Text Classification – CNN with LSTM. Medium. <https://anandsarank.medium.com/cnn-with-lstm-for-text-classification-53d18e5f7f5c>

УДК 004.93

DOI <https://doi.org/10.32782/IT/2023-2-6>

Сергій ПРИХОДЬКО

доктор технічних наук, професор, завідувач кафедри програмного забезпечення автоматизованих систем, Національний університет кораблебудування імені адмірала Макарова, просп. Героїв України, 9, м. Миколаїв, Україна, 54007, sergiy.prykhodko@nuos.edu.ua

ORCID: 0000-0002-2325-018X

Scopus Author ID: 55225622100

Артем ТРУХОВ

аспірант кафедри програмного забезпечення автоматизованих систем, Національний університет кораблебудування імені адмірала Макарова, просп. Героїв України, 9, м. Миколаїв, Україна, 54007, timka357@gmail.com

ORCID: 0000-0002-7160-8609

Бібліографічний опис статті: Приходько, С., Трухов, А. (2023). Побудова правил прийняття рішень для розпізнавання облич на основі квадрата відстані Махаланобіса для нормалізованих даних. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 50–58, doi: <https://doi.org/10.32782/IT/2023-2-6>

ПОБУДОВА ПРАВИЛ ПРИЙНЯТТЯ РІШЕНЬ ДЛЯ РОЗПІЗНАВАННЯ ОБЛИЧ НА ОСНОВІ КВАДРАТА ВІДСТАНІ МАХАЛАНОБІСА ДЛЯ НОРМАЛІЗОВАНИХ ДАНИХ

Розпізнавання облич є одним із завдань розпізнавання образів, яке стає все більш популярним через широке застосування в комп'ютерному зорі, системах безпеки та ін. Низька ймовірність ідентифікації особи за обличчям може мати негативні наслідки. Тому існує потреба в розробці та вдосконаленні методів розпізнавання облич.

Один із широко використовуваних методів розпізнавання образів базується на застосуванні вирішальних правил на основі квадрата відстані Махаланобіса. Квадрат відстані Махаланобіса використовується для побудови еліпсоїда прогнозування. Але суттєвим обмеженням його використання є необхідність виконання припущення про нормальність розподілу багатовимірних даних, порушення якого як правило призводить до зменшення ймовірності розпізнавання.

Метою роботи є підвищення ймовірності розпізнавання облич за рахунок побудови правил прийняття рішень на основі квадрата відстані Махаланобіса для десятивимірних нормалізованих даних характеристик обличчя.

Для отримання вектору характеристик обличчя з використанням бібліотеки Dlib було розроблено програму мовою Python. Для оцінювання відхилення від нормального розподілу даних було використано тест Mardia. Досліджено, що отримані вибірки характеристик обличчя мають розподіл, який відхиляється від нормального, тому була виконана нормалізація за допомогою відомого одновимірного перетворення у вигляді десяткового логарифму. На основі квадрата відстані Махаланобіса побудовані правила прийняття рішень у вигляді десятивимірних еліпсоїдів прогнозування для початкових та нормалізованих даних. Вирішальні правила, побудовані для нормалізованих даних, показали більшу ймовірність розпізнавання облич.

Результати доводять, що нормалізація збільшує ймовірність розпізнавання облич у разі суттєвого відхилення багатовимірного розподілу характеристик обличчя від нормального. Також було з'ясовано, що у разі існування високої кореляції між характеристиками обличчя, застосування одновимірних нормалізуючих перетворень не завжди призводить до добрих результатів розпізнавання. У цьому разі потрібно використовувати багатовимірні нормалізуючі перетворення, такі як Бокса-Кокса або Джонсона.

Ключові слова: розпізнавання облич, відстань Махаланобіса, еліпсоїд прогнозування, нормалізація, Dlib, багатовимірний нормальний розподіл.

Sergiy PRYKHODKO

Dr.Sc., Professor, Head of the Department of Software for Automated Systems, Admiral Makarov National University of Shipbuilding, Heroiv Ukrainy ave., 9, Mykolaiv, Ukraine, 54007, sergiy.prykhodko@nuos.edu.ua

ORCID: 0000-0002-2325-018X

Scopus Author ID: 55225622100

Artem TRUKHOV

Postgraduate Student at the Department of Software for Automated Systems, Admiral Makarov National University of Shipbuilding, Heroiv Ukrainy ave., 9, Mykolaiv, Ukraine, 54007, timka357@gmail.com

ORCID: 0000-0002-7160-8609

To cite this article: Prykhodko, S., Trukhov, A. (2023). Pobudova pravyl pryiniattia rishen dlia rozpoznavannia oblych na osnovi kvadrata vidstani Mahalanobisa dlia normalizovanykh danykh [Construction of decision rules for face recognition based on the squared Mahalanobis distance for normalized data]. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 50–58, doi: <https://doi.org/10.32782/IT/2023-2-6>

CONSTRUCTION OF DECISION RULES FOR FACE RECOGNITION BASED ON THE SQUARED MAHALANOBIS DISTANCE FOR NORMALIZED DATA

Face recognition is one of the tasks of pattern recognition, which is becoming more and more popular due to its wide application in computer vision, security systems, etc. The low probability of identifying a person by face can have negative consequences. Therefore, there is a need to develop and improve face recognition methods.

One of the widely used pattern recognition methods is based on the application of decision rules based on the squared Mahalanobis distance. The squared Mahalanobis distance is used to construct the prediction ellipsoid. But a significant limitation of its use is the need to fulfill the assumption about the normality of the distribution of multidimensional data, the violation of which usually leads to a decrease in the probability of recognition.

The aim of the work is to increase the probability of face recognition by building decision rules based on the squared Mahalanobis distance for ten-dimensional normalized data of face characteristics.

A Python program was developed to obtain a vector of face characteristics using the Dlib library. The Mardia test was used to assess the deviation from the normal distribution of the data. It was investigated that the received samples of facial characteristics have a distribution that deviates from normal, so normalization was performed using the well-known one-dimensional transformation in the form of a decimal logarithm. Based on the squared Mahalanobis distance, decision rules are built in the form of ten-dimensional prediction ellipsoids for initial and normalized data. Decision rules built for normalized data showed a higher probability of recognizing faces.

The results prove that normalization increases the probability of face recognition in case of significant deviation of the multidimensional distribution of face characteristics from normal. It was also found that in the case of high correlation between facial features, the application of one-dimensional normalizing transformations does not always lead to good recognition results. In this case, it is necessary to use multivariate normalizing transformations, such as Box-Cox or Johnson.

Key words: face recognition, Mahalanobis distance, prediction ellipsoid, normalization, Dlib, multivariate normal distribution.

Вступ. Розпізнавання облич є одним із завдань розпізнавання образів, яке стає все більш популярним через широке застосування в комп'ютерному зорі, системах безпеки та ін. Процес розпізнавання обличчя передбачає автоматичну ідентифікацію людини за відмінними рисами, включаючи форму очей, носа, рота та інші характеристики (Beham, 2013).

Ймовірність розпізнавання обличчя за його характеристиками суттєво залежить від правила прийняття рішення, або вирішального правила, яке застосовується. Правила прийняття рішень є математичними виразами або алгоритмами, які визначають, до якого класу розпізнавання належить образ, застосовуючи заданий алфавіт (Довбиш, 2015).

Статистичні методи розпізнавання образів, як правило, використовують вирішальні правила, які базуються на припущенні про нормальність розподілу (Webb, 2011; Sen, 2015). Реальні дані, як правило, мають багатовимірний розподіл, який відхиляється від нормального, що призводить до помилкових результатів розпізнавання.

Актуальність роботи. Технологія розпізнавання облич стає все більш важливою в різних сферах. Низька ймовірність ідентифікації осіб може мати негативні наслідки з точки зору безпеки, зручності та персоналізації. Тому існує потреба в розробці та вдосконаленні методів розпізнавання облич.

Аналіз останніх досліджень. Статистичні методи розпізнавання образів проводять

аналіз характеристик об'єкту, які описуються векторами для створення статистичної моделі обличчя. Для побудови вирішальних правил існують різні підходи, однак один із найпоширеніших методів базується на використанні відстані Махаланобіса (Nader, 2014; Hamill, 2016), значення квадрату якого для i -ої точки даних обчислюється за наступною формулою:

$$D^2(x_i) = (x_i - \bar{x})^T S^{-1} (x_i - \bar{x}), \quad (1)$$

де $\bar{x}$ – вектор вибірових середніх;

x_i – вектор даних в i -ій точці;

S – вибіркова коваріаційна матриця даних.

Квадрат відстані Махаланобіса має широке застосування в статистиці та аналізі багатовимірних даних для розпізнавання, класифікації і виявлення викидів. Квадрат відстані Махаланобіса приблизно відповідає розподілу Хі-квадрат з k ступенями свободи, що дорівнює числу характеристик, та враховує кореляцію між змінними, тобто, враховується розкид даних відносно центру (Etherington, 2021; Brereton, 2014). Це можна використовувати для побудови еліпсоїда прогнозування, який описує простір допустимих значень (Ghorbani, 2019), таким чином усі спостереження одного класу повинні лежати в межах цього еліпсоїда, а іншого класу – поза межами. На основі (1) правило прийняття рішення набуває вигляду (Brereton 2021; Van Aelst, 2009):

$$D^2(x_i) \leq \chi^2(q, k), \quad (2)$$

де q – рівень значущості;

k – кількість ступенів свободи.

Побудова еліпсоїда базується на припущенні про багатовимірний нормальний розподіл даних (Gallego, 2013). Однак реальні дані можуть мати відмінний від нормального багатовимірний розподіл, що призводить до зниження ймовірності розпізнавання. Для вирішення певних практичних задач, які базуються на використанні відстані Махаланобіса, у випадку негаусівських даних застосовується нормалізація (Meshkani, 2011; Prykhodko, 2018; Prykhodko, 2020). Її застосування дозволяє здійснювати розв'язання відповідних задач для даних, багатовимірний розподіл яких відхиляється від нормального.

Метою роботи є підвищення ймовірності розпізнавання облич за рахунок побудови правил прийняття рішень на основі квадрата відстані Махаланобіса для десятивимірних нормалізованих даних характеристик обличчя.

Виклад основного матеріалу дослідження. Вибір набору даних має вирішальне

значення для побудови правил прийняття рішень. Для дослідження було обрано відомий набір даних Pins Face Recognition (Burak, 2019), який використовується при будові вирішальних правил (Saib, 2021). З метою забезпечення надійності та точності дослідження, було проведено фільтрацію набору даних, після чого здійснено додаткове завантаження фотографій з відкритих ресурсів та розширення набору. Всього було отримано 400 високоякісних фотографій, по 200 для кожного з двох облич.

Формування вектору характеристик.

Одним з найважливіших завдань розпізнавання обличчя є пошук ключових точок, таких як очі, ніс та рот. Коли орієнтири виявлені, їх можна використовувати для обчислення відстаней між різними частинами обличчя. Потім на основі цих відстаней формуються характеристики для представлення об'єкту, які можна порівняти з характеристиками іншого обличчя для розпізнавання (Harry, 2014).

Для створення векторів характеристик була створена програма мовою Python з використанням бібліотеки комп'ютерного зору Dlib. Dlib – це потужний інструмент з відкритим вихідним кодом, яка надає набір методів для широкого спектру завдань комп'ютерного зору, включаючи виявлення обличчя, пошук ключових точок та виділення характеристик. Бібліотека використовує надійний набір із 68 ключових точок для точного представлення контурів об'єкту (Boyo, 2018; Khabaralak, 2021).

Алгоритм роботи програми складається з кількох кроків, які виконуються послідовно для досягнення бажаних результатів. Спочатку програма знаходить обличчя на вхідному зображенні за допомогою алгоритму пошуку Dlib. На другому кроці видаляється задній фон, а лице вирівнюється таким чином, щоб очі були на одному рівні. Така обробка допомагає усунути деякі спотворення, викликані положенням обличчя на вхідному зображенні (Haghpanah, 2022; Amato, 2018). На третьому кроці програма отримує набір характеристик із вирівняного зображення. Кожна риса – це відстань між орієнтирами на обличчі, визначеними бібліотекою Dlib.

Формуючи вектор характеристик необхідно керуватися кінцевими цілями та обмеженнями інструментів, що використовуються. Зазвичай бібліотеки комп'ютерного зору мають свої власні недоліки. У даному випадку, використання Dlib для розпізнавання країв обличчя може призвести до спотворень в векторі характеристик, особливо коли голова повернута. Проаналізувавши наявні роботи (Osman, 2015; Bijarnia, 2016; Juhong, 2017; Li, 2018; Gaber, 2021), було

вирішено зосередитися на 17 найбільш інформативних ключових точках, що використовуються для вимірювання 15 різних відстаней обличчя. На основі цих відстаней було створено вектор, який містить 10 характеристик (рис. 1).

Для мінімізації впливу положення обличчя на зображенні та його віддаленості використовується нормування за допомогою ділення кожної характеристики вектору на базову відстань між очима (Johnston, 2018; Wu, 2018).

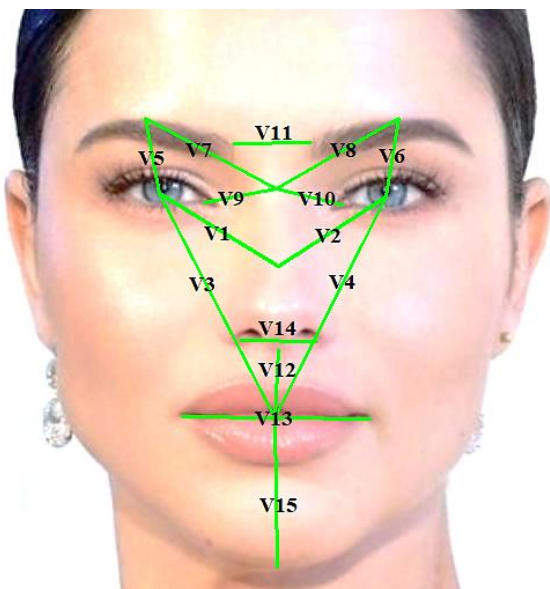


Рис. 1. Відстані між ключовими точками обличчя, які використовуються для розпізнавання

Вектор характеристик складається з 10 змінних, при тому симетричні дистанції усереднюються (Nestor, 2012) (таблиця 1).

Таблиця 1

Опис вектору характеристик

№ характеристики	Формула	Опис
1	$(v1 + v2) / 2$	Усереднена відстань від очей до середини носа
2	$(v3 + v4) / 2$	Усереднена відстань від очей до центру рота
3	$(v5 + v6) / 2$	Усереднена відстань від очей до центру брів
4	$(v7 + v8) / 2$	Усереднена відстань від брів до верху носа
5	$(v9 + v10) / 2$	Усереднена відстань від куточків очей до верху носа
6	$v11$	Відстань між бровами
7	$v12$	Відстань між носом і серединою рота
8	$v13$	Ширина рота
9	$v14$	Ширина носу
10	$v15$	Відстань від рота до підборіддя

Характеристики отриманих вибірок. Для кожної людини було використано по 200 фотографій, з яких 100 використовувалися для побудови правил прийняття рішень, а ще 100 – для тестування та оцінювання якості розпізнавання. Після отримання векторів характеристик було сформовано по одній навчальній вибірці для кожної особи.

Вектор вибірових середніх значень першої людини для побудови вирішальних правил складається з наступних елементів $\bar{x} = [0,6269; 1,2482; 0,3175; 0,6744; 0,3004; 0,3921; 0,3571; 0,8588; 0,4196; 0,7886]$. У таблиці 2 представлено коваріаційну матрицю вибірки першої людини, таблиця 3 містить діапазони значень характеристик.

Таблиця 2

Коваріаційна матриця вибірки першої людини

0,0016	0,0013	-0,0006	0,0001	-0,0001	0,0006	-0,0005	-0,0003	0,0001	-0,0016
0,0013	0,0030	-0,0002	0,0002	-0,0001	0,0002	0,0009	-0,0012	0,0001	0,0001
-0,0006	-0,0002	0,0013	0,0007	0,0000	0,0008	0,0004	0,0004	-0,0001	0,0007
0,0001	0,0002	0,0007	0,0008	0,0000	0,0012	0,0001	0,0004	-0,0001	0,0002
-0,0001	-0,0001	0,0000	0,0000	0,0001	-0,0001	0,0000	-0,0001	0,0000	0,0001
0,0006	0,0002	0,0008	0,0012	-0,0001	0,0027	-0,0003	0,0005	-0,0003	-0,0004
-0,0005	0,0009	0,0004	0,0001	0,0000	-0,0003	0,0013	-0,0007	-0,0002	0,0015
-0,0003	-0,0012	0,0004	0,0004	-0,0001	0,0005	-0,0007	0,0054	0,0010	0,0016
0,0001	0,0001	-0,0001	-0,0001	0,0000	-0,0003	-0,0002	0,0010	0,0007	0,0000
-0,0016	0,0001	0,0007	0,0002	0,0001	-0,0004	0,0015	0,0016	0,0000	0,0055

Таблиця 3

Діапазони характеристик вибірки першої людини

№	1	2	3	4	5	6	7	8	9	10
Min	0,5467	1,0959	0,2497	0,6123	0,2800	0,2733	0,2772	0,7142	0,3672	0,6209
Max	0,7049	1,4012	0,4268	0,7411	0,3271	0,5064	0,4868	1,0782	0,4835	0,9104

Вектор вибірових середніх значень другої людини для побудови вирішальних правил дорівнює $\bar{x} = [0,6330; 1,1629; 0,3010; 0,6259; 0,2994; 0,3670; 0,3015; 0,8011; 0,3584; 0,6290]$. Коваріаційна матриця вибірки другої людини зображена в таблиці 4, діапазони характеристик – в таблиці 5.

Оцінювання багатовимірного відхилення розподілу даних від нормального. Для оцінювання відхилення багатовимірного розподілу даних від нормального використано тест Мардіа. Він заснований на аналізі багатовимірних асиметрії B_1 та ексцесу B_2 даних (Mardia, 1970), які є показниками того, наскільки дані відхиляються від нормального розподілу та обчислюються за наступними формулами:

$$B_1 = \frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n \{(\mathbf{x}_i - \bar{\mathbf{x}})^T \mathbf{S}^{-1} (\mathbf{x}_j - \bar{\mathbf{x}})\}^3, \quad (3)$$

$$B_2 = \frac{1}{n} \sum_{j=1}^n \{(\mathbf{x}_j - \bar{\mathbf{x}})^T \mathbf{S}^{-1} (\mathbf{x}_j - \bar{\mathbf{x}})\}^2. \quad (4)$$

Величина асиметрії, розрахована за допомогою формули (3), помножена на $n/6$, розподілена за законом Хі-квадрат з $p(p+1)(p+2)/6$ ступенями свободи, де p – це кількість змінних, в даному випадку 10, а величина ексцесу, розрахована за допомогою формули (4), розподілена згідно із законом нормального розподілу із математичним сподіванням $p(p+2)$ і дисперсією $8p(p+2)/n$ (Mardia, 1974).

Відповідно до тесту Мардіа багатовимірні розподіли даних вибірок відхиляються від нормального закону, тому що значення тестових статистик розподілів вибірок першої та другої людини для багатовимірної асиметрії

$B_1 n/6$ 279,98 та 289,20 перевищують значення квантиля розподілу Хі-квадрат, що становить 277,77 для 220 ступенів свободи та рівня значущості 0,005. Натомість, тестові статистики розподілів вибірок першої та другої людини для багатовимірного ексцесу B_2 122,46 та 122,35 не перевищують значення квантиля розподілу Гауса, яке становить 127,97 для математичного сподівання 120 і дисперсії 9,6 та рівня значущості 0,005.

На основі (2) вирішальне правило для початкових даних, за допомогою якого виконується розпізнавання облич та віднесення їх до класу 1 або 2, де значення квантиля розподілу Хі-квадрат становить 25,19 для 10 ступенів свободи та рівня значущості 0,005, матиме вигляд:

$$(\mathbf{x}_i - \bar{\mathbf{x}})^T \mathbf{S}^{-1} (\mathbf{x}_i - \bar{\mathbf{x}}) \leq \chi^2(10, 0,995), \quad (5)$$

Для перевірки якості розпізнавання вирішального правила, побудованого для кожної вибірки, тестовий набір складається зі 100 тестових векторів характеристик даної вибірки та 200 векторів характеристик іншої вибірки. Для вирішального правила (5), побудованого на основі даних вибірки першої людини, відсоток помилок першого роду становить 4 %, а другого – 12,5 %. Для вирішального правила (5), побудованого на основі вибірки другої людини, відсоток помилок першого роду становить 8 %, а другого – 4,5 %.

Така ймовірність помилок розпізнавання обумовлена тим, що використання квадрата відстані Махаланобіса базується на припущенні про нормальність розподілу даних. Розподіл обраних вибірок відхиляється від нормального, відповідно виникає необхідність використання нормалізуючого перетворення.

Таблиця 4

Коваріаційна матриця вибірки другої людини

0,0010	0,0009	-0,0006	-0,0001	-0,0001	-0,0001	-0,0005	0,0000	0,0003	-0,0011
0,0009	0,0020	0,0000	0,0003	-0,0002	0,0001	0,0005	0,0000	0,0005	0,0002
-0,0006	0,0000	0,0011	0,0005	0,0000	0,0002	0,0007	-0,0001	-0,0001	0,0012
-0,0001	0,0003	0,0005	0,0004	0,0000	0,0002	0,0003	-0,0003	-0,0001	0,0003
-0,0001	-0,0002	0,0000	0,0000	0,0001	0,0000	0,0000	0,0000	-0,0001	0,0001
-0,0001	0,0001	0,0002	0,0002	0,0000	0,0007	0,0003	0,0000	0,0000	0,0000
-0,0005	0,0005	0,0007	0,0003	0,0000	0,0003	0,0012	-0,0002	-0,0001	0,0015
0,0000	0,0000	-0,0001	-0,0003	0,0000	0,0000	-0,0002	0,0028	0,0010	0,0013
0,0003	0,0005	-0,0001	-0,0001	-0,0001	0,0000	-0,0001	0,0010	0,0010	0,0002
-0,0011	0,0002	0,0012	0,0003	0,0001	0,0000	0,0015	0,0013	0,0002	0,0042

Таблиця 5

Діапазони характеристик вибірки другої людини

	1	2	3	4	5	6	7	8	9	10
Min	0,5573	1,0703	0,2400	0,5778	0,2729	0,3109	0,2261	0,6987	0,2933	0,4681
Max	0,7085	1,2873	0,3958	0,6698	0,3217	0,4316	0,3996	1,0007	0,4604	0,8095

Нормалізація даних була здійснена за допомогою перетворення у вигляді десятичного логарифму, вектор нормалізованих величин $\mathbf{z}$ визначається наступним чином:

$$\mathbf{z} = \lg(\mathbf{x}). \quad (6)$$

В результаті застосування перетворення (6) нормалізовані дані не відхиляються від багатовимірного нормального розподілу, адже для нормалізованих даних значення тестових статистик розподілів вибірок першої та другої людини для багатовимірної асиметрії B_1 277,45 та 277,05 не перевищують критичного значення 277,77; тестові статистики розподілів вибірок першої та другої людини для багатовимірного ексцесу B_2 120,95 та 121,36 менше за критичне значення 127,97.

На основі (2) вирішальне правило для нормалізованих даних матиме вигляд:

$$(\mathbf{z}_i - \bar{\mathbf{z}})^T \mathbf{S}_z^{-1} (\mathbf{z}_i - \bar{\mathbf{z}}) \leq \chi^2(10, 0,995), \quad (7)$$

де $\bar{\mathbf{z}}$ – вектор вибірових середніх нормалізованих даних;

$\mathbf{S}_z$ – вибіркова коваріаційна матриця нормалізованих даних.

Використання правила прийняття рішення (7) для нормалізованих даних першої людини

дозволило знизити відсоток помилок першого роду з 4 % до 2 %, другого роду – з 12,5 % до 3,5 %. Застосування вирішального правила (7) для нормалізованих даних другої людини дало можливість знизити рівень помилок першого роду з 8 % до 5 %, але, на жаль, помилки другого роду зросли на 1 %. Це пояснюється тим, що одновимірне перетворення (6) не враховує кореляцію між характеристиками та передбачає, що змінні є незалежними (Fischer, 2016). Вибіркові коваріаційні матриці $\mathbf{S}_z$ для правил прийняття рішень для нормалізованих даних (7) наведені в таблицях 6 та 7.

Висновки. Удосконалено правило прийняття рішення для розпізнавання облич на основі застосування квадрата відстані Махаланобіса для нормалізованих даних з 10 характеристик за допомогою перетворення десятичного логарифма, що дозволяє підвищити ймовірність відповідного розпізнавання у разі задовільної нормалізації.

В подальшому планується побудова правил прийняття рішень для розпізнавання облич із застосуванням квадрату відстані Махаланобіса для нормалізованих даних на основі багатовимірних перетворень, які дозволять врахувати кореляцію між характеристиками, за рахунок чого підвищити ймовірність розпізнавання.

Таблиця 6

Коваріаційна матриця нормалізованої вибірки першої людини

0,0007	0,0003	-0,0006	0,0000	-0,0001	0,0004	-0,0005	-0,0001	0,0001	-0,0006
0,0003	0,0004	-0,0001	0,0000	-0,0001	0,0001	0,0004	-0,0002	0,0000	0,0000
-0,0006	-0,0001	0,0025	0,0007	0,0000	0,0012	0,0007	0,0002	-0,0002	0,0006
0,0000	0,0000	0,0007	0,0003	0,0000	0,0008	0,0001	0,0001	-0,0001	0,0001
-0,0001	-0,0001	0,0000	0,0000	0,0001	-0,0001	0,0000	-0,0001	0,0000	0,0001
0,0004	0,0001	0,0012	0,0008	-0,0001	0,0034	-0,0005	0,0003	-0,0004	-0,0002
-0,0005	0,0004	0,0007	0,0001	0,0000	-0,0005	0,0019	-0,0005	-0,0003	0,0011
-0,0001	-0,0002	0,0002	0,0001	-0,0001	0,0003	-0,0005	0,0013	0,0005	0,0004
0,0001	0,0000	-0,0002	-0,0001	0,0000	-0,0004	-0,0003	0,0005	0,0007	0,0000
-0,0006	0,0000	0,0006	0,0001	0,0001	-0,0002	0,0011	0,0004	0,0000	0,0018

Таблиця 7

Коваріаційна матриця нормалізованої вибірки другої людини

0.0005	0.0002	-0.0006	0.0000	-0.0001	-0.0001	-0.0005	0.0000	0.0003	-0.0005
0.0002	0.0003	0.0000	0.0001	-0.0001	0.0000	0.0003	0.0000	0.0002	0.0001
-0.0006	0.0000	0.0023	0.0005	0.0000	0.0003	0.0014	-0.0001	-0.0002	0.0012
0.0000	0.0001	0.0005	0.0002	0.0000	0.0001	0.0003	-0.0001	-0.0001	0.0002
-0.0001	-0.0001	0.0000	0.0000	0.0002	0.0000	0.0001	0.0000	-0.0002	0.0001
-0.0001	0.0000	0.0003	0.0001	0.0000	0.0010	0.0004	0.0000	0.0000	0.0000
-0.0005	0.0003	0.0014	0.0003	0.0001	0.0004	0.0024	-0.0001	-0.0002	0.0015
0.0000	0.0000	-0.0001	-0.0001	0.0000	0.0000	-0.0001	0.0008	0.0006	0.0005
0.0003	0.0002	-0.0002	-0.0001	-0.0002	0.0000	-0.0002	0.0006	0.0015	0.0001
-0.0005	0.0001	0.0012	0.0002	0.0001	0.0000	0.0015	0.0005	0.0001	0.0020

ЛІТЕРАТУРА:

1. Довбиш, А. С., І. В. Шелехов. Основи теорії розпізнавання образів : навч. посіб. : у 2-х ч. Суми : Сумський державний університет, 2015. 109 с.
2. Amato, Giuseppe, et al. A comparison of face verification with facial landmarks and deep features. 10th International Conference on Advances in Multimedia (MMEDIA). 2018.
3. Beham, M. Parisa, and S. Mohamed Mansoor Roomi. A review of face recognition methods. International Journal of Pattern Recognition and Artificial Intelligence, 2013.
4. Bijarnia, Saroj, and Preeti Singh. Age invariant face recognition using minimal geometrical facial features. Advanced Computing and Communication Technologies: Proceedings of the 9th ICACCT, 2015. Springer Singapore, 2016.
5. Boyko, Nataliya, Oleg Basystiuk, and Nataliya Shakhovska. Performance evaluation and comparison of software for face recognition, based on dlib and opencv library. 2018 IEEE Second International Conference on Data Stream Mining & Processing (DSMP). IEEE, 2018.
6. Brereton, Richard G. The chi squared and multinormal distributions. *Journal of Chemometrics* 29.1, 2015 : 9–12.
7. Brereton, Richard G. Multivariate classification models. *Journal of Chemometrics* 35.12, 2021 : e3332.
8. Burak. Dataset Pins Face Recognition. 2019. URL: <https://www.kaggle.com/datasets/hereisburak/pins-face-recognition> (дата звернення: 08.02.2023).
9. Etherington, Thomas R. Mahalanobis distances for ecological niche modelling and outlier detection: implications of sample size, error, and bias for selecting and parameterising a multivariate location and scatter method. *PeerJ* 9, 2021 : e11436.
10. Fischer, Christoph. Comparing the logarithmic transformation and the box-cox transformation for individual tree basal area increment models. *Forest Science*, 2016.
11. Gaber, Amira, et al. SVM classification of facial functions based on facial landmarks and animation Units. *Biomedical Physics & Engineering Express* 7.5, 2021.
12. Gallego, Guillermo, et al. On the Mahalanobis distance classification criterion for multidimensional normal distributions. IEEE Transactions on Signal Processing, 2013.
13. Ghorbani, Hamid. Mahalanobis distance and its application for detecting multivariate outliers. Facta Universitatis, Series: Mathematics and Informatics, 2019.
14. Haghpahan, Mohammad A., et al. Real-time facial expression recognition using facial landmarks and neural networks. 2022 International Conference on Machine Vision and Image Processing (MVIP). IEEE, 2022.
15. Hamill, Patrick, et al. An AERONET-based aerosol classification using the Mahalanobis distance. *Atmospheric Environment* 140, 2016: 213-233.
16. Happy, S. L., and Aurobinda Routray. Automatic facial expression recognition using features of salient facial patches. *IEEE transactions on Affective Computing* 6.1, 2014: 1–12.
17. Johnston, Benjamin, and Philip de Chazal. A review of image-based automatic facial landmark identification techniques. *EURASIP Journal on Image and Video Processing* 2018.1, 2018: 1–23.
18. Juhong, Aniwat, and Chuchart Pintavirooj. Face recognition based on facial landmark detection. 2017 10th Biomedical Engineering International Conference (BMEiCON). IEEE, 2017.
19. Khabaralak, Kostiantyn, and Larysa Koriashkina. Fast facial landmark detection and applications: A survey. arXiv preprint arXiv:2101.10808, 2021.
20. Li, Dahua, et al. Facial expression recognition based on electroencephalogram and facial landmark localization. *Technology and Health Care* 27.4, 2019: 373–387.
21. Mardia, Kanti V. Measures of multivariate skewness and kurtosis with applications. *Biometrika*, 1970, 57.3: 519–530.
22. Mardia, Kanti V. Applications of some measures of multivariate skewness and kurtosis in testing normality and robustness studies. *Sankhyā: The Indian Journal of Statistics, Series B*, 1974, 115-128.
23. Meshkani, Seyed Ahmad, et al. The application of geochemical pattern recognition to regional prospecting: A case study of the Sanandaj–Sirjan metallogenic zone, Iran. *Journal of geochemical exploration* 108.3, 2011: 183–195.
24. Nader, Patric, Paul Honeine, and Pierre Beuseroy. Mahalanobis-based one-class classification. 2014 IEEE International Workshop on Machine Learning for Signal Processing (MLSP). IEEE, 2014.
25. Nestor, Adrian, Jean M. Vettel, and Michael J. Tarr. Internal representations for face detection: An application of noise-based image classification to BOLD responses. *Human brain mapping* 34.11, 2013: 3101–3115.

26. Osman Ali, Amal Seralkhatem, et al. Age-invariant face recognition using triangle geometric features. *International Journal of Pattern Recognition and Artificial Intelligence* 29.05, 2015: 1556006.
27. S. Prykhodko, N. Prykhodko, L. Makarova and A. Pukhalevych, Application of the squared mahalanobis distance for detecting outliers in multivariate non-Gaussian data, 2018 14th International Conference on Advanced Trends in Radioelectronics, Telecommunications and Computer Engineering (TCSET), Lviv-Slavske, Ukraine, 2018, pp. 962–965, doi: 10.1109/TCSET.2018.8336353.
28. Prykhodko, Sergiy, et al. Application of transformed prediction ellipsoids for outlier detection in multivariate non-Gaussian data. 2020 IEEE 15th International Conference on Advanced Trends in Radioelectronics, Telecommunications and Computer Engineering (TCSET). IEEE, 2020.
29. Saib, Yaaseen Muhammad, and Sameerchand Pudaruth. Is Face Recognition with Masks Possible?. *International Journal of Advanced Computer Science and Applications* 12.7, 2021.
30. Sen, Sumen, Norou Diawara, and Khan M. Iftekharruddin. Statistical pattern recognition using Gaussian copula. *Journal of Statistical Theory and Practice* 9, 2015: 768–777.
31. Van Aelst, Stefan, and Peter Rousseeuw. Minimum volume ellipsoid. *Wiley Interdisciplinary Reviews. Computational Statistics* 1.1, 2009: 71–82.
32. Andrew R. Webb, Keith D. Copsey. *Statistical Pattern Recognition*. Third Edition. John Wiley & Sons, 2011.
33. Wu, Yue, and Qiang Ji. Facial landmark detection: A literature survey. *International Journal of Computer Vision* 127, 2019: 115–142.

REFERENCES:

1. Dovbysh, A. S., & Shelekhov, I. V. (2015). *Osnovy teorii rozpoznavannia obraziv: navch. posib.: u 2-kh ch [Basics of pattern recognition theory: teaching. manual: in the 2nd part]*. Sumy: Sumskyi derzhavnyi universytet [in Ukrainian].
2. Amato, G., Falchi, F., Gennaro, C., & Vairo, C. (2018, April). A comparison of face verification with facial landmarks and deep features. In 10th International Conference on Advances in Multimedia (MMEDIA) (pp. 1–6).
3. Beham, M. P., & Roomi, S. M. M. (2013). A review of face recognition methods. *International Journal of Pattern Recognition and Artificial Intelligence*.
4. Bijarnia, S., & Singh, P. (2016). Age invariant face recognition using minimal geometrical facial features. In *Advanced Computing and Communication Technologies: Proceedings of the 9th ICACCT, 2015* (pp. 71–77). Springer Singapore.
5. Boyko, N., Basystiuk, O., & Shakhovska, N. (2018, August). Performance evaluation and comparison of software for face recognition, based on dlib and opencv library. In 2018 IEEE Second International Conference on Data Stream Mining & Processing (DSMP) (pp. 478–482). IEEE.
6. Brereton, R. G. (2015). The chi squared and multinormal distributions. *Journal of Chemometrics*, 29(1), (pp 9–12).
7. Brereton, R. G. (2021). Multivariate classification models. *Journal of Chemometrics*, 35(12), e3332.
8. Burak. (2019). Pins Face Recognition. URL: <https://www.kaggle.com/datasets/hereisburak/pins-face-recognition> (available: 08.02.2023).
9. Etherington, T. R. (2021). Mahalanobis distances for ecological niche modelling and outlier detection: implications of sample size, error, and bias for selecting and parameterising a multivariate location and scatter method. *PeerJ*, 9, e11436.
10. Fischer, C. (2016). Comparing the logarithmic transformation and the box-cox transformation for individual tree basal area increment models. *Forest Science*, 62(3), 297–306.
11. Gaber, A., Taher, M. F., Wahed, M. A., & Shalaby, N. M. (2021). SVM classification of facial functions based on facial landmarks and animation Units. *Biomedical Physics & Engineering Express*, 7(5), 055008.
12. Gallego, G., Cuevas, C., Mohedano, R., & Garcia, N. (2013). On the Mahalanobis distance classification criterion for multidimensional normal distributions. *IEEE Transactions on Signal Processing*, 61(17), 4387–4396.
13. Ghorbani, H. (2019). Mahalanobis distance and its application for detecting multivariate outliers. *Facta Universitatis, Series: Mathematics and Informatics*, 583–595.
14. Haghpanah, M. A., Saeedizade, E., Masouleh, M. T., & Kalhor, A. (2022, February). Real-time facial expression recognition using facial landmarks and neural networks. In 2022 International Conference on Machine Vision and Image Processing (MVIP) (pp. 1–7). IEEE.
15. Hamill, P., Giordano, M., Ward, C., Giles, D., & Holben, B. (2016). An AERONET-based aerosol classification using the Mahalanobis distance. *Atmospheric Environment*, 140, 213–233.

16. Happy, S. L., & Routray, A. (2014). Automatic facial expression recognition using features of salient facial patches. *IEEE transactions on Affective Computing*, 6(1), 1–12.
17. Johnston, B., & Chazal, P. D. (2018). A review of image-based automatic facial landmark identification techniques. *EURASIP Journal on Image and Video Processing*, 2018(1), 1–23.
18. Juhong, A., & Pintavirooj, C. (2017). Face recognition based on facial landmark detection. In 2017 10th Biomedical Engineering International Conference (BMEiCON) (pp. 1–4). IEEE.
19. Khabarлак, K., & Koriashkina, L. (2021). Fast facial landmark detection and applications: A survey. *arXiv preprint arXiv:2101.10808*.
20. Li, D., Wang, Z., Gao, Q., Song, Y., Yu, X., & Wang, C. (2019). Facial expression recognition based on electroencephalogram and facial landmark localization. *Technology and Health Care*, 27(4), 373–387.
21. Mardia, K. V. (1970). Measures of multivariate skewness and kurtosis with applications. *Biometrika*, 57(3), 519–530.
22. Mardia, K. V. (1974). Applications of some measures of multivariate skewness and kurtosis in testing normality and robustness studies. *Sankhyā: The Indian Journal of Statistics, Series B*, 115–128.
23. Meshkani, S. A., Mehrabi, B., Yaghubpur, A., & Alghalandis, Y. F. (2011). The application of geochemical pattern recognition to regional prospecting: A case study of the Sanandaj–Sirjan metallogenic zone, Iran. *journal of geochemical exploration*, 183–195.
24. Nader, P., Honeine, P., & Beuseroy, P. (2014, September). Mahalanobis-based one-class classification. In 2014 IEEE International Workshop on Machine Learning for Signal Processing (MLSP) (pp. 1–6). IEEE.
25. Nestor, A., Vettel, J. M., & Tarr, M. J. (2013). Internal representations for face detection: An application of noise-based image classification to BOLD responses. *Human brain mapping*, 34(11), 3101–3115.
26. Osman Ali, A. S., Asirvadam, V. S., Malik, A. S., Eltoukhy, M. M., & Aziz, A. (2015). Age-invariant face recognition using triangle geometric features. *International Journal of Pattern Recognition and Artificial Intelligence*, 29(05), 1556006.
27. Prykhodko, S., Prykhodko, N., Makarova, L., & Pukhalevych, A. (2018, February). Application of the squared Mahalanobis distance for detecting outliers in multivariate non-Gaussian data. In 2018 14th International Conference on Advanced Trends in Radioelectronics, Telecommunications and Computer Engineering (TCSET) (pp. 962–965). IEEE.
28. Prykhodko, S., Makarova, L., Prykhodko, K., & Pukhalevych, A. (2020, February). Application of transformed prediction ellipsoids for outlier detection in multivariate non-Gaussian data. In 2020 IEEE 15th International Conference on Advanced Trends in Radioelectronics, Telecommunications and Computer Engineering (TCSET) (pp. 359–362). IEEE.
29. Saib, Y. M., & Pudaruth, S. (2021). Is Face Recognition with Masks Possible?. *International Journal of Advanced Computer Science and Applications*, 12(7).
30. Sen, S., Diawara, N., & Iftekharuddin, K. M. (2015). Statistical pattern recognition using Gaussian copula. *Journal of Statistical Theory and Practice*, 9, 768–777.
31. Van Aelst, S., & Rousseeuw, P. (2009). Minimum volume ellipsoid. *Wiley Interdisciplinary Reviews: Computational Statistics*, 1(1), 71–82.
32. Andrew R. Webb., & Keith D. Copsey. (2011). *Statistical Pattern Recognition*. Third Edition. John Wiley & Sons.
33. Wu, Y., & Ji, Q. (2019). Facial landmark detection: A literature survey. *International Journal of Computer Vision*, 127, 115–142.

УДК 519.8

DOI <https://doi.org/10.32782/IT/2023-2-7>

Олексій СЕРГЄЄВ

аспірант кафедри системного аналізу та управління, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, м. Дніпро, Україна, 49005, serhieiev.o.s@ntmu.one

ORCID: 0000-0001-5781-5540

Scopus-Author ID: 57220190743

Світлана УС

кандидат фізико-математичних наук, доцент, професор кафедри системного аналізу та управління, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, м. Дніпро, Україна, 49005, us.s.a@ntmu.one

ORCID: 0000-0003-0311-9958

Scopus-Author ID: 55603096000

Бібліографічний опис статті: Сергєєв, О., Ус, С. (2023). Аналіз сучасних підходів до розв'язання дискретних та неперервних багатоетапних задач розміщення. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 59–70, doi: <https://doi.org/10.32782/IT/2023-2-7>

АНАЛІЗ СУЧАСНИХ ПІДХОДІВ ДО РОЗВ'ЯЗАННЯ ДИСКРЕТНИХ ТА НЕПЕРЕРВНИХ БАГАТОЕТАПНИХ ЗАДАЧ РОЗМІЩЕННЯ¹

Стрімкий розвиток логістичних процесів між регіонами і державами є основною з причин ускладнення і подовження ланцюгів матеріальних потоків. Робота спрямована на дослідження моделей розміщення підприємств та їх застосування на практиці. Метою роботи є огляд актуальних проблем і підходів до розв'язання задач розміщення підприємств з подальшим аналізом методів їх розв'язання та визначення перспективних напрямків подальшого розвитку. Окреслено різні аспекти моделювання багатоетапних задач розміщення, включаючи вплив географічного положення, інфраструктури, доступності робочої сили, попиту та інших факторів на ефективність виробництва. Досліджено методи моделювання, включаючи лінійне та нелінійне програмування, генетичні алгоритми, аналіз ієрархій та нечітку логіку. Розглянуто вплив невизначеності та ризиків на процес прийняття рішень про розміщення підприємств. Особливу увагу приділено аналізу впливу розміщення підприємств на довкілля та сталий розвиток.

В роботі розглянуто загальні математичні постановки практичних задач, що можуть бути зведені до багатоетапних задач розміщення, запропоновано класифікацію методів та підходів до розв'язання задач цього типу. Проведено огляд актуальних наукових робіт для точних, евристичних, метаевристичних, багатокритеріальних, стохастичних, інтегрованих та континуальних методів до розв'язання багатоетапної задачі розміщення та окреслено сильні та слабкі сторони кожного з підходів. Окремо зазначено проблему розмірності, що виникає при розв'язанні задач дискретної постановки, коли кількість об'єктів, що розміщуються, є великою. Проаналізовано обмеження, що виникають у різних задачах з предметної області.

Авторами зазначено, що наявні методи до розв'язання є ефективними, однак, з метою урахування можливого масштабування проблеми, перспективним є дослідження поєднання неперервних методів з іншими підходами до розв'язання, такими як метаевристика або стохастика, для подальшого покращення їхньої продуктивності при застосуванні в практичних сценаріях.

Ключові слова: багатоетапна задача розміщення, генетичні алгоритми, неперервні задачі розміщення, оптимальне розбиття множин.

¹ Робота є складовою частиною досліджень за держбюджетною науковою темою 0123U100011 «Задачі аналізу, моделювання та оптимізації технологічних процесів у складних системах різної природи», що виконуються в НТУ «Дніпровська політехніка».

Oleksii SERHIEIEV

Postgraduate Student at the Department of System Analysis and Control, Dnipro University of Technology, Dmytro Yavornytskyi ave., 19, Dnipro, Ukraine, 49005, serhieiev.o.s@nmu.one

ORCID: 0000-0001-5781-5540

Scopus-Author ID: 57220190743

Svitlana US

Candidate of Physics and Mathematics Science, Associate Professor, Professor at the Department of System Analysis and Control, University of Technology, Dmytro Yavornytskyi ave., 19, Dnipro, Ukraine, 49005, us.s.a@nmu.one

ORCID: 0000-0003-0311-9958

Scopus-Author ID: 55603096000

To cite this article: Serhieiev, O., Us, S. (2023). Analiz suchasnykh pidkhodiv do rozviazannia dyskretnykh ta neperervnykh bahatoetapnykh zadach rozmishchennia [Analysis of modern approaches to solving discrete and continuous multi-stage allocation problems]. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 59–70, doi: <https://doi.org/10.32782/IT/2023-2-7>

ANALYSIS OF MODERN APPROACHES TO SOLVING DISCRETE AND CONTINUOUS MULTI-STAGE ALLOCATION PROBLEMS

The rapid development of logistics processes between regions and countries is the main reason for the complexity and elongation of material flow chains. The paper aims to consider various enterprise location models and their application in practice. The purpose of the paper is to review current problems and approaches to solving the issues of enterprise location, followed by an analysis of specific methods and identification of promising areas for further development. Various aspects of the multi-stage location problem are outlined, including the impact of geographic location, infrastructure, labor availability, demand and other factors on production efficiency. Various modeling techniques are investigated, including linear and nonlinear programming, genetic algorithms, hierarchy analysis and fuzzy logic. The publication discusses the impact of uncertainty and risks on the decision-making process of enterprise location. Particular attention is paid to analyzing the relationship between the location of enterprises and the reduction of negative environmental impact and sustainable development.

The paper considers general mathematical formulations of practical problems that can be reduced to multi-stage location problems and proposes a classification of methods and approaches to solving problems of this type. Relevant scientific works for exact, heuristic, metaheuristic, multi-criteria, stochastic, integrated and continuum methods for solving a multi-stage location problem is reviewed. The pros and cons of each approach are mentioned. The authors also emphasize the problem of dimensionality that arises when solving problems in discrete statement when the number of objects to be placed is large. The limitations arising in various problems from the subject area were analyzed.

The authors note that the existing solution approaches are effective. However, to consider the possible scaling of the problem, it is promising to study the combination of continuous methods with other solution approaches, such as metaheuristics or stochastic methods, further to improve their performance and practical scenarios applicability.

Key words: multi-stage allocation problem, genetic algorithms, continuous allocation problems, optimal partitioning of sets.

Вступ. Логістичний підхід вимагає об'єднання всіх складових частин процесу у єдину систему, щоб забезпечити ефективне управління матеріальними потоками. В сучасному економічному контексті транспортно-виробничі системи, що забезпечують процеси кругообігу матеріальних ресурсів, можуть бути досить складними та об'ємними. Ці системи включають багато економічних агентів та посередників, розташованих в різних регіонах та територіях. Крім того, вони характеризуються відмінностями у розмірах потреб різних клієнтів та іншими факторами. Іншими словами, логістичний підхід до управління матеріальними потоками спрямова-

ний на оптимізацію руху товарів та матеріалів від їх джерел до кінцевих пунктів призначення в найбільш ефективний та результативний спосіб. Цей підхід передбачає інтегроване бачення всього логістичного процесу, включаючи всі дії та учасників, що беруть участь в переміщенні, зберіганні та розподілі матеріалів.

На практиці менеджери логістики повинні координувати та інтегрувати різні елементи логістичного процесу, такі як транспортування, управління запасами, складське господарство та обробка замовлень, серед інших. Також вони повинні враховувати різні компроміси, такі як баланс між витратами на транспортування та

рівнями запасів та компроміси між швидкістю та вартістю транспортування.

Логістика матеріальних потоків здійснюється не тільки в рамках одного самостійного підприємства, але й за його межами, в інших, подібних до нього або пов'язаних з ним структурах, а це визначає таку властивість транспортно-виробничих процесів, як їх багатоетапність.

Прикладами задач розміщення підприємств є велика кількість практичних задач, що виникають у сьогоденні. Наведемо деякі з них.

Задача покращення просторового планування громадських служб охорони здоров'я через розробку моделей локації-алокації та доступності. Зокрема, розглядається визначення оптимального розміщення лікарень та інших медичних закладів з урахуванням таких факторів як попит населення, доступність та відстань до інших закладів охорони здоров'я. Дослідження (Polo, Acosta, Ferreira, & Dias, 2015) було проведено на прикладі організації медичного обслуговування в Лісабоні (Португалія), де застосування запропонованих методів дозволило покращити якість охорони здоров'я та ефективність витрат на її надання.

Оптимальне розміщення складів (You & Xiao, 2019). Запропоновано нову модель змішаного цілочисельного лінійного програмування для розв'язання проблеми розміщення складів з використанням лінеаризацій евклідової відстані.

Розміщення об'єктів утилізації відходів. У роботі (Yadav & Bhurjee, 2017) пропонується нова модель розміщення об'єктів для систем управління твердими побутовими відходами з урахуванням невизначеності вихідних даних. Автори використовують множинне лінійне програмування та метод мінімаксного критерію. Розв'язано задачу мінімізації затрат та максимізації покриття потреб в обслуговуванні населення.

Розміщення заводів та виробництв. Публікація (Paul, Chowdhury, & Ahsan, 2021) пропонує новий підхід до розміщення виробничих підприємств з використанням нечіткої логіки та систем виведення висновків. Дослідження є корисним для менеджерів з прийняття рішень в галузі виробництва, оскільки воно дозволяє ефективно використовувати нечітку логіку та системи виведення висновків для розв'язання складних проблем розміщення виробничих підприємств. Запропонована модель може бути застосована у різних галузях, включаючи автомобільну, електронну та харчову промисловість, де виробництво пов'язане з багатьма факторами, які впливають на ефективність виробництва.

Постановка завдання в предметній області. Проблема розміщення об'єктів є поширеною проблемою в дослідженні операцій та логістиці, яка полягає у визначенні оптимального розміщення об'єктів (таких як склади, фабрики, розподільчі центри та ін.) для обслуговування заданої множини точок попиту. Метою є мінімізація витрат при одночасному задоволенні попиту споживачів та вимог до рівня обслуговування.

Одним з головних викликів проблеми розміщення об'єктів є велика кількість задіяних змінних, включаючи кількість об'єктів, що підлягають розміщенню, розмір об'єктів, транспортні витрати, а також обсяги попиту. Ще однією проблемою є необхідність балансування суперечливих цілей, таких як мінімізація транспортних витрат при одночасному забезпеченні швидких термінів доставки. Невизначеність обсягів попиту, мінливість транспортних витрат у зв'язку з дорожнім рухом або погодними умовами, та неочікувані перебої у ланцюгу поставок можуть створювати перешкоди у логістичному процесі. Крім того, рішення про розміщення об'єктів можуть мати значні екологічні та соціальні наслідки, які також необхідно враховувати. Наприклад, розміщення об'єкта в густонаселеному районі може мати негативний вплив на місцевих жителів, тоді як розміщення об'єкта у віддаленому районі може мати негативний вплив на навколишнє середовище.

Для розв'язання задач розміщення в літературі запропоновано різні математичні методи оптимізації: лінійне і нелінійне програмування, цілочисельне програмування, направлені пошуки, та інші. Зараз актуальним є застосування штучного інтелекту і машинного навчання, з використанням алгоритмів кластеризації і моделей глибокого навчання, для розв'язання проблем розміщення об'єктів.

Математична модель. При огляді багатоетапних задач розміщення, існує велика кількість різних практичних завдань, що можуть бути зведені до задач такого типу. Наведемо загальну математичну постановку, що ілюструє проблему.

Припустимо, маємо такі множини: I – множина потенційних підприємств; J – множина етапів чи рівнів в логістичному процесі; K – множина споживачів.

Також відомо: f_i – фіксована вартість за розміщення підприємства $i \in I$; c_{ij} – вартість транспортування одиниці продукту з підприємства $i \in I$ на етапі $(j-1) \in J$ до підприємства $i \in I$ на етапі $j \in J$; h_{ij} – витрати на утримання одиниці продукції за період на підприємстві $i \in I$ на етапі $j \in J$;

d_{ik} – попит споживача $k \in K$ на етапі $j \in J$; q_i – потужність підприємства i ; x_{ij} – булева змінна, яка дорівнює 1, якщо підприємство $i \in I$ розташоване на етапі $j \in J$; y_{ijk} – кількість продукту, що відвантажено з підприємства $i \in I$ на етапі $(j-1) \in J$ до споживача $k \in K$ на етапі $j \in J$; z_{ij} – потужність підприємства $i \in I$ на етапі $j \in J$.

Цільова функція: метою є зменшення загальних витрат мережі ланцюгів поставок, які є сумою витрат на основні засоби, транспортні витрати та витрати на утримання. Це можна виразити так: мінімізувати

$$\sum_{i \in I} \sum_{j \in J} f_i x_{ij} + \sum_{i \in I} \sum_{j \in J \setminus \{1\}} \sum_{i' \in I} c_{ij} x_{ij} x'_{i,j+1} + \sum_{i \in I} \sum_{j \in J} \sum_{k \in K} h_{ij} z_{ij}, \quad (1)$$

за таких обмежень:

1. Загальна пропозиція на кожному об'єкті не може перевищувати потужність об'єкта:

$$\sum_{j \in J} z_{ij} \leq q_i x_{ij} \quad \forall i \in I, j \in J. \quad (2)$$

2. Загальний попит на кожному етапі від клієнта має бути задоволений:

$$\sum_{i \in I} y_{ijk} = d_{jk} \quad \forall j \in J, k \in K. \quad (3)$$

3. Кількість продукції, що перевозиться між підприємствами, має задовольняти рівняння балансу потоків:

$$\sum_{i \in I} y_{i,j-1,k} - \sum_{i \in I} y_{i,j,k} = 0, \quad \forall j \in J \setminus \{1\}, k \in K. \quad (4)$$

4. Кількість товару, що зберігається на складі, має задовольняти рівняння балансу запасів:

$$z_{ij} - z_{i,j-1} + \sum_{i' \in I} y_{i',j-1,k} - \sum_{i' \in I} y_{i',j,k} = 0, \quad \forall i \in I, j \in J \setminus \{1\}, k \in K. \quad (5)$$

Додатково, залежно від постановки задачі в предметній області, можуть виникати обмеження, які залежать від конкретної проблеми та її контексту.

Найбільш часто у задачах розміщення об'єктів використовуються перелічені нижче типи обмежень.

– Стосовно пропускної здатності: об'єкти мають обмежену пропускну здатність для обслуговування клієнтів, і ці обмеження необхідно враховувати, щоб гарантувати, що об'єкти зможуть задовольнити попит з боку клієнтів.

– Щодо відстані: підприємства мають бути розташовані на певній відстані від клієн-

тів, яких вони обслуговують, щоб мінімізувати транспортні витрати та забезпечити своєчасне надання послуг.

– Щодо рівня обслуговування: підприємства повинні забезпечувати певний рівень обслуговування клієнтів, наприклад, мінімальний рівень якості, оперативності та надійності.

– На підприємства можуть поширюватися різні обмеження щодо розміщення, наприклад, правила зонування, екологічні норми або інші юридичні чи логістичні обмеження.

– Витрати: об'єкти мають бути розташовані таким чином, щоб мінімізувати загальні витрати на обслуговування клієнтів, беручи до уваги такі фактори, як транспортні витрати, витрати на експлуатацію об'єкта та інші відповідні витрати.

– Обмеження щодо підключення до мережі: підприємства можуть потребувати підключення до більшої мережі або транспортної інфраструктури для забезпечення ефективного надання послуг.

– Попит на послуги може піддаватися різним обмеженням, таким як сезонність, мінливість або інші фактори, що обмежують обсяг попиту, який можна задовольнити в будь-який момент часу.

– Обмеження щодо спільного використання потужностей: для оптимізації надання послуг та мінімізації витрат може виникнути потреба у спільному використанні потужностей або ресурсів декількох підприємств.

Комбінації цих обмежень використовуються для опису проблем розміщення об'єктів та побудови відповідних математичних постановок для задач, розв'язання яких потребує комплексних методів оптимізації.

Класифікація методів розв'язання. З розглянутих вище практичних прикладів, загальної моделі (1)–(5) та можливих застосувань задач випливає, що підходить до розв'язування цих задач відрізняються залежно від предметної області чи обмежень, викликаних кількістю змінних. В сучасних дослідженнях використовуються такі **методи** для розв'язання двоетапних задач розміщення:

– *Точні*, що передбачають розв'язання проблеми розміщення об'єкта шляхом пошуку оптимального рішення за допомогою математичних методів оптимізації, таких як лінійне та нелінійне програмування, цілочисельне програмування та змішане цілочисельне програмування. Вони можуть бути обчислювально інтенсивними, але при цьому забезпечують гарантоване оптимальне рішення.

– *Метаевристичні* методи з використанням алгоритмів високого рівня для управління

пошуком наближеного до оптимального розв'язку. Вони включають генетичні алгоритми, імітаційне відпалювання, табу-пошук та оптимізацію мурашиних колоній та можуть бути обчислювально інтенсивними. Однак при цьому дозволяють швидко знаходити ефективні розв'язки.

– *Евристичні*, що передбачають використання алгоритмів апроксимації або евристик для пошуку наближених до оптимальних рішень проблеми розміщення об'єкта. Ці підходи часто швидші за точні методи, але можуть не гарантувати оптимального рішення.

– *Методи багатоцільової оптимізації*, які передбачають оптимізацію проблеми розміщення об'єкта з декількома конфліктуючими цілями, такими як мінімізація транспортних витрат і максимізація задоволення потреб клієнтів. Вони можуть допомогти збалансувати конфліктуючі цілі та знайти рішення, які є оптимальними з точки зору кількох цілей.

– *Стохастичні методи*, що передбачають моделювання невизначеності в проблемі розміщення об'єктів, таких як мінливість попиту, транспортні витрати та перебої в ланцюгу поставок та можуть допомогти розробити надійні логістичні плани, що можуть адаптуватися до несподіваних змін у навколишньому середовищі.

– *Інтегровані*, в яких застосовується інтеграція проблем розміщення об'єктів з іншими проблемами оптимізації ланцюга поставок, такими як управління запасами, планування виробництва і маршрутизація перевезень. Вони допомагають оптимізувати весь ланцюг поставок і знайти рішення, оптимальні з точки зору багатьох вимірів.

Розглянемо наявні дослідження, результати яких можна віднести до однієї в вищезазначених категорій.

Точні методи. Задача розміщення потужностей з єдиним джерелом постачання (Holmberg, Rönnqvist, & Yuan, 1999), для розв'язання якої запропоновано точний алгоритм, коли кожен клієнт обслуговується одним об'єктом. Для методів цього типу характерним є розгляд більш спрощеної математичної моделі ніж наведена раніше загальна математична постановка, наприклад:

$$\min \sum_{i=1}^m \sum_{j=1}^n c_{ij} x_{ij} + \sum_{i=1}^m f_i y_i, \quad (6)$$

за обмежень:

$$\sum_{j=1}^n a_j x_{ij} \leq b_i y_i, \forall i, \quad (7)$$

$$x_{ij} - y_i \leq 0, \forall i, j, x_{ij}, y_i \in \{1, 0\}, \forall i, j, \quad (8)$$

де m – кількість потенційних підприємств, n – кількість клієнтів, a_j – попит клієнта j ,

b_j – ємність підприємства j , f_j – фіксована вартість відкриття підприємства i , c_{ij} – вартість обслуговування i -м підприємством j -го користувача. Тобто розглядається одноетапна задача розміщення та вартість розміщення того чи іншого підприємства враховується у цільовій функції. Також застосовуються обмеження (2) та (4) із загальної математичної постановки.

Алгоритм ґрунтується на підході з розгалуженням і обмеженнями з підзадачею релаксації Лагранжа. Автори демонструють ефективність свого алгоритму, порівнюючи його з іншими точними та евристичними методами на наборі тестових прикладів. Результати показують, що запропонований алгоритм здатний оптимально розв'язувати задачі розмірності до 30×200 .

Проблема розміщення виробничих потужностей (Christensen & Klose, 2021) з диференційованими опуклими виробничими витратами, яка є варіантом класичної задачі розміщення потужностей, де вартість виробництва на кожному підприємстві моделюється як опукла функція від його виробничої потужності. Автори пропонують швидкий точний метод, заснований на підході гілок і цін, який використовує структуру задачі та опуклість виробничих витрат. Стверджується, що алгоритм є ефективним при розв'язуванні задач, вхідні дані яких мають не більше ніж 1000 клієнтів та 100 об'єктів. Запропонований метод порівнюється з іншими точними та евристичними підходами, показуючи, що він перевершує інші точні методи і дає розв'язки, які знаходяться в межах декількох відсотків від найкращих відомих рішень.

Метаевристичні методи. У роботі (Mitsuo, Fulya, & Lin, 2006) розглядається двоетапна транспортна задача у якій за мету взято мінімізацію логістичних витрат, враховуючи вартість розміщення центрів дистрибуції та витрат на транспортування між підприємствами, центрами дистрибуції та клієнтами. Пропонується розглянути таку математичну модель:

$$\min Z = \sum_{i=1}^I \sum_{j=1}^J t_{ij} x_{ij} + \sum_{j=1}^J \sum_{k=1}^K c_{jk} y_{jk} + \sum_{j=1}^J g_j z_j, \quad (9)$$

при обмеженнях

$$\sum_{j=1}^J x_{ij} \leq a_i, \sum_{k=1}^K y_{jk} \leq b_j z_j, \forall i, j, \quad (10)$$

$$\sum_{j=1}^J z_j \leq W, \forall j, k, \quad (11)$$

$$\sum_{j=1}^J y_{jk} \geq d_k, \forall j, k, \quad (12)$$

$$\sum_{i=1}^I \sum_{j=1}^J x_{ij} = \sum_{j=1}^J \sum_{k=1}^K y_{jk}, \forall i, j, k, \quad (13)$$

$$x_{ij}, y_{jk} \geq 0, z_j = \{0, 1\}, \forall i, j, k. \quad (14)$$

Тут I – кількість заводів ($i = \overline{1, I}$); J – кількість центрів дистрибуції ($j = \overline{1, J}$); K – кількість споживачів ($k = \overline{1, K}$); a_i – потужність заводу i ; b_j – потужність центру дистрибуції; d_k – попит k -го споживача; t_{ij} – одиниця вартості перевезення з заводу i до центра дистрибуції j ; c_{jk} – одиниця вартості перевезення з центра дистрибуції j до споживача k ; g_j – фіксована вартість відкриття j -го центра дистрибуції; w – верхня межа загальної кількості відчинених центрів дистрибуції; x_{ij} – кількість перевезеного ресурсу з заводу i до центру дистрибуції j ; y_{jk} – кількість перевезеного ресурсу з центру дистрибуції j до споживача k ; z_j : 0–1 – булева змінна, що відповідає 1, коли центр дистрибуції j відчинений.

Тут розглянуто стандартну дискретну модель з додатковими обмеженнями на максимальну кількість розміщення центрів дистрибуції. Автори пропонують застосовувати генетичний алгоритм для розв'язання задачі. План перевезень закодовано з використанням пріоритетів. Запропоновано вживання зваженого кросоверу, який гарантує коректність отриманої хромосоми та дозволяє не проводити процедуру її відновлення. Використано алгоритм мутації вставками або обміном. В роботі автори розв'язують декілька модельних задач та порівнюють отримані розв'язки з відомими результатами розв'язання задач невеликої розмірності із застосуванням повного перебору.

У роботі (Raj & Rajendran, 2012) досліджуються два різні випадки двоетапної транспортної задачі:

1) задачі з фіксованою платою за використання маршруту та платою за перевезення одиниці товару платою;

2) задачі з фіксованою платою за перевезення одиниці товару та врахуванням вартості відчинених центрів дистрибуції.

Для обох випадків, авторами запропоновано використання генетичного алгоритму з кодуванням плану перевезень у вигляді матриці. Алгоритми для розв'язування обох типів задач наведені у роботі та проведені числові експерименти. Цікавою особливістю запропонованого генетичного підходу є генерація першої популяції, що робиться з врахуванням різних алгоритмів побудови вихідних планів.

Інший підхід до розв'язання двоетапної транспортної задачі з фіксованою платою за використання маршруту запропоновано у (Cosma, Pop, & Sabo, An efficient solution approach for solving the two-stage supply chain problem with fixed costs associated to the routes, 2019). Розглянута модель є дискретною та відповідає стандартній постановці двоетапної задачі з фіксованою платою. Запропонований алгоритм є ітераційним процесом, в якому розглядаються декілька варіантів розподілу розв'язків, з яких обираються найкращі. Після кожного отриманого розв'язку проводиться направлений пошук кращих варіантів, за допомогою додаткових обмежень на відповідні ланцюги постачання. Розв'язок, отриманий в результаті роботи алгоритму, зберігається у випадку, коли він є кращим за всі попередні варіанти. Додатковим інструментом, що зменшує ймовірність потрапляння у локальний мінімум є процедура перемішування клієнтів при побудові нових розв'язків. У програмній реалізації застосована процедура хешу для ефективного видалення дублікатів. Автори наводять результати декількох експериментів, які свідчать про те, що запропонований алгоритм працює не гірше для будь-яких вхідних даних, та працює краще для даних невеликої розмірності.

Багатоступенева реверсивна логістична мережа розглянута у (Lee, Gen, & Rhee, 2009). Автори формулюють математичну модель переробної системи перепрофілюного відпрацювання як триетапну модель мережі логістики для мінімізації загальної вартості відправлення зворотної логістики та фіксованої вартості відкриття центрів демонтажу та обробки. Також в роботі враховується багатоетапність, багатопродуктовість та деякі умови для центрів демонтажу та обробки відповідно. Для розв'язання цієї проблеми застосовують генетичний алгоритм з пріоритетним кодуванням для першого та другого етапу з використанням методу зваженого кросоверу. В роботі наведені обчислювальні експерименти з доведенням ефективності запропонованого алгоритму.

Публікація (Khan, Pal, & Maiti, 2018) описує розв'язання задачі комівояжера за допомогою ройових методів. Застосовуються поєднання методу рою частинок та поєднується з генетичним алгоритмом, використовуючи вихідні дані одного для вхідних даних іншого з алгоритмів. Цікавим в роботі є дослідження ефективності застосування методу рулетки та багатоточкового циклічного кросоверу. При тестуванні з використанням вибірок великих розмірів, алгоритм демонструє 100 % успіх для чіткого випадку.

Двоетапна транспортна задача з обмеженням на кількість відчинених центрів дистрибуції наведена у (Cosma, Pop, & Danciulescu, A Parallel Algorithm for Solving a Two-Stage Fixed-Charge Transportation Problem, 2022). Пропонується застосування генетичного алгоритму для пошуку ефективного розв'язку. Автори роботи застосовують модифікований генетичний алгоритм з можливістю паралелізації процесу оцінки хромосом. Основна ідея цього алгоритму складається з двох кроків: обирання перспективних центрів дистрибуції та розв'язання декількох підзадач, в яких беруть участь тільки обрані центри дистрибуції. В роботі пропонується розділяти центри дистрибуції на «хороші» та «погані» з метою їх подальшого включення у відповідні «хороші» та «погані» плани перевезень. Саме процедура оцінки та класифікації в алгоритмі розроблена із застосування пулу потоків, що прискорює отримання результатів. В публікації наведений псевдокод реалізації відповідних процедур та аналіз впливу параметрів алгоритму (якість та швидкість) на кількість необхідних ітерацій для отримання найкращого розв'язку.

В (Salman, 2018) пропонується розв'язати двоетапну транспортну задачу з використанням табу-пошуку та процедури кодування. Математична модель відповідає дискретній постановці двоетапної транспортної задачі з обмеженням на кількість підприємств. Автор роботи наголошує на ідеї використання списку заборонених розв'язків, щоб врахувати можливі потрапляння в локальні мінімуми, з метою отримання кращого розв'язку. У роботі пропонується розглядати закодовані плани перевезень, що модифікуються на кожній ітерації i , якщо вони містять кращий розв'язок, то заносяться у лист заборонених розв'язків. У результаті, в роботі стверджується, що результати, отримані за допомогою табу-пошуку є кращими ніж ті, що отримані із застосуванням генетичного алгоритму.

Евристичні методи є схожими на модель (9) – (14), що була розглянута для метаевристичних підходів.

У статті (Buson, Roberti, & Toth, 2014) розглядається модель двоетапної транспортної задачі з фіксованою платою за використання маршруту. Автори пропонують ітерований евристичний алгоритм локального пошуку, що спирається на мінімізацію затрат на кожній ітерації з подальшою фазою перезапуску. Мінімізація витрат досягається за допомогою застосування процедури нижнього обмеження при розв'язанні проблеми з трьома індексами при обмеженнях з дійсними нерівностями. Запропо-

нований метод був протестований на двох еталонних прикладах з літератури.

Планування лікарняних мереж в умовах невизначеності (Mestre, Oliveira, & Barbosa-Róvoa, 2014) за допомогою набору дискретних сценаріїв, які охоплюють майбутні можливі імплементації системи. Було розроблено дві математичні моделі, які відображають різні припущення щодо рішень, які необхідно приймати за відсутності повної інформації про невизначені параметри. Незважаючи на те, що моделі представляють різні підходи до розв'язання проблеми невизначеності, вони мають однакову комплексну структуру, розглядаючи питання, які можуть виникнути в реальних умовах використання лікарень, що працюють в ієрархічній мережі і надають різні послуги. Розглядаються компроміси між покращенням доступом та мінімальними витратами, які розв'язуються за допомогою методу багатоцільового програмування з обмеженнями, що надає особі, яка ухвалює рішення, альтернативні варіанти для планування. Моделі були розроблені в контексті системи охорони здоров'я, що базується на структурі Національної служби охорони здоров'я (НСОЗ), а для перевірки їхньої застосовності було використано тематичне дослідження на прикладі португальської НСОЗ.

Задача обґрунтованого планування та оптимізації розміщення укриттів (He & Xie, 2022) з метою зменшення втрат від стихійних лих та покращення сталого розвитку міст. Розглядається послідовний розв'язок двокритеріальної задачі: що базується на логіці послідовних рішень для максимізації економічної стійкості та соціальної корисності.

Стохастичні методи. Розглянемо наступну математичну модель.

Нехай задані:

I – множина потенційних підприємств; J – множина клієнтів; K – множина можливих сценаріїв; f_i – фіксована вартість за розміщення підприємства $i \in I$; c_{ij} – вартість транспортування одиниці продукту з підприємства $i \in I$ до клієнта $j \in J$; p_{jk} – ймовірність сценарію K ; d_{jk} – попит клієнта j в сценарії k ; $x_{ij} \in [0,1]$ – неперервна змінна, що вказує на частку попиту клієнта j , що обслуговується підприємством i ; $y_{ik} \in \{0,1\}$ – булева змінна, що відповідає чи відчинено підприємство i у сценарії k ; $w_{ijk} \in [0,1]$ – неперервна змінна, що показує частку попиту клієнта j у сценарії k , що обслуговується з об'єкта i .

Мінімізувати

$$\sum_{i \in I} \sum_{k \in K} f_i y_{ik} + \sum_{i \in I} \sum_{j \in J} \sum_{k \in K} c_{ij} w_{ijk} p_{jk}, \quad (15)$$

за таких обмежень

– на потужність:

$$\sum_{j \in J} \sum_{k \in K} d_{jk} w_{ijk} \leq \sum_{j \in J} x_{ij} \quad \forall i \in I, \quad (16)$$

– на задоволення попиту для кожного споживача j та сценарію k :

$$\sum_{i \in I} w_{ijk} \leq 1, \quad \forall j \in J, k \in K, \quad (17)$$

– на розподіл ресурсів:

$$w_{ijk} \leq x_{ij} y_{jk}, \quad \forall i \in I, j \in J, k \in K, \quad (18)$$

– природні:

$$x_{ij}, w_{ijk} \geq 0, y_{jk} \in \{0, 1\} \quad \forall i \in I, j \in J, k \in K. \quad (19)$$

Робота (Bieniek, 2015) зосереджується на проблемі розміщення об'єктів зі стохастичним попитом, коли попит клієнтів є невизначеним і моделюється як випадкова величина. Автор пропонує двоетапний стохастичний підхід для розв'язання проблеми та використовує генерацію скінченної кількості сценаріїв і розв'язує задачу для кожного з них, застосовуючи змішане лінійне програмування. Результати роботи демонструють ефективність запропонованого методу на низці числових прикладів, показуючи, що він гарантує розв'язки, стійкі до невизначеності у попиті, і перевершує деякі евристичні підходи.

Інтеграція стохастичної моделі інвентаризації з моделлю мережі ланцюга поставок розглянута у (You & Grossmann, 2008). Цю задачу можна сформулювати як широкомасштабну комбінаторну оптимізаційну модель, яка включає нелінійні умови. У роботі проаналізовано властивості базової моделі та пропонуються методи розв'язання для спільного проектування мережі ланцюга постачання та моделі управління запасами для даного продукту. Модель сформульована як задача нелінійного цілочисельного програмування. Переформулювавши її як задачу змішаного цілочисельного нелінійного програмування і, використовуючи пов'язану модель опуклої релаксації для ініціалізації, автори пропонують евристичний метод для швидкого отримання якісних рішень. Крім того, розроблено алгоритм декомпозиції на основі релаксації Лагранжа для отримання глобальних або майже глобальних оптимальних рішень.

Інтегровані методи. Застосовуються до математичних моделей вигляду (1) – (5), однак підходи до розв'язання є комбінацією різних з розглянутих вище типів.

Вже відома двоетапна транспортна задача з фіксованою платою за використання маршруту в публікації (Calvetea, Galé, & Toth, 2018). Для розв'язання цієї задачі, автори застосовують перехід до іншої форми задачі, що схожа на двоетапну транспортну задачу з врахуванням вартості перевезення одиниці товару. Вартість товару представлена у вигляді умовного виразу залежно від того використовується цей маршрут чи ні. В якості алгоритму для розв'язання застосовують генетичний алгоритм. Хромосому кодують з використанням матричного представлення. В роботі доведено робастність запропонованого алгоритму з використанням випадково згенерованих сутностей. Для кожної з цих сутностей, розв'язок був знайдений за раціональний час.

В публікації (Arabzad, Ghorbani, & Zolfani, 2015) досліджується проблема location-allocation з метою розроблення ланцюга поставок. Розглянуто ланцюг постачання з декількома постачальниками, продуктами, заводами та клієнтами. Для ефективного розв'язання проблеми невизначеності параметрів попиту і витрат застосовано сценарний підхід. Формулювання являє собою стійке багатоцільове частково-цілочисельне лінійне програмування контексті якого одночасно враховуються дві суперечливі цілі: (1) мінімізація загальних витрат ланцюга поставок, включаючи витрати на сировину, транспортні витрати та витрати на створення заводів, і (2) мінімізація загального рівня погіршення якості, спричиненого альтернативними варіантами перевезень. Застосовано метод зваженої суми та розв'язано одноцільову задачу частково-цілочисельного програмування.

Прикладами комбінування декількох підходів можуть бути роботи (Khan, Pal, & Maiti, 2018), (Buson, Roberti, & Toth, 2014), (Mestre, Oliveira, & Barbosa-Póvoa, 2014), опис яких наведено вище.

Розглянувши підходи до розв'язання задач розміщення, можна зробити висновок, що увагу дослідників більш захоплює використання метаевристичного підходу. Використовуючи цей підхід, авторам вдається досягти ефективних рішень за оптимальний час. Однак у дискретному випадку знайти ефективний розв'язок за адекватний час вдається лише для задач невеликих та середніх розмірностей. При застосуванні точних методів, розмірності задач обмежуються ще більше (розглядають невеликі задачі). Враховуючи це, слід звернути увагу на дослідження, де розв'язки для задач можуть бути отримані для більшої кількості об'єктів.

В монографії (Кісельова, 2018) містяться засади теорії оптимального розбиття множин, що можуть бути застосовані при розв'язанні багатоетапних задач розміщення, у випадку коли користувачі або ресурс розподілені неперервно в даній області. У роботі наведені основні методи негладкої оптимізації, а саме субградієнтні методи з розтягом простору в напрямі субградієнту та описані задачі оптимального розбиття множин, які, будучи задачами нескінченновимірної оптимізації, зводяться до декількох задач негладкої, але скінченновимірної оптимізації.

Дослідженню багатоетапних задач розміщення присвячена робота (Us & Stanina, 2017). Розглядається клас задач, в яких існує кілька груп розміщуваних об'єктів, кожна з яких має свою множину можливих місць розміщення (в окремому випадку ці множини можуть збігатися) і регламентовано зв'язки між об'єктами, наприклад задана певна ієрархія між об'єктами розміщення. У роботі зазначається відсутність досліджень випадку континуальних багатоетапних задач і наводяться приклади з предметної області. Також наведена модель двоетапної континуальної задачі розміщення. Континуальність розглядається для кожного з етапів, іноді комбінована з дискретними випадками.

Загальні засади створення нових математичних моделей для процесів двоетапного виробництва та обґрунтування методів їх розв'язування з урахуванням неперервності розподілення ресурсів та наявності декількох етапів виробництва описано в монографії (Станіна, Ус, & Коряшкіна, 2021). Особливу увагу приділено експериментальним дослідженням та аналізу результатів модельних задач. Наведено приклади практичного застосування запропонованих моделей у задачах оптимізації двоетапного розподілу матеріального потоку в паливно-енергетичному комплексі. Наприклад, потрібно знайти такі розбиття множини Ω на N вимірних за Лебегом підмножин $(\Omega_1, \dots, \Omega_N)$ (серед яких можуть бути і порожні) і обсяги перевезень $V_{11}, \dots, V_{NM}$, які забезпечують

$$\min_{\{\Omega_1, \dots, \Omega_N\}, \{V_{11}, \dots, V_{NM}\}} F(\{\Omega_1, \dots, \Omega_N\}, \{V_{11}, \dots, V_{NM}\}), \quad (20)$$

при обмеженнях:

$$\int_{\Omega_i} \rho(x) dx = b'_i, \forall i = \overline{1, N}, \quad (21)$$

$$\sum_{j=1}^M v_{ij} = b'_i, \forall i = \overline{1, N}, \quad (22)$$

$$\sum_{i=1}^N v_{ij} = b''_j, \forall j = \overline{1, M}, \quad (23)$$

$$(\Omega_1, \dots, \Omega_N) \in \sum_{\Omega}, \quad (24)$$

$$\sum_{\Omega} = \{(\Omega_1, \dots, \Omega_N) : \bigcup_{i=1}^N \Omega_i = \Omega, \text{mes}(\Omega_i \cap \Omega_j) = \emptyset, i \neq j, i, j = \overline{1, N}\}, \quad (25)$$

$$v_{ij} \geq 0, i = \overline{1, N}, j = \overline{1, M}, \quad (26)$$

де

$$F(\{\Omega_1, \dots, \Omega_N\}, \{V_{11}, \dots, V_{NM}\}) = \sum_{i=1}^N \int_{\Omega_i} c'_i(x, \tau'_i) \rho(x) dx + \sum_{i=1}^N \sum_{j=1}^M c''_{ij}(\tau'_i, \tau''_j), \quad (27)$$

– b'_i – потужність i -го підприємства I-го етапу;
 – b''_j – потужність j -го підприємства II-го етапу;
 – $c'_i(x, \tau'_i)$ – вартість доставки одиниці сировини з точки $x \in \Omega$ до i -го підприємства I етапу;
 – $c''_{ij}(\tau'_i, \tau''_j)$ – вартість доставки одиниці сировини від i -го підприємства I етапу до j -го підприємства II етапу;
 – $\rho(x)$ – кількість ресурсів в точці $x \in \Omega$;
 – $\tau'_i = (\tau'_{i1}, \tau'_{i1})$ – координати i -го підприємства i -го етапу;
 – v_{ij} – обсяг продукції, що постачається від i -го підприємства I етапу до j -го підприємства II етапу: $v_{ij} \geq 0, i = \overline{1, N}; j = \overline{1, M}$.

Авторами отримані розв'язки для задач більших розмірностей, ніж для тих, що зазвичай розв'язанні за допомогою метаевристичних методів

Підсумуємо розгляд різних класів задач та наведемо переваги та недоліки кожних з них у таблиці 1.

Таблиця 1

Недоліки та переваги кожного з класу методів

№	Клас методу	Переваги	Недоліки
1	2	3	4
1	Точні	- Гарантують оптимальне рішення; - використовуються для широкого спектру типів задач розміщення об'єктів	- Вимагають значних обчислень, особливо для великих задач; - мають проблеми з масштабуванням для обробки змін параметрів або прийняття рішень в реальному часі
2	Евристичні	- Швидші за точні методи; - працюють з великими проблемами та динамічними середовищами	- Не гарантують оптимального рішення; - не підтримують усі типи проблем розміщення об'єктів

Закінчення таблиці 1

1	2	3	4
3	Мета-евристичні	- Швидко знаходять високоякісні рішення; - застосовуються зі складними типами проблем та динамічними середовищами	- Можуть не гарантувати оптимального рішення; - вимагають значних обчислень для великих задач
4	Багато-цільова оптимізація	- Дозволяють збалансувати кілька конфліктуючих цілей; - допомагають визначити компроміси та знайти оптимальні рішення щодо декількох цілей	- Складні для імплементації та аналізу; - потребують додаткових критеріїв прийняття рішень та участі зацікавлених сторін
5	Стохастичні	- Працюють з невизначеністю та мінливістю проблеми; - допомагають розробити більш надійні та гнучкі плани розміщення об'єктів	- Можуть бути більш складними для реалізації та аналізу; - вимагають додаткових даних та обчислювальних ресурсів для моделювання невизначеності та мінливості
6	Інтегровані	- Дозволяють оптимізувати весь ланцюг постачання та знайти рішення, які є оптимальними за багатьма параметрами; - допомагають визначити можливості для синергії та економії витрат	- Складні для імплементації та аналізу; - необхідні додаткові дані та обчислювальні ресурси для інтеграції декількох оптимізаційних задач
7	Континуальні	Менше залежать від розмірності задачі	Більш складна програмна реалізація алгоритму

Висновки. Задача розміщення об'єктів є важливою темою в галузі дослідження операцій. У цій роботі представлено огляд проблеми, її передумови та формулювання та обговорено різні підходи до розв'язання задачі розміщення об'єктів, включаючи точні методи, евристичні методи, метаевристичні методи та стохастичні методи. Кожен з цих підходів має свої сильні і слабкі сторони, і вибір методу залежить від конкретних характеристик задачі, що розглядається. В цілому, проблема розміщення об'єктів є складною і багатогранною та потребує ретельного вивчення та інноваційних рішень.

Подальші дослідження можуть бути зосереджені на побудові нових моделей практичних задач, що можуть бути зведені до багатоетапних задач з неперервно розподіленими ресурсами. Алгоритми розв'язання таких задач можуть краще справлятися з невизначеністю і складнощами реального світу. Крім того, можна дослідити потенційні переваги поєднання неперервних методів з іншими підходами до розв'язання, такими як метаевристика або стохастичні методи, для подальшого покращення продуктивності та застосовності в практичних сценаріях.

ЛІТЕРАТУРА:

1. Location-Allocation and accessibility models for improving the spatial planning of public health services / G. Polo et al. *Plos one*. 2015. Vol. 10, no. 3. P. 1–14. URL: <https://doi.org/10.1371/journal.pone.0119190> (date of access: 22.04.2023).
2. Optimal mathematical programming for the warehouse location problem with Euclidean distance linearization / M. You et al. *Computers & industrial engineering*. 2019. Vol. 136. P. 70–79. URL: <https://doi.org/10.1016/j.cie.2019.07.020> (date of access: 22.04.2023).
3. A facility location model for municipal solid waste management system under uncertain environment / V. Yadav et al. *Science of the total environment*. 2017. Vol. 603–604. P. 760–771. URL: <https://doi.org/10.1016/j.scitotenv.2017.02.207> (date of access: 22.04.2023).
4. An advanced decision-making model for evaluating manufacturing plant locations using fuzzy inference system / S. K. Paul et al. *Expert systems with applications*. 2022. Vol. 191. P. 116378. URL: <https://doi.org/10.1016/j.eswa.2021.116378> (date of access: 22.04.2023).
5. Holmberg K., Rönnqvist M., Yuan D. An exact algorithm for the capacitated facility location problems with single sourcing. *European journal of operational research*. 1999. Vol. 113, no. 3. P. 544–559. URL: [https://doi.org/10.1016/s0377-2217\(98\)00008-3](https://doi.org/10.1016/s0377-2217(98)00008-3) (date of access: 22.04.2023).
6. Christensen T. R. L., Klose A. A fast exact method for the capacitated facility location problem with differentiable convex production costs. *European journal of operational research*. 2020. URL: <https://doi.org/10.1016/j.ejor.2020.11.048> (date of access: 22.04.2023).

7. Gen M., Altıparmak F., Lin L. A genetic algorithm for two-stage transportation problem using priority-based encoding. *OR spectrum*. 2006. Vol. 28, no. 3. P. 337–354. URL: <https://doi.org/10.1007/s00291-005-0029-9> (date of access: 22.04.2023).
8. Antony Arokia Durai Raj K., Rajendran C. A genetic algorithm for solving the fixed-charge transportation model: two-stage problem. *Computers & operations research*. 2012. Vol. 39, no. 9. P. 2016–2032. URL: <https://doi.org/10.1016/j.cor.2011.09.020> (date of access: 22.04.2023).
9. Cosma O., Pop P., Sabo C. An efficient solution approach for solving the two-stage supply chain problem with fixed costs associated to the routes. *Procedia computer science*. 2019. Vol. 162. P. 900–907. URL: <https://doi.org/10.1016/j.procs.2019.12.066> (date of access: 22.04.2023).
10. Lee J.-E., Gen M., Rhee K.-G. Network model and optimization of reverse logistics by hybrid genetic algorithm. *Computers & industrial engineering*. 2009. Vol. 56, no. 3. P. 951–964. URL: <https://doi.org/10.1016/j.cie.2008.09.021> (date of access: 22.04.2023).
11. Khan I., Pal S., Maiti M. K. A hybrid PSO-GA algorithm for traveling salesman problems in different environments. *International journal of uncertainty, fuzziness and knowledge-based systems*. 2019. Vol. 27, no. 05. P. 693–717. URL: <https://doi.org/10.1142/s0218488519500314> (date of access: 22.04.2023).
12. Cosma O., Pop P. C., Dănciulescu D. A parallel algorithm for solving a two-stage fixed-charge transportation problem. *Informatica*. 2020. P. 1–26. URL: <https://doi.org/10.15388/20-infor432> (date of access: 22.04.2023).
13. Salman M. Using tabu search algorithm for solving two-stage transportation problem based on decoding procedure. *Sci.Int*. 2018. Vol. 30, no. 4. P. 643–651.
14. Buson E., Roberti R., Toth P. A reduced-cost iterated local search heuristic for the fixed-charge transportation problem. *Operations research*. 2014. Vol. 62, no. 5. P. 1095–1106. URL: <https://doi.org/10.1287/opre.2014.1288> (date of access: 22.04.2023).
15. Mestre A. M., Oliveira M. D., Barbosa-Póvoa A. P. Location–allocation approaches for hospital network planning under uncertainty. *European journal of operational research*. 2015. Vol. 240, no. 3. P. 791–806. URL: <https://doi.org/10.1016/j.ejor.2014.07.024> (date of access: 22.04.2023).
16. He L., Xie Z. Optimization of urban shelter locations using bi-level multi-objective location-allocation model. *International journal of environmental research and public health*. 2022. Vol. 19, no. 7. P. 4401. URL: <https://doi.org/10.3390/ijerph19074401> (date of access: 22.04.2023).
17. Bieniek M. A note on the facility location problem with stochastic demands. *Omega*. 2015. Vol. 55. P. 53–60. URL: <https://doi.org/10.1016/j.omega.2015.02.006> (date of access: 22.04.2023).
18. You F., Grossmann I. E. Mixed-Integer nonlinear programming models and algorithms for large-scale supply chain design with stochastic inventory management. *Industrial & engineering chemistry research*. 2008. Vol. 47, no. 20. P. 7802–7817. URL: <https://doi.org/10.1021/ie800257x> (date of access: 22.04.2023).
19. A matheuristic for the two-stage fixed-charge transportation problem / H. I. Calvete et al. *Computers & operations research*. 2018. Vol. 95. P. 113–122. URL: <https://doi.org/10.1016/j.cor.2018.03.007> (date of access: 22.04.2023).
20. Hashemkhani Zolfani S., Arabzad S. M., Ghorbani M. A multi-objective robust optimization model for a facility location-allocation problem in a supply chain under uncertainty. *Engineering economics*. 2015. Vol. 26, no. 3. URL: <https://doi.org/10.5755/j01.ee.26.3.4287> (date of access: 22.04.2023).
21. Кісельова О. Становлення та розвиток теорії оптимального розбиття множин. Теоретичні і практичні застосування. Дніпро : Ліра, 2018. 532 р.
22. Us S., Stanyina O. Location-allocation problems. *Комп'ютерне моделювання: аналіз, управління, оптимізація*. 2017. № 1. С. 73–79.
23. Станіна О., Ус С., Коряшкіна Л. Моделі та методи розв'язання задач оптимального розміщення двоетапного виробництва з неперервно розподіленим ресурсом. Дніпро, 2021.

REFERENCES:

1. Polo, G., Acosta, C. M., Ferreira, F., & Dias, R. A. (2015). Location-Allocation and accessibility models for improving the spatial planning of public health services. *Plos One*, 10(3), Stattia e0119190. <https://doi.org/10.1371/journal.pone.0119190>
2. You, M., Xiao, Y., Zhang, S., Yang, P., & Zhou, S. (2019). Optimal mathematical programming for the warehouse location problem with Euclidean distance linearization. *Computers & Industrial Engineering*, 136, 70–79. <https://doi.org/10.1016/j.cie.2019.07.020>
3. Yadav, V., Bhurjee, A. K., Karmakar, S., & Dikshit, A. K. (2017). A facility location model for municipal solid waste management system under uncertain environment. *Science of the Total Environment*, 603–604, 760–771. <https://doi.org/10.1016/j.scitotenv.2017.02.207>

4. Paul, S. K., Chowdhury, P., Ahsan, K., Ali, S. M., & Kabir, G. (2022). An advanced decision-making model for evaluating manufacturing plant locations using fuzzy inference system. *Expert Systems With Applications*, 191, 116378. <https://doi.org/10.1016/j.eswa.2021.116378>
5. Holmberg, K., Rönnqvist, M., & Yuan, D. (1999). An exact algorithm for the capacitated facility location problems with single sourcing. *European Journal of Operational Research*, 113(3), 544–559. [https://doi.org/10.1016/s0377-2217\(98\)00008-3](https://doi.org/10.1016/s0377-2217(98)00008-3)
6. Christensen, T. R. L., & Klose, A. (2020). A fast exact method for the capacitated facility location problem with differentiable convex production costs. *European Journal of Operational Research*. <https://doi.org/10.1016/j.ejor.2020.11.048>
7. Gen, M., Altıparmak, F., & Lin, L. (2006). A genetic algorithm for two-stage transportation problem using priority-based encoding. *Or Spectrum*, 28(3), 337–354. <https://doi.org/10.1007/s00291-005-0029-9>
8. Antony Arokia Durai Raj, K., & Rajendran, C. (2012). A genetic algorithm for solving the fixed-charge transportation model: Two-stage problem. *Computers & Operations Research*, 39(9), 2016–2032. <https://doi.org/10.1016/j.cor.2011.09.020>
9. Cosma, O., Pop, P., & Sabo, C. (2019). An efficient solution approach for solving the two-stage supply chain problem with fixed costs associated to the routes. *Procedia Computer Science*, 162, 900–907. <https://doi.org/10.1016/j.procs.2019.12.066>
10. Lee, J.-E., Gen, M., & Rhee, K.-G. (2009). Network model and optimization of reverse logistics by hybrid genetic algorithm. *Computers & Industrial Engineering*, 56(3), 951–964. <https://doi.org/10.1016/j.cie.2008.09.021>
11. Khan, I., Pal, S., & Maiti, M. K. (2019). A hybrid PSO-GA algorithm for traveling salesman problems in different environments. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 27(05), 693–717. <https://doi.org/10.1142/s0218488519500314>
12. Cosma, O., Pop, P. C., & Dănculescu, D. (2020). A parallel algorithm for solving a two-stage fixed-charge transportation problem. *Informatica*, 1–26. <https://doi.org/10.15388/20-infor432>
13. Salman, M. (2018). Using tabu search algorithm for solving two-stage transportation problem based on decoding procedure. *Sci.Int.*, 30(4), 643–651.
14. Buson, E., Roberti, R., & Toth, P. (2014). A reduced-cost iterated local search heuristic for the fixed-charge transportation problem. *Operations Research*, 62(5), 1095–1106. <https://doi.org/10.1287/opre.2014.1288>
15. Mestre, A. M., Oliveira, M. D., & Barbosa-Póvoa, A. P. (2015). Location-allocation approaches for hospital network planning under uncertainty. *European Journal of Operational Research*, 240(3), 791–806. <https://doi.org/10.1016/j.ejor.2014.07.024>
16. He, L., & Xie, Z. (2022). Optimization of urban shelter locations using bi-level multi-objective location-allocation model. *International Journal of Environmental Research and Public Health*, 19(7), 4401. <https://doi.org/10.3390/ijerph19074401>
17. Bieniek, M. (2015). A note on the facility location problem with stochastic demands. *Omega*, 55, 53–60. <https://doi.org/10.1016/j.omega.2015.02.006>
18. You, F., & Grossmann, I. E. (2008). Mixed-Integer nonlinear programming models and algorithms for large-scale supply chain design with stochastic inventory management. *Industrial & Engineering Chemistry Research*, 47(20), 7802–7817. <https://doi.org/10.1021/ie800257x>
19. Calvete, H. I., Galé, C., Iranzo, J. A., & Toth, P. (2018). A matheuristic for the two-stage fixed-charge transportation problem. *Computers & Operations Research*, 95, 113–122. <https://doi.org/10.1016/j.cor.2018.03.007>
20. Hashemkhani Zolfani, S., Arabzad, S. M., & Ghorbani, M. (2015). A multi-objective robust optimization model for a facility location-allocation problem in a supply chain under uncertainty. *Engineering Economics*, 26(3). <https://doi.org/10.5755/j01.ee.26.3.4287>
21. Kiselova, O. (2018). *Stanovlennia ta rozvytok teorii optymalnoho rozbytta mnozhyn. Teoretychni i praktychni zastosuvannia*. Lira [in Ukrainian].
22. Us, S., & Stanyna, O. (2017). Location-allocation problems. *Komputerne Modeluvanna: Analiz, Upravlinna, Optimizacia*, 1, 73–79.
23. Stanina, O., Us, S., & Koriashkina, L. (2021). *Modeli ta metody rozviazannia zadach optymalnoho rozmishchennia dvoetapnoho vyrobnytstva z neperervno rozpodilenyim resursom* [in Ukrainian].

УДК 517.91:532.2

DOI <https://doi.org/10.32782/IT/2023-2-8>**Світлана СИВАШ**

кандидат фізико-математичних наук, доцент, доцент кафедри «Математика, фізика та астрономія», Одеський національний морський університет, вул. Мечникова, 34, м. Одеса, Україна, 65029, rusboris@ukr.net

ORCID: 0000-0002-9726-7865

Галина СОКОЛОВСЬКА

старший викладач кафедри «Математика, фізика та астрономія», Одеський національний морський університет, вул. Мечникова, 34, м. Одеса, Україна, 65029, sokolovskahalyna7@gmail.com

ORCID: 0000-0001-8161-1660

Бібліографічний опис статті: Сиваш, С., Соколовська, Г. (2023). Наближене розв'язання задачі стаціонарної теплопровідності з екстремальною граничною умовою. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 71–75, doi: <https://doi.org/10.32782/IT/2023-2-8>

НАБЛИЖЕНЕ РОЗВ'ЯЗАННЯ ЗАДАЧІ СТАЦІОНАРНОЇ ТЕПЛОПРОВІДНОСТІ З ЕКСТРЕМАЛЬНОЮ ГРАНИЧНОЮ УМОВОЮ

Широке коло прикладних задач математичної фізики, астрофізики приводить до пошуку розв'язків задач мінімізації квадратичних функціоналів. Зокрема, задачі мінімізації квадратичних функціоналів виду $\| (Au)(x) - g(x) \|^2 = \int_{R_n} \rho(x) |(Au)(x) - g(x)|^2 dx \rightarrow \inf$ з операторами згортки

$(Au)(x) = \int_{\Omega} k(x-s)u(s)ds$, $\Omega \subset R_n$ часто зустрічаються в теорії механізмів, лінійних електричних та радіотехнічних ланцюгів, оптимальних фільтрів, систем регулювання. Некоректні задачі для лінійних рівнянь також приводяться до розв'язку задач мінімізації квадратичних функціоналів.

Відомо, що екстремальні задачі допускають розв'язки у явному вигляді лише в деяких окремих випадках. Тому побудова та обґрунтування методів їх наближеного розв'язання має значний теоретичний та практичний інтерес. Застосування чисельних методів відкриває можливості побудови розв'язків нових екстремальних задач для рівнянь математичної фізики та алгоритмізації цього процесу.

У роботі поставлена задача стаціонарної теплопровідності з екстремальною граничною умовою. За допомогою перетворення Фур'є та формул Сохоцького вона зводиться до розв'язання матричної задачі Рімана на дійсній осі з непозитивною системою часткових індексів. Доведена еквівалентність цих задач з точки зору їх розв'язності та формули, які виражають залежність розв'язку екстремальної задачі від розв'язків відповідної задачі Рімана. На основі дослідження задачі Рімана встановлено умови нормальної розв'язності екстремальної задачі. Наближені розв'язки екстремальної задачі будуються на основі наближених розв'язків задачі Рімана. Запропоновано та обґрунтовано проєкційний метод їх знаходження. Проведена оцінка збіжності наближених розв'язків екстремальної задачі до її точного розв'язку.

Досліджено також винятковий випадок екстремальної задачі та побудовано відповідні наближені розв'язки.

Окрім того, проведено чисельний експеримент для конкретних значень параметрів задачі, результати якого повністю узгоджуються з теоретичними висновками. Це підтверджує високу ефективність запропонованого методу побудови наближених розв'язків екстремальних задач математичної фізики. Отримані результати можуть бути використані при розв'язанні задач теорії пружності та термопружності, теплопровідності та інших прикладних задач.

Ключові слова: задача Рімана, крайова задача, частковий індекс, розв'язність, нетеровість, наближений розв'язок.

Svitlana SYVASH

Candidate of Physical and Mathematical Sciences, Associate Professor, Associate Professor at the Department "Mathematics, Physics and Astronomy", Odesa National Maritime University, Mechnikova str., 34, Odesa, Ukraine, 65029, rusboris@ukr.net
ORCID: 0000-0002-9726-7865

Galyna SOKOLOVSKA

Senior Lecturer at the Department "Mathematics, Physics and Astronomy", Odesa National Maritime University, Mechnikova str., 34, Odesa, Ukraine, 65029, sokolovskahalyna7@gmail.com
ORCID: 0000-0001-8161-1660

To cite this article: Syvash, S., Sokolovska, G. (2023). Nablyzhene rozvyazannya zadachi statsionarnoyi teploprovodnosti z ekstremalnoyu hranychnoyu umovoyu [Approximate solution of the stationary heat conduction problem with an extreme boundary condition]. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 71–75, doi: <https://doi.org/10.32782/IT/2023-2-8>

APPROXIMATE SOLUTION OF STATIONARY HEAT CONDUCTIVITY PROBLEM WITH EXTREME BOUNDARY CONDITION

A wide range of applied problems of mathematical physics and astrophysics leads to the search for solutions to problems of minimization of quadratic functionals. In particular, problems of minimization of quadratic functionals of the form $\| (Au)(x) - g(x) \|^2 = \int_{R_n} \rho(x) |(Au)(x) - g(x)|^2 dx \rightarrow \inf$ with convolution operators $(Au)(x) = \int_{\Omega} k(x-s)u(s)ds$, $\Omega \subset R_n$ often found in the theory of mechanisms, linear electric and radio engineering circuits, optimal filters, control systems. Incorrect problems for linear equations also lead to the solution of minimization problems of quadratic functionals.

It is known that extremal problems admit solutions in an explicit form only in some individual cases. Therefore, the construction and substantiation of methods for their approximate solution is of significant theoretical and practical interest. The use of numerical methods opens up the possibility of constructing solutions to new extreme problems for the equations of mathematical physics and algorithmizing this process.

The paper presents the problem of stationary thermal conductivity with an extreme boundary condition. With the help of the Fourier transform and Sohotsky formulas, it is reduced to the solution of the Riemann matrix problem on the real axis with a non-positive system of partial indices. The equivalence of these problems is proved from the point of view of their solvability and formulas expressing the dependence of the solution of the extremal problem on the solutions of the corresponding Riemann problem. Based on the study of the Riemann problem, the conditions for the normal solvability of the extremal problem have been established. Approximate solutions of the extremal problem are constructed on the basis of approximate solutions of the Riemann problem. A projection method for finding them is proposed and substantiated. Convergence of approximate solutions of the extreme problem to its exact solution was evaluated.

The exceptional case of the extremal problem was also studied and the corresponding approximate solutions were constructed.

In addition, a numerical experiment was conducted for specific values of the parameters of the problem, the results of which are fully consistent with the theoretical conclusions. This confirms the high efficiency of the proposed method of constructing approximate solutions of extreme problems of mathematical physics. The obtained results can be used in solving problems of the theory of elasticity and thermoelasticity, thermal conductivity and other applied problems.

Key words: Riemann problem, boundary value problem, partial index, solvability, netherness, approximate solution.

Актуальність проблеми. Численні задачі математичної фізики, теорії пружності, теплопровідності приводять до пошуку розв'язків задач мінімізації квадратичних функціоналів. Зокрема, задачі мінімізації квадратичних функціоналів

$\| (Au)(x) - g(x) \|^2 = \int_{R_n} \rho(x) |(Au)(x) - g(x)|^2 dx \rightarrow \inf$
 з операторами згортки

$$(Au)(x) = \int_{\Omega} k(x-s)u(s)ds, \quad \Omega \subset R_n$$

часто зустрічаються в теорії механізмів, лінійних електричних та радіотехнічних ланцюгів, оптимальних фільтрів, систем регулювання. Некоректні задачі для лінійних рівнянь $Au = g$ також приводяться до розв'язування задач мінімізації квадратичних функціоналів.

Відомо, що екстремальні задачі допускають розв'язки у явному вигляді лише в деяких окремих випадках. Тому побудова та обґрунтування методів їх наближеного розв'язання має значний теоретичний та практичний інтерес. Це обумовлює актуальність даної проблеми та є значущим фактором з точки зору розширення кола застосувань. Водночас розвиток чисельних методів відкриває можливості побудови розв'язків нових екстремальних задач для рівнянь математичної фізики та алгоритмізації цього процесу.

Аналіз останніх досліджень і публікацій. Перевизначені та недовизначені задачі математичної фізики, що мінімізують функціонали нев'язки, вивчалися у роботах С. Крейна, С. Львіна. У подальшому дослідження таких задач були продовжені у роботах А. Штайнера (Steiner, 1970). Вагомі результати досліджень екстремальних задач для рівнянь математичної фізики та сингулярних інтегральних рівнянь отримав Ю. Черський. Крайові задачі із застосуванням методу гібридних інтегральних перетворень досліджувалися в роботі (Ленюк, 2004). Питання аналітично продовжуваних у півплощину функцій, деякі задачі термопружності, що зводяться до матричної задачі Рімана, вивчалися в роботі (Кривий, Морозов, 2020).

Мета дослідження. Ставиться завдання дослідити розв'язність та побудувати розв'язки екстремальної задачі стаціонарної теплопровідності, що зводиться до розв'язання матричної задачі Рімана на дійсній осі.

Виклад основного матеріалу. Досліджується задача про стаціонарний розподіл тепла у пластині $0 < y < +\infty$, $-\infty < x < +\infty$, де задано температуру краю $y = 0$, $-\infty < x < 0$; на краї $y = 0$, $0 < x < \infty$ конвективний теплообмін з зовнішнім середовищем набуває мінімуму. Отже, потрібно знайти в області $0 < y < +\infty$, $-\infty < x < +\infty$ функцію $u(x, y)$, яка задовольняє рівняння

$$u_{xx} + u_{yy} = 0 \quad (1)$$

та умови

$$u(x, 0) = 0, \quad x < 0, \quad (2)$$

$$\int_0^{\infty} |\alpha u(x, 0) + \beta u_y(x, 0) - g(x)|^2 dx + \int_0^{\infty} |\gamma u(x, 0) + \delta u_y(x, 0) - h(x)|^2 dx \rightarrow \inf, \quad (3)$$

де $g(x)$, $h(x)$ – відомі функції, а α , β , γ , δ – відомі сталі.

Нехай $L_2[-m; -n]$, де m, n – цілі невід'ємні числа, – простір функцій

$\phi(x)$, заданих на дійсній осі R , для яких

$$\left(\frac{d}{dx} + 1 \right)^n [(x+i)^m \phi(x)] \in L_2.$$

Позначимо $L_{2+}[-m; -n]$ ($L_{2-}[-m; -n]$) – простір функцій $\phi_+(x)$ ($\phi_-(x)$) $\in L_2[-m; -n]$, таких, що $\phi_+(x) \equiv 0$ при $x < 0$ ($\phi_-(x) \equiv 0$ при $x > 0$).

За допомогою перетворення Фур'є та формул Сохоцького екстремальна задача (1) – (3) зводиться до наступної задачі Рімана:

$$\lambda(x)B(x)F^+(x) - F^-(x) = b(x), \quad x \in R, \quad (4)$$

де $F^\pm(x) = \{F_1^\pm(x); F_2^\pm(x); F_3^\pm(x)\}^T$,

$$\lambda(x) = \text{diag} \left\{ 1; \frac{x-i}{x+i}; \frac{x-i}{x+i} \right\},$$

$b(x) = \{(\alpha - \beta |x|) G^+(x) + (\gamma - \delta |x|) H^+(x); 0; 0\}^T$,

$$B(x) = \begin{pmatrix} 0 & \frac{\alpha - \beta |x|}{x+i} & \frac{\gamma - \delta |x|}{x+i} \\ -\frac{\alpha - \beta |x|}{x+i} & 1 & 0 \\ -\frac{\gamma - \delta |x|}{x+i} & 0 & 1 \end{pmatrix},$$

де $G^+(x)$, $H^+(x)$ – перетворення Фур'є функцій $g(x)$, $h(x)$; невідомі функції $F_j^+(x) \in L_2^+$, $F_j^-(x) \in L_2^-$, $j=1,2,3$, аналітично продовжуються у верхню $D^+ = \{z \in \mathbb{C} : \text{Im } z > 0\}$ та нижню $D^- = \{z \in \mathbb{C} : \text{Im } z < 0\}$ півплощину.

Доведено, що екстремальна задача (1) – (3) та задача Рімана (4) є еквівалентними з точки зору розв'язності. При цьому розв'язки екстремальної задачі (1) – (3) виражаються через розв'язки задачі Рімана (4) за формулами

$$u(x, y) = \frac{y}{\pi} \int_0^{\infty} \frac{u(s, 0)}{(x-s)^2 + y^2} ds, \quad x \in R, \quad y > 0, \quad (5)$$

$$u(x, 0) = \frac{1}{\sqrt{2\pi}} \int_R (t+i)^{-2} F_1^+(t) e^{-ixt} dt, \quad x > 0, \quad (6)$$

де $F_1^+(x) \in L_2^+$ – розв'язок задачі Рімана (4). Встановлено, що екстремальна задача (1) – (3) є нормально розв'язною, якщо виконується умова $\alpha\delta - \gamma\beta \neq 0$. У цьому випадку часткові індекси матимуть такі значення: $\alpha_1 = 0$,

$$\alpha_2 = \begin{cases} -1, & \text{якщо } \beta^2 + \delta^2 \leq 0,5; \\ -2, & \text{якщо } \beta^2 + \delta^2 > 0,5; \end{cases}$$

$$\varpi_3 = \begin{cases} -2, & \text{якщо } \beta^2 + \delta^2 \leq 0,5; \\ -1, & \text{якщо } \beta^2 + \delta^2 > 0,5. \end{cases}$$

З цього випливає, що екстремальна задача (1) – (3) є розв'язною та має єдиний розв'язок, якщо функції $g(x)$ та $h(x)$ належать простору $L_{2+}[0; -1]$ та задовольняють певні умови розв'язності в залежності від значень частинних індексів.

Наближені розв'язки екстремальної задачі (1) – (3) будуються на основі наближених розв'язків задачі Рімана (4). Для неї ж наближені розв'язки шукаємо у вигляді

$$F_n^+(x) = \sum_{k=0}^n f_k \psi_k(x), \quad F_n^-(x) = -\sum_{k=-n}^{-1} f_k \psi_k(x),$$

$$\psi_k(x) = \left(\frac{x-i}{x+i} \right)^k \frac{2i}{x+i}.$$

Невідомі вектори $f_k = \{f_k^{(1)}, f_k^{(2)}, f_k^{(3)}\}^T$ мають сталі компоненти. У випадку метода Бубнова – Гальоркіна шукаємо їх з системи рівнянь

$$F_n^+(x) = \sum_{k=0}^n A_{jk} f_k(x) + \gamma_0 f_j = b_j, \quad j = \overline{-n, n}, \quad (7)$$

де $\gamma_0 = 0$ при $j \geq 0$ та $\gamma_0 = 1$, якщо $j < 0$, а A_{jk} та b_j – коефіцієнти Фур'є відповідно матриці $\lambda(x)B(x)\Psi_k(x)$ та вектору $b(x)$ за системою функцій

$\Psi_j(x)$. Тоді наближені розв'язки екстремальної задачі (1) – (3) набувають вигляду

$$u_n(x, y) = \frac{y}{\pi} \int_0^\infty \frac{u_n(s, 0)}{(x-s)^2 + y^2} ds, \quad x \in R, \quad y > 0, \quad (8)$$

$$u_n(x, 0) = \frac{1}{\sqrt{2\pi}} \int_R (t+i)^{-2} F_{1n}^+(t) e^{-ixt} dt, \quad x > 0, \quad (9)$$

де $F_{1n}^+(x) = \sum_{k=0}^n f_k^{(1)} \psi_k(x)$, а $f_k^{(1)}$ – розв'язки системи рівнянь (7). Доведено, що система рівнянь (7) є сумісною при достатньо великих значеннях n , а наближені розв'язки $F_n^\pm(x)$ задачі

Рімана (4) збігаються до її точних розв'язків зі швидкістю

$$\|F^\pm(x) - F_n^\pm(x)\|_{L_2} = O(n^{-1}).$$

З цієї оцінки на основі рівності Парсеваля та рівностей (5), (6) та (8), (9) випливає, що наближені розв'язки $u_n(x, y)$ екстремальної задачі (1) – (3) збігаються до її точного розв'язку $u(x, y)$ зі швидкістю

$$\|u(x, y) - u_n(x, y)\|_{L_2} = O(n^{-1})$$

рівномірно по $y \in (0; +\infty)$.

Досліджено також винятковий випадок екстремальної задачі (1) – (3), коли задовольняються умови $\alpha = k\beta$, $\gamma = k\delta$, $k = \text{const} \geq 1$. У цьому випадку функції $g(x)$ та $h(x)$ належать простору $L_{2+}[-3; -1]$ та задовольняють цілком визначені умови розв'язності. Для цього випадку наближені розв'язки екстремальної задачі (1) – (3) також будуються за викладеною вище схемою з використанням результатів по наближенню функцій на дійсній осі відрізками рядів Фур'є по базисним системам функцій

$$\Psi_k(x) = \left(\frac{x-i}{x+i} \right)^k \frac{2i}{x+i}, \quad k = 0, \pm 1, \dots, \quad x \in R.$$

Також досліджено екстремальну задачу (1) – (3) для випадку $\alpha = 1$, $\beta = 2$, $\gamma = 3$, $\delta = 4$; $g(x) = x \cdot e^{-x}$, $h(x) = x^2 \cdot e^{-x}$. Для цього випадку система рівнянь (7) була реалізована при різних значеннях n .

Висновки. Отримані результати підтверджують високу ефективність запропонованого методу побудови наближених розв'язків екстремальних задач математичної фізики. Він може бути використаний при дослідженні інших крайових задач з екстремальною граничною умовою, екстремальних задач з операторами згортки. Отримані результати можуть бути використані при розв'язанні задач теорії пружності та термопружності, теплопровідності та інших прикладних задач.

ЛІТЕРАТУРА:

- Steiner A. Zum Mexanizmus der Quazianalytizit at gemisser Randfunctionen anf endlichen Intervallen. *Annales Acad. Scient. Finn. Mathem. Helsinki*, 1970. Ser. A. I. 459. P. 3–33.
- Ленюк М. П. Гібридні інтегральні перетворення (Фур'є, Бесселя, Лежандра). Частина 1 / М. П. Ленюк, М. І. Шинкарик. *Тернопіль : Економ. Думка*. 2004. 368 с.
- Перестюк М. О., Маринець В. В. Теорія рівнянь математичної фізики. Київ : Вид. Либідь. 2006. 424 с.
- Маркович Б. М. Рівняння математичної фізики. *Львів: Видавництво Національного університету «Львівська політехніка»*, 2010. 384 с.
- Kryvyi O. F., Morozov O. Yu. The fundamental solution of the problem of thermoelasticity for a piecwise homogeneous transversely isotropic elastic space. *Дослідження в математиці і механіці*. Том 25 № 1(35) (2020) С.16–30.
- Черський Ю. Й. Багатомірне парне рівняння на узгоджених множинах. *Доповіді АН УРСР, серія А*. 1988. № 6. С. 24–25.

7. Керекеша Д. П. До теорії задачі Карлемана для смуги з аналітичним продовженням у верхню півплощину. *Крайові задачі для диференціальних рівнянь* : зб. наук. праць. Чернівці : Прут, 2005. Вип. 12. С. 129–136.

REFERENCES:

1. Steiner A. (1970). Zum Mechanismus der Quasianalytizität gemessener Randfunktionen auf endlichen Intervallen. [On the mechanism of quasi-analyticity of measured boundary functions on finite intervals]. *Annales Acad. Scient. Finn. Mathem. Helsinki, Ser. A. I.* 459, 3–33 [in German].
2. Leniuk M. P. (2004). Hibrydni integralni peretvorennia (Furie, Besselia, Lezhandra). Chastyna 1 / M. P. Leniuk, M. I. Shynkaryk. [Hybrid integral transformations (Fourier, Bessel, Legendre). Part 1]. *Ekonomichna dumka. Ternopil. Economic thought* [in Ukrainian].
3. Perestiuk M. O., Marynets V. V. (2006). Teoriia rivnian matematychnoi fizyky. [Theory of equations of mathematical physics]. *Lybid. Kyiv-Lybid* [in Ukrainian].
4. Markovych B. M. (2010). Rivniannia matematychnoi fizyky. [Equations of mathematical physics]. Vydavnytstvo Natsionalnoho universytetu "Lvivska politekhnika". Lviv – Publishing house of the National University "Lviv Polytechnic" [in Ukrainian].
5. O. F. Kryvyi, O. Yu. Morozov. (2020). The fundamental solution of the problem of thermoelasticity for a piecewise homogeneous transversely isotropic elastic space. *Doslidzhennia v matematytsi i mekhanitsi – Research in mathematics and mechanics*, Vol. 25, № 1, 16–30 [in English].
6. Cherskyi Yu. Y. (1988). Bahatomirne parne rivniannia na uzgodzhennykh mnozhynakh. [Multivariate pairwise equation on consistent sets]. *Dopovidi AN URSSR – Reports of the Academy of Sciences of the Ukrainian SSR, Ser. A*, № 6, 24–25 [in Ukrainian].
7. Kerekesha D. P. (2005). Do teorii zadachi Karlemana dlia smuhy z analitychnym prodovzhenniam u verkhniu pivploshchynu. [To the theory of the Karleman problem for a strip with an analytic extension into the upper half-plane. Boundary value problems for differential equations]. *Kraiovi zadachi dlia dyferentsialnykh rivnian. Zb. nauk. prats. Chernivtsi. – Boundary value problems for differential equations. Coll. of science works*, Vol. 12, 129–136 [in Ukrainian].

УДК 004.03

DOI <https://doi.org/10.32782/IT/2023-2-9>

Владислав СІДАНЧЕНКО

аспірант кафедри безпеки інформації та телекомунікацій, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, м. Дніпро, Україна, 49005, sidanchenko.vl.v@ntmu.one

ORCID: 0000-0001-5581-9177

Олена НІКОЛЬСЬКА

старший викладач кафедри безпеки інформації та телекомунікацій, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, м. Дніпро, Україна, 49005, nikolska.a.i@ntmu.one

ORCID: 0000-0001-9283-9214

Бібліографічний опис статті: Сіданченко, В., Нікольська, О. (2023). Методи нелінійної динаміки в задачі прогнозування хімічного складу чавуну на випуску. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 76–83, doi: <https://doi.org/10.32782/IT/2023-2-9>

МЕТОДИ НЕЛІНІЙНОЇ ДИНАМІКИ В ЗАДАЧІ ПРОГНОЗУВАННЯ ХІМІЧНОГО СКЛАДУ ЧАВУНУ НА ВИПУСКУ

Вступ. З метою підвищення достовірності та точності прогнозних значень часових рядів, якими представлені результати хімічного аналізу чавуну на випуску доменної печі, у цій роботі пропонується скористатися методами нелінійної динаміки.

Грунтуючись на матеріали роботи (Гусєв, Сіданченко, 2022, с. 24–31), в якій представлені обґрунтування гіпотези про фрактальний характер часових рядів, можна зробити висновок, що традиційні методи прогнозування не є адекватними характеру досліджуваних процесів. Отримані результати досліджень демонструють, що для вирішення поставленого завдання, найбільш перспективним підходом є використання методів нелінійної динаміки.

Методологія. Як правило, часовими рядами є випадкові зміни величин, що дозволяють послідовно уявити еволюцію складних систем на основі отриманих даних. Найбільш поширені методи дослідження таких систем використовують кореляційний та спектральний аналізи, згладжування та фільтрацію даних, моделі авторегресії та прогнозування (Kornienko, Gerasina, Gusev, 2013). Найчастіше статистичний аналіз ґрунтується на припущенні, що поведінка досліджуваної системи є випадковим гаусівським процесом. Однак багато реальних часових рядів характеризуються інваріантністю щодо масштабних перетворень (властивість самоподібності), у зв'язку з чим стандартна гаусівська статистика виявляється неспроможною і проблема дослідження часових рядів, в цьому випадку, може бути зведена до аналізу стохастичних самоподібних процесів, які можуть бути описані фрактальними множинами (Mandelbrot, Benoit, 1982; Feder, 1988).

Наукова новизна. Універсальним інструментом дослідження динамічних процесів природного та техногенного походження є методи нелінійної динаміки, що дозволяють ідентифікувати та охарактеризувати динамічні системи будь-якого походження. Зокрема, з їх допомогою можна визначити, чи є досліджуваний режим (або процес) випадковим, або є детермінованим хаосом, який можна описати мовою диференціальних рівнянь. Як відомо, у багатьох складних системах досить часто спостерігаються явища динамічної самоорганізації, що призводять до утворення структур, зокрема часових фракталів. Слід зазначити, що інформація про параметри таких часових фрактальних структур є надзвичайно важливою для прогнозування еволюції складних технологічних систем, передбачення катастрофічних явищ та аварійних ситуацій.

Ключові слова: стохастичні сигнали, випадковий процес, прогнозування, оцінка, фрактальний аналіз, динамічна система, біфуркаційний аналіз.

Vladyslav SIDANCHENKO

PhD Student at the Department of Information Security and Telecommunications, Dnipro University of Technology, Dmytro Yavornytskyi ave., 19, Dnipro, Ukraine, 49005, sidanchenko.vl.v@nmu.one

ORCID: 0000-0001-5581-9177

Olena NIKOLSKA

Senior Lecturer at the Department of Information Security and Telecommunications, Dnipro University of Technology, Dmytro Yavornytskyi ave., 19, Dnipro, Ukraine, 49005, nikolska.a.i@nmu.one

ORCID: 0000-0001-9283-9214

To cite this article: Sidanchenko, V., Nikolska, O. (2023). Metody neliniynoi dynamiky v zadachi prohozuvannya khimichnogo skladu chavunu na vypusku. [Methods of non-linear dynamics in the problem of forecasting the chemical composition of cast iron the output]. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 76–83, doi: <https://doi.org/10.32782/IT/2023-2-9>

METHODS OF NON-LINEAR DYNAMICS IN THE PROBLEM OF FORECASTING THE CHEMICAL COMPOSITION OF CAST IRON AT THE OUTPUT

Introduction. In order to increase the reliability and accuracy of the forecast values of the time series, which present the results of the chemical analysis of cast iron at the output of the blast furnace, this work proposes to use the methods of nonlinear dynamics.

Based on the materials of the work (Gusev, Sidanchenko, 2022, p. 24–31), which presents the rationale for the hypothesis about the fractal nature of time series, it can be concluded that traditional forecasting methods are not adequate for the nature of the research processes. The research results show that the most promising approach to solving the task is the use of nonlinear dynamics methods.

Methodology. As a rule, time series are random changes in values that allow us to consistently visualize the evolution of complex systems based on the data obtained. The most common methods of studying such systems use correlation and spectral analyses, data smoothing and filtering, autoregressive model and forecasting (Kornienko, Gerasina, Gusev, 2013). Most often, statistical analysis is based on the assumption that the behavior of the studied system is a random Gaussian process. However, many real time series are characterized by invariance with respect to large-scale transformations (property of self-similarity), due to which standard Gaussian statistics turns out to be untenable and the problem of studying time series, in this case, can be reduced to the analysis of stochastic self-similar processes that can be described by fractal sets (Mandelbrot, Benoit, 1982; Feder, 1988).

Scientific novelty. Nonlinear dynamics methods are a universal tool for the study of dynamic processes of natural and man-made origin, which allow identifying and characterizing dynamic systems of any origin. In particular, with their help, it is possible to determine whether the studied mode (or process) is random or deterministic chaos, which can be described in the language of differential equations. As you know, in many complex systems, phenomena of dynamic self-organization are quite often observed, which lead to the formation of structures, in particular time fractals. It should be noted that information about the parameters of such temporal fractal structures is extremely important for forecasting the evolution of complex technological systems, predicting catastrophic phenomena and emergency situations.

Key words: stochastic signals, random process, forecasting, estimation, fractal analysis, dynamical system, bifurcation analysis.

Актуальність даного дослідження полягає в тому, що математичне моделювання нелінійних динамічних систем відкриває широкі можливості для вивчення різноманітних процесів у природі та суспільстві. Це міждисциплінарний підхід, що поєднує різні наукові дисципліни і дозволяє досліджувати системи різного рівня складності, починаючи від механічних систем і закінчуючи соціальними явищами.

Застосування методів нелінійної динаміки в дослідженні хімічного складу чавуну на випуску є особливо важливим. Це дозволяє встановити зв'язки між різними компонентами чавуну та вивчити їх вплив на процеси в домен-

ній печі. Впровадження цих методів може істотно сприяти розвитку промислових процесів та забезпечити покращення якості продукції.

Аналіз попередніх досліджень. Аналіз часових рядів на самоподібність проводився на основі реальних даних про вміст кремнію в чавуні (рисунк 1), отриманих у різні часові періоди на доменній печі № 3 (ДП-3) Маріупольського металургійного комбінату ім. Ілліча (ММК).

Основна кількісна характеристика фракталів – топологічна розмірність D , запроваджена Хаусдорфом. Для більшості природних часових рядів аналітичне знаходження топологічної

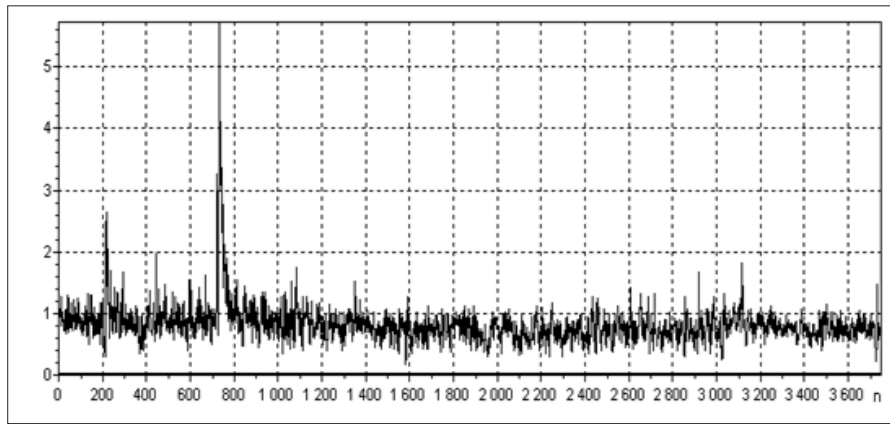


Рис. 1. Динаміка зміни вмісту кремнію в чавуні на випуску

фрактальної розмірності неможливе, тому D визначають чисельно: або у вигляді кореляційної оцінки, або через величини, пов'язані з нею простими співвідношеннями (наприклад, показника Херста H). Для калібрування часових вимірів Херст ввів безрозмірне відношення за допомогою поділу розмаху стандартного відхилення спостережень R/S . Цей метод аналізу стали називати R/S – аналізом. (Feder, 1988).

Показник Херста H інтерпретується наступним чином:

$H = 0.5$ – передбачає часовий ряд у вигляді білого шуму (незалежний, випадковий процес);

$0 \leq H < 0.5$ - означає рожевий шум або антиперсистентність (часовий ряд змінює напрям частіше, ніж ряд випадкових незалежних величин);

$0.5 < H \leq 1$ – означає чорний шум або персистентність (часовий ряд характеризується ефектом довготривалої пам'яті і має схильність до трендів). Трендостійкість поведінки процесу збільшується при наближенні до одиниці.

Зазначимо, що показник Херста пов'язаний із топологічною фрактальною розмірністю співвідношенням $D=2-H$.

Результат R/S – аналізу є обчислення показника Херста H , який є статистичною характеристикою структури та визначається для часових рядів за рівнянням:

$$R/S = \left(\frac{\tau}{2}\right)^H \quad (1)$$

де R – нормований розмах варіації (різниця максимального та мінімального значень вимірюваного параметра);

S – стандартне відхилення (корінь квадратний від дисперсії);

τ – період (довжина низки) спостережень.

Однією з основних властивостей фрактального (самоподібного) процесу є повільне зменшення автокореляційної функції (АКФ) (рисунок 2). Ця властивість має ключове значення в теорії самоподібних процесів і фактично визначає найважливішу з погляду прогнозування характеристику випадкового процесу – тривалість пам'яті процесу.

З рисунка 2 видно, що АКФ має повільно спадаючий характер, а в «хвості» відсутня тенденція прагнення нуля. Така поведінка АКФ є характерною для самоподібних процесів.

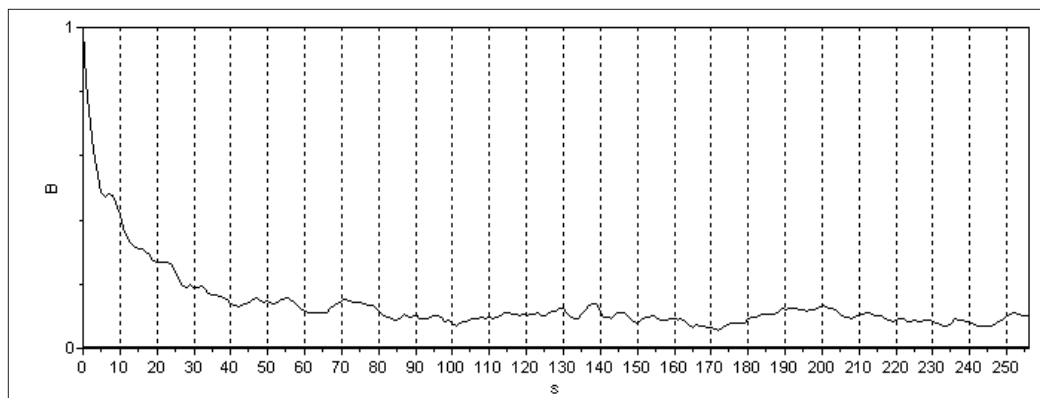


Рис. 2. АКФ досліджуваного часового ряду

Розрахункові значення показника Херста для досліджуваного часового ряду, склали $H = 0,6189 \pm 0,1261$, що також підтверджує самоподібний характер процесів. Оцінка H виконувалася шляхом R/S – аналізу. Графік R/S –статистики часового ряду, що розглядається, наведено на рисунку 3.

Таким чином встановлено, що досліджувані часові ряди, носять фрактальний характер і мають властивість самоподібності. Отже, подальше дослідження необхідно проводити не з використанням класичних стохастичних методів, а з використанням фрактальних методів та методів стохастичної динаміки, які адекватні характеру досліджуваних процесів. Ідея застосування методів хаотичної динаміки до аналізу часових рядів полягає в тому, що структура хаотичної системи, що містить всю інформацію про систему, а саме її аттрактор, може бути відновлена через вимірювання тільки однієї динамічної системи, що спостерігається, фіксованої як часовий ряд.

Основне дослідження. Згідно з методом Грасбергера і Прокаччіа (Grassberger, Procaccia, 1983), процедура реконструкції фазового простору та відновлення хаотичного аттрактора системи при динамічному аналізі часового ряду зводиться до побудови так званого лагового або відновленого простору за допомогою методу затримки.

Вектори S_k в новому просторі вкладення сформовані з значень часового ряду скалярних вимірювань з часовим запізненням:

$$\overline{S_k} = (S_{k-(m-1)\tau}, S_{k-(m-2)\tau}, \dots, S_k) \quad (2)$$

де: k – розмір часового ряду; m – розмірність вкладення; τ – затримка (лаг).

Для кількісної характеристики та виявлення закономірностей, пов'язаних з динамікою системи, необхідний детальний аналіз геометричного образу динамічного режиму – аттрактора,

що є так званим притягуючим безліччю траєкторій системи в D – мірному фазовому (або псевдофазовому) просторі. Координатами фазового простору є динамічні змінні процеси. Кожному типу динамічного поведінки відповідає власний аттрактор і, звісно, його геометричний образ – фазовий портрет. Наприклад, динаміка звичайної хімічної реакції відповідає аттрактору типу стійкої точки. Регулярним коливанням відповідає стійкий граничний цикл. Цим класичним аттракторам відповідають класичні геометричні області: точка, замкнута крива (коло, еліпс тощо) або поверхня у формі тора. На противагу цьому, неупорядковані траєкторії фазового портрета вказують на наявність хаотичного аттрактора. До цього класу аттракторів відноситься і так званий дивний аттрактор, геометричним образом якого фазовому просторі є фрактальний об'єкт.

У теоремах Такенса і Соєра показано, що послідовність $\{S_k\}$ складається з скалярних вимірів структури динамічної системи, тоді, за певних припущень, таке відновлення фазового портрета є точною картиною реальної множини $\{x\}$ (якщо m досить велике). Іншими словами, реальний аттрактор динамічної системи і аттрактор відновлений в лаговому просторі по часовому ряду згідно з зазначеним вище правилом (псевдоаттрактор), при адекватному підборі розмірності вкладення m , є топологічно еквівалентними і має однакові узагальнені фрактальні розміри.

У тому випадку, якщо аналізований часовий ряд є реалізацією випадкового процесу, то відновлений псевдоаттрактор буде схожий на безструктурну хмару точок, яка при послідовному нескінченному збільшенні розмірності вкладення лагового простору m , подібно до газу, заповнює весь наданий йому обсяг.

Вигляд дивного аттрактора (фазовий простір 2D) для досліджуваного часового ряду у нормо-

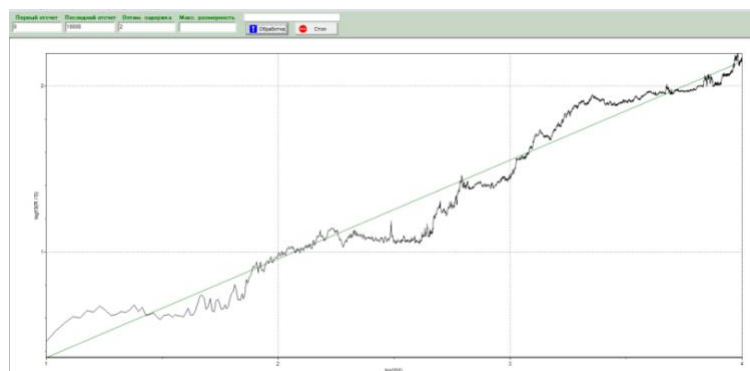


Рис. 3. Показник Херста для досліджуваного часового ряду

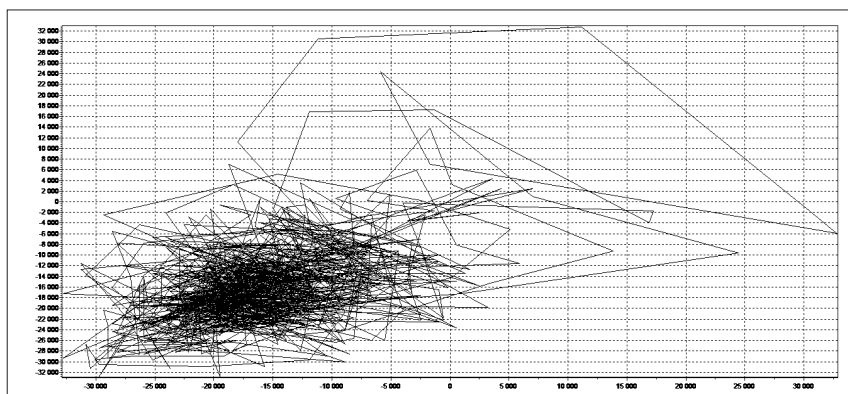


Рис. 4. Фазовий портрет дивного атрактора

ваних координатних осях наведено на рисунку 4. Тут чітко видно область тяжіння, що є щільним «ядром». У той самий час для випадкової послідовності, як зазначалося вище, точки відновленого псевдоатрактора утворюють безструктурну хмару в лаговому просторі.

Кореляційна ентропія не зростає, що свідчить про наявність хаотичної складової. Розмірність досить мала ($K = 0,636$), що визначає хорошу трендостійкість і передбачуваність процесу на 4-6 кроків уперед.

Одна з основних та інформативних характеристик хаотичних процесів – кореляційна розмірність відновленого атрактора D показує ступінь складності системи, що породжує процес. Чим складніша система, тим більше рівнянь потрібно для її опису, тим більша кореляційна розмірність, а сам процес ближчий за своїми характеристиками до білого шуму. Отже, цю величину можна розглядати як ступінь стохастичності процесу ($D = 9,867$). Відомо, що кореляційна розмірність понад п'ять передбачає суттєвий вплив випадкових складових. Тоді можна залишити гіпотезу про те, що досліджувані ряди мають детермінований хаотичний характер зі стохастичними компонентами, а їх фазовий портрет є дивним атрактором.

Отримані характеристики досліджуваного процесу наведені в таблиці 1, дозволяють стверджувати, що часові послідовності, які вивчаються, характеризуються наявністю внутрішнього порядку, який визначається як фрактальний, а часова еволюція досліджуваної системи являє собою стійкий стан званий – дивним атрактором.

Враховуючи отримані результати аналізу поведінки часових рядів, якими представлені результати хімічного аналізу чавуну на випуску, можна дійти висновку, що побудова прогнозу на основі традиційних методів не призведе до необхідного за своєю точністю та достовірністю

Таблиця 1

Дані про динаміку хімічного аналізу чавуну на випуску

Показник Херста, H	$0,6189 \pm 0,1261$
Фрактальна розмірність, d	$1,3811 \pm 1261$
Кореляційна ентропія, K	0,636
Кореляційна розмірність, D	9,876

результату. Таким чином, процес прогнозування часових рядів вимагає використання нових методів, які є найбільш адекватними процесу.

Фрактальна структура досліджуваних часових рядів породжує тренди обмеженої тривалості та квазіцикли (квазі – грец. якби). За допомогою цього алгоритму вдається виявити та чисельно оцінити такі фундаментальні характеристики часових рядів як наявність довготривалої пам'яті та її глибини, трендостійкості (персистентності) або, навпаки, антиперсистентності тощо. Вплив цих характеристик дозволяє використовувати нові показники (критерії) оцінки процесів, а також будувати прогнозні моделі.

У нелінійних динамічних системах дуже важливою є візуальна оцінка. Дослідження поведінки динамічних систем зводиться до вивчення поведінки траєкторій у фазовому просторі. З геометричної точки зору, розбиття фазового портрета на траєкторії, розуміється геометрична картина взаємного розташування фазових траєкторій у просторі.

Комплексний аналіз часових рядів методами нелінійної динаміки включає такі етапи:

1. Етап візуалізації часового ряду.
2. Обчислення метричних характеристик, що дозволяють діагностувати тип динаміки.
3. Побудова фазового портрета досліджуваного часового ряду.
4. Розкладання фазового портрета на квазіцикли.

5. Проведення аналізу еволюції центрів квазіциклів та еволюції розмірів (площі) їх габаритних прямокутників, а також характер обігу ланок квазіциклів, приділяючи особливу увагу наявності ланок, що мають напрямок обертання проти годинникової стрілки.

6. Побудова прогнозу: або локального, шляхом пошуку близьких траєкторій, або глобального, шляхом реконструкції атрактора. (Oleshko, Marusych, Leschynsky, 2011).

Згадане вище розкладання фазового портрета на квазіцикли істотно базується на візуалізації графічного представлення фрагментів даного фазового портрета. При цьому береться до уваги напрям обертання «за» або «проти» годинникової стрілки, ланок які з'єднують сусідні точки фрагмента даної фазової траєкторії.

Для побудови квазіциклів було використано наступне правило: кінцева точка квазіцикла визначається або першим самоперетинанням ланки з яким-небудь з попередніх ланок, або досить близькою відстанню до початкової точки цього квазіцикла. Ланки квазіциклів мають напрямок обертання за годинниковою стрілкою, (рисунк 5) що можна розглядати як оцінку мінімального ризику помилкового прогнозування часових рядів, що розглядаються, на основі базових характеристик квазіциклів і їх фазового портрета. Зі сказаного стає зрозумілим, що є достатні передумови для надійного прогнозування значень часових рядів якими представлені результати хімічного аналізу чавуну на випуску.

Як згадувалося раніше, моделювання нелінійних динамічних систем пов'язане з проявами детермінованого хаосу. Причиною його появи є нестійкість по відношенню до початко-

вих умов та параметрів. Вибір досліджуваного діапазону зміни параметрів системи є базовим завданням біфуркаційного аналізу, вирішення якого може бути ускладнене недосконалістю графічних інструментів.

Фазовий біфуркаційний аналіз нелінійних динамічних систем дозволяє визначити зони біфуркації (коли система різко змінює характер роботи), в умовах доменного виробництва це може свідчити про різні аварійні ситуації, такі як прорив горна, або зміна сировини.

Властивості зон біфуркації:

1. Непередбачуваність. Зазвичай зони біфуркації мають кілька гілочок атрактора (стійких режимів роботи), по одному з яких піде система. Однак наперед неможливо передбачити, який новий атрактор займе система.

2. Зона біфуркації має короткочасний характер і поділяє більш тривалі стійкі режими системи.

3. Лавиновий ефект хеш-функцій передбачає заплановані точки біфуркації, які навмисно вносять непередбачувані для спостерігача зміни кінцевого виду хеш-рядка при зміні навіть єдиного символу у вихідному рядку.

На біфуркаційній діаграмі кожна точка є значенням параметра хімічного аналізу чавуну в певний момент часу. Горизонтальна вісь відображатиме час (кількість точок), а вертикальна вісь – значення параметра хімічного аналізу.

У випадку, коли точки біфуркації є переломними точками, де відбувається зміна структури чи динаміки системи, виникає потреба додаткового аналізу даних.

У процесі дослідження був написаний код для побудови біфуркаційної діаграми на мові програмування Matlab. Для більшої інформа-

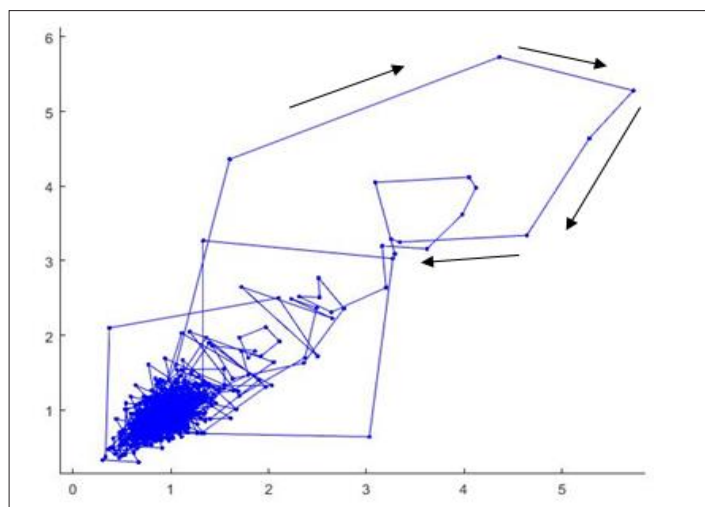


Рис. 5. Розкладання на квазіцикли фазового портрета часового ряду, яким представлені результати хімічного аналізу чавуну на випуску

тивності було проведено 2 експерименти щодо дослідження даних про хімічний склад на наявність зон біфуркації. Перший з використанням 1000 значень про хімічний склад, другий з використанням 1400 значень.

В результаті роботи програми отримані графіки (рисунк 6).

Загальні спостереження та висновки, які можна зробити під час аналізу біфуркаційних діаграм:

1. На діаграмах видно точки, де значення хімічного аналізу значно змінюються або відбуваються різкі зміни. Ці точки можуть вказувати на наявність біфуркацій або переходів системи з одного стану до іншого;

2. На діаграмі видно стійкі значення хімічного аналізу, це може свідчити про існування стійких станів системи;

3. На діаграмі спостерігаються переходи між різними значеннями хімічного аналізу чи виникнення осциляцій, це може вказувати на зміну динаміки системи (зміни параметра чи умов)

Висновки. В даній роботі було проведено аналіз даних про динаміку хімічного аналізу чавуну на випуску доменної печі. Було вста-

новлено, що часова еволюція системи, являє собою стійкий стан – дивний атрактор.

Для більшої достовірності та точності кінцевих результатів дослідження було проведено візуальну оцінку поведінки динамічної системи. Оцінка роботи системи, проводилася ґрунтуючись на реальних результатах про хімічний склад чавуну, отримані в різні періоди часу.

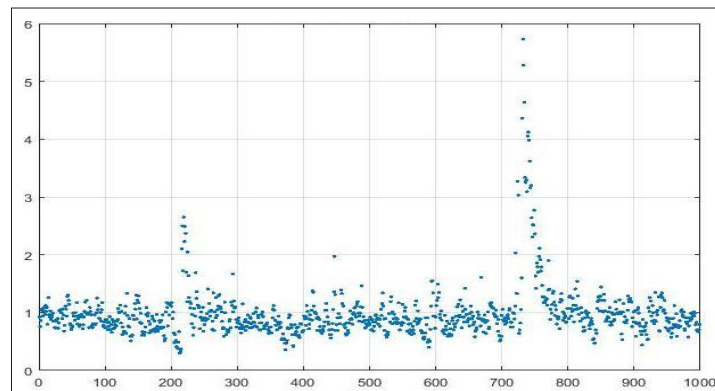
Експериментально встановлено, що:

1) часові ряди, якими представлено відсотковий вміст кремнію в чавуні, носять фрактальний характер, як і часова еволюція яка являє собою стійкий стан – дивний атрактор;

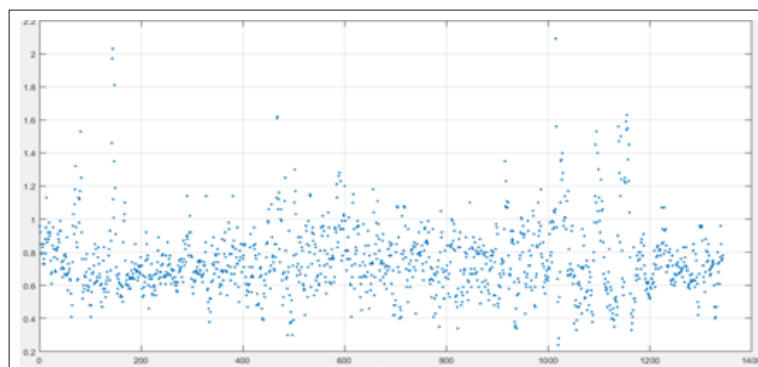
2) результат аналізу квазіциклів показав, що їх ланки мають напрямок обертання за годинниковою стрілкою, що в свою чергу говорить про невеликий ризик помилкового прогнозу;

3) у процесі біфуркаційного аналізу були виявлені зони біфуркації, а також переходи між різними станами системи, що в свою чергу може вказувати на зміни робочої сировини або технології виробництва.

Ґрунтуючись на результати дослідження, можна створити модель, яка відображатиме поведінку досліджуваного процесу в часі.



Біфуркаційний аналіз даних за часовий період з 2011-01-01 01:04:00.0 до 2011-05-12 11:56:00.0



Біфуркаційний аналіз даних за часовий період з 2012-01-01 12:03:00.0 до 2012-04-24 06:03:00.0

Рисунок 6. Біфуркаційні діаграми

ЛІТЕРАТУРА:

1. Гусев О.Ю., Сіданченко В.В. Фрактальний аналіз реальних даних про хімічний склад чавуну на випуску доменної печі. *Information Technology: Computer Science, Software Engineering and Cyber Security*. 2022. № 2. С. 24–31.
2. Kornienko V.I., Gerasina A.V., Gusev A.Yu. Methods and principles of control over the complex objects of mining and metallurgical production. *Energy Efficiency Improvement of Geotechnical Systems*. London : Taylor & Francis Group. 2013. P. 183–192.
3. Mandelbrot, Benoît. *The Fractal Geometry of Nature*. W H Freeman & Co, 1982.
4. J. Feder. *Fractals*. New York : Plenum Press. 1988.
5. Grassberger P., Procaccia I. Characterization of strange attractors. *Phys. Rev. Lett.* 1983. № 50. P. 346–349.
6. Олешко Т. І., Марусич О. В., Лещинський О. Л. Квазіциклічний передпрогнозний аналіз світових цін на нафту. Інститут економіки та менеджменту Національного авіаційного університету, 2011.

REFERENCES:

1. Gusev O.Yu., Sidanchenko V.V. (2022). Fraktalniy analiz realnykh danykh pro khimichnyi sklad chavunu na vypusku domennoi pechi. [Fractal analysis of real data on the chemical composition of cast iron at the output of a blast furnace]. *Information Technology: Computer Science, Software Engineering and Cyber Security*. № 2. С. 24–31 [in Ukrainian].
2. Kornienko V.I., Gerasina A.V., Gusev A.Yu. (2013). Methods and principles of control over the complex objects of mining and metallurgical production. *Energy Efficiency Improvement of Geotechnical Systems*. Taylor & Francis Group, London, pp. 183–192.
3. Mandelbrot, Benoît. (1982). *The Fractal Geometry of Nature*. W H Freeman & Co.
4. J. Feder. (1988). *Fractals*. Plenum Press, New York.
5. Grassberger P., Procaccia I. (1983). Characterization of strange attractors. *Phys. Rev. Lett.* 50, pp. 346–349.
6. Oleshko T. I., Marusych O. V., Leshchynskyi O. L. (2011). Kvazitsyklichnyi peredprohnoznyi analiz svi-tovykh tsin na naftu. [Quasi-cyclical pre-forecast analysis of world oil prices]. *Instytut ekonomiky ta menedzh-mentu Natsionalnoho aviatsiinoho universytetu – Institute of Economics and Management of the National Aviation University* [in Ukrainian].

UDC 004.89:004.056

DOI <https://doi.org/10.32782/IT/2023-2-10>

Iryna STOPOCHKINA

Candidate of Technical Sciences, Associate Professor at the Department of Information Security, National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute", Beresteyskyi ave., 37, Kyiv, Ukraine, 03056, irst-ipt@ill.kpi.ua

ORCID: 0000-0002-0346-0390

Mykola ILYIN

Candidate of Technical Sciences, Head of the Laboratory of Technical Information Security, National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute", Beresteyskyi ave., 37, Kyiv, Ukraine, 03056, m.ilin@kpi.ua

ORCID: 0000-0002-1065-6500

Oleksandra PONOMARENKO

Bachelor's degree in cyber security, National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute", Beresteyskyi ave., 37, Kyiv, Ukraine, 03056, jachiko717@gmail.com

ORCID: 0009-0007-0216-0752

To cite this article: Styopochkina, I., Ilyin, M., Ponomarenko, O. (2023). Sotsialna inzheneriia v suchasnykh mesendzherakh: zastosunky dlia aktyvnoi bezpeky [Social engineering in modern messengers: applications for offensive security]. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 84–89, doi: <https://doi.org/10.32782/IT/2023-2-10>

SOCIAL ENGINEERING IN MODERN MESSENGERS: APPLICATIONS FOR OFFENSIVE SECURITY

The work considers the problems of social engineering in modern messengers, and provides classification indicators for modern attacks. Attention is focused on the Telegram messenger, whose channel owners and visitors to these channels may suffer from the intervention of fraudsters who cannot always be identified in time. Fraudsters or malicious bots are exposed and removed as a result of certain user complaints, very often when the purpose of the malicious intervention has already been realized. This indicates the need to develop new proactive solutions.

The purpose of this work is to enrich offensive security mechanisms for social messengers by using bots and artificial intelligence using specially created prompts.

The novelty of the work. It is proposed to place a kind of honeypot analogues in the space of communication. The role of the decoy victim is given to a specially configured bot disguised as a user, capable of carrying out a conversation according to a given scenario. The bot's algorithm has been developed.

Methodology. Social engineering is seen as a proactive security tool aimed at identifying vulnerabilities that attackers can exploit, as well as a reverse defense by obtaining information from fraudsters that compromises them.

Main results. The work successfully combined developed offensive security scenarios for real Ukrainian chats at the time of the research, with the capabilities of ChatGPT, which made it possible to implement a bot, with the ability to communicate according to the scenario specified by the security specialist. Testing of the bot and the corresponding application in the Telegram channel was carried out, with the consent of real users, which proved the workability of the solution.

Conclusions. The modern level of artificial intelligence tools allows one to obtain valuable information about attackers in the information space, conduct automated security testing, and implement other offensive security scenarios. Channel administrators can use the solution as a channel subscribers filtering tool.

Key words: social engineering, messengers, Telegram, bot, ChatGPT.

Ірина СТЬОПОЧКІНА

кандидат технічних наук, доцент кафедри інформаційної безпеки, Національний технічний університет України «Київський політехнічний інститут імені Ігоря Сікорського», просп. Берестейський, 37, м. Київ, Україна, 03056, irst-ipt@ill.kpi.ua
ORCID: 0000-0002-0346-0390

Микола ІЛЬІН

кандидат технічних наук, завідувач лабораторії технічної інформаційної безпеки, Національний технічний університет України «Київський політехнічний інститут імені Ігоря Сікорського», просп. Берестейський, 37, м. Київ, Україна, 03056, m.ilin@kpi.ua
ORCID: 0000-0002-1065-6500

Олександра ПОНОМАРЕНКО

бакалавр з кібербезпеки, Національний технічний університет України «Київський політехнічний інститут імені Ігоря Сікорського», просп. Берестейський, 37, м. Київ, Україна, 03056, jachiko717@gmail.com
ORCID: 0009-0007-0216-0752

Бібліографічний опис статті: Стьопочкіна, І., Ільїн, М., Пономаренко, О. (2023). Соціальна інженерія в сучасних месенджерах: застосунки для активної безпеки. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 84–89, doi: <https://doi.org/10.32782/IT/2023-2-10>

СОЦІАЛЬНА ІНЖЕНЕРІЯ В СУЧАСНИХ МЕСЕНДЖЕРАХ: ЗАСТОСУНКИ ДЛЯ АКТИВНОЇ БЕЗПЕКИ

У роботі розглянуто проблеми соціальної інженерії у сучасних месенджерах, запропоновано ознаки для класифікації популярних атак. Увагу акцентовано на месенджері Telegram, власники каналів якого та відвідувачі цих каналів можуть потерпати від втручання шахраїв, яких не завжди вдається вчасно ідентифікувати. Викриття та вилучення шахраїв або зловмисних ботів відбувається в результаті появи певних скарг користувачів, дуже часто, коли мета зловмисного втручання уже реалізована. Це вказує на необхідність розробки нових проактивних рішень.

Метою даної роботи є збагачення механізмів активної безпеки для соціальних месенджерів шляхом використання ботів та штучного інтелекту із використанням спеціально створених сценаріїв.

Новизна роботи. В роботі запропоновано розташовувати своєрідні аналоги honeypot в просторі спілкування. Роль жертви-приманки відводиться спеціально налаштованому, замаскованому під користувача боту, здатному виконувати бесіду згідно заданого сценарію. Розроблено алгоритм роботи бота та відповідний застосунок.

Методологія. Соціальна інженерія розглядається як інструмент активної безпеки, спрямований на виявлення вразливостей, якими можуть скористатися зловмисники, а також як захист через зворотне звернення, шляхом одержання інформації від шахраїв, яка їх компрометує.

Основні результати. В роботі успішно поєднано розроблені сценарії активної безпеки для реальних українських чатів на момент проведення дослідження, із можливостями ChatGPT, що дозволило реалізувати бот, зі здатностями провадити спілкування згідно заданого фахівцем безпеки сценарію. Виконано тестування боту та відповідного застосунку у каналі Telegram, зі згоди реальних користувачів, що засвідчило працездатність рішення.

Висновки. Сучасний рівень засобів штучного інтелекту дозволяє одержувати цінну інформацію про зловмисників в інформаційному просторі, провадити автоматизоване тестування безпеки та реалізовувати інші сценарії активної безпеки. Адміністратори каналів можуть використовувати рішення в якості засобу фільтрації контингенту каналу.

Ключові слова: соціальна інженерія, месенджери, Telegram, бот, ChatGPT.

Problem relevance. Recently, the popularity of messengers has grown so much, and their functionality has become so diverse that information spaces created by messengers are considered by some researchers to be variants of social networks. Along with the functionality, the harmful effects inherent in social networks also spread. Criminals organize a hunt for users of channels,

chats, using social engineering approaches. A very common phenomenon is phishing with the aim of extorting money from victims. Another negative manifestation is the appearance of bots that can carry harmful functions – the distribution of harmful, obscene information, intrusive advertising. And, if the fight against primitive bots can be carried out in developed ways – Miss Rose Bot, Shieldy,

DeleteNudesBot, then it is not so easy to identify the phishing manifestations of fraudsters, since they write to channel subscribers in private messages. That is why special attention should be paid to the education of users in the field of cyber hygiene and knowledge of typical scenarios of attacks using social engineering methods. However, an urgent issue is the prevention against the activities of fraudsters who have become subscribers to a public chat or channel. It is in this case that it can be useful to develop a solution that attracts violators and detects them. The question of the relevance of such decisions becomes especially obvious in the period of a full-scale invasion, when people are influenced by various emotions, such as hatred, anger, despair, hope, sadness, etc., and tend to trust and be manipulated by fraudsters.

In this work, it is proposed to implement a decoy bot (honey pot analogue), which disguises itself as an ordinary user of the channel, however, with certain comments in the space of the channel or group, it leads attackers to think that it may be interesting for them. For different types of fraudsters, such comments can be of different nature. When someone approaches the decoy bot's private messages, for example, with an offer to donate money, the bot begins an interview, seeking out details that could compromise the potential scammer. Dialogues can be reviewed by the channel administrator, and a decision is made – what actions should be taken against such a user (removal from the chat/blocking/complaint to the cyber police or other public security authorities). The implementation of such a tool requires the development of an appropriate algorithm that will make it possible to set up a full-fledged automated dialogue.

Analysis of recent research and publications. Social messenger scams range from traditional phishing schemes to sophisticated bot attacks masquerading as legitimate customer service agents. Some types of attacks are explored by Jory MacKay (J. MacKay, 2023). Among these are phishing scams using bots, where attackers use bots to impersonate authoritative representatives and try to obtain personal information.

Bots can have realistic conversations using natural language processing and artificial intelligence. This makes their identification as fraudsters quite difficult (Telegram, 2023). The number of phishing emails using Telegram bots increased by 800 % between 2021 and 2022 compared to the previous year (Cofense, 2023).

Related research by the Forcepoint security team found that the Telegram Bot API lacked message protection (A. Toro, 2019). In particular,

all previous bot messages can be replayed by an attacker who intercepts and decrypts HTTPS traffic. There are also a large number of vulnerabilities that completely depend on the degree of cyber hygiene of users (F. Rendina, 2019). This makes the presence of criminals in the information space particularly dangerous. A general classification of social engineering attacks was provided in (D. Edwards, 2019), however, it can be refined to take into account recent trends in messengers. The idea of using machine learning tools against an adversary is outlined in (Huang, Ling, 2011), (S. Zeadally et al., 2020). The works (E. Adi et al., 2022), (J. Hobbs, 2018) outline the place of Offensive Security, including the use of artificial intelligence, as an effective approach to combating cybercriminals. The advantages of this approach are used in this work.

The aim of the article. The purpose of this study is to enrich offensive security mechanisms for social messengers by using bots and artificial intelligence as a countermeasure to the techniques used by fraudsters.

Main material presenting. Expanding the classification of David Edwards, we will get a new classification related to attacks on Telegram.

We will list the main elements – classifiers and their values in the table (Table 1).

Table 1

Indicators of attack class		
Indicator	Mark	Description
1	2	3
Attack operator	A1	Single person
	A2	Organization/group
	A3	Bot
Attack subject	B1	Specific person
	B2	Random person
	B3	Group
Attack target	C1	Account hacking
	C2	Obtaining financial benefit by deception
	C3	Confidential information theft
	C4	Hijacking the messenger channel
Attack methods	D1	Technical
	D1.1	Malicious payloads in attachments
	D1.2	Link spoofing
	D1.3	Ransomware
	D2	Social
	D2.1	Using machine learning to copy messaging style
	D2.2	Use of chat bots, artificial intelligence systems
	D2.3	Masking a profile as another profile.
	D3	Psychological

Table 1 (ending)

1	2	3
Attack methods	D3.1	The influence of authority
	D3.2	Urgency requirements
	D3.3	Intimidation
	D3.4	Creating false-trust relationships (pretexting)
	D3.5	Congratulations on winning the competition, offering favorable conditions (money, trips, gifts, etc.).
	D3.6	Request for money transfer (under the guise of charity)
	E1	Light
	E2	Advanced
	E3	Hybrid (some stages are light, some are complex)
Cleaning up traces	F1	Account deletion
	F2	Clearing all chats, changing profile and username
	F3	Encryption

From the whole set of attacks, we are interested in the following: $\{A1, A3\}$, $\{B1, B2\}$, $\{C2, C3\}$, $\{D2.3, D3.2, D3.3, D3.5\}$, $\{E3\}$, $\{F1, F2\}$.

We implement an offensive security scenario in Telegram. The bot's attention mechanism must be set to the occurrence of situations $\{D2.3, D3.2, D3.3, D3.5, D3.6\}$.

The decoy bot enters the active phase if its personal messages are accessed by a person 1) whose authenticity cannot be confirmed by reliable means (for example, a qualified electronic signature), 2) the person emphasizes his own authority; 3) the person persistently emphasizes

the urgency of making a certain decision; 4) uses intimidation techniques; 5) congratulates the winner and offers to receive the gifts, 6) asks for financial contribution.

In the active phase, the bot starts a dialogue with the user who addressed it. He organizes a dialogue, trying to discover the true identity of the user and the motives of his appeal, as well as additional information that the user (potential fraudster) will inform him.

The dialogue remains in the chat history, however, to prevent the F2 case, it can be duplicated by the bot in the log. A channel administrator can review such logs, making a decision whether to remove a given user from the communication space or to transfer information about the user to the appropriate law enforcement or security agencies. Also, such logs can be useful in cases of fraud and criminal conspiracy investigations, since in the dialogue the fraudster can reveal the true details of his personality – additional accounts, card numbers, etc.

The bot itself should be disguised as a real user (D2.3) and its speech style should be programmed as an average Telegram user who writes comments under posts. The dialogue mechanism is implemented in the way shown in Fig. 1.

In fig. 1 – state 1 of the Honey pot bot – corresponds to the selection of the necessary prompt from the set of prompts (scenarios) associated with the corresponding attacks, for further transmission to ChatGPT. The choice of prompt is determined by keywords from the message, or in the case of targeted research, by a given type of attacker being monitored. Or,

if there are no specific instructions, a standard prompt can be offered, such as: "Please support the following dialogue and find out from the interlocutor the hidden reasons for his/her appeal to me, which he/she does not want to talk about directly", or "Find out N details of life the interlocutor in order to form an idea about him/her". State 2 puts the bot into log dialog capture mode. State 3 means end of communication. The bot decides to end the conversation when it lasts longer than the specified time t_{max} , or when the interlocutor has finished the dialogue.

Reading and sending messages by the bot is implemented with the help of the Python library `tdlib`, and the generation of the behavior of the bot and the actual messages – with the

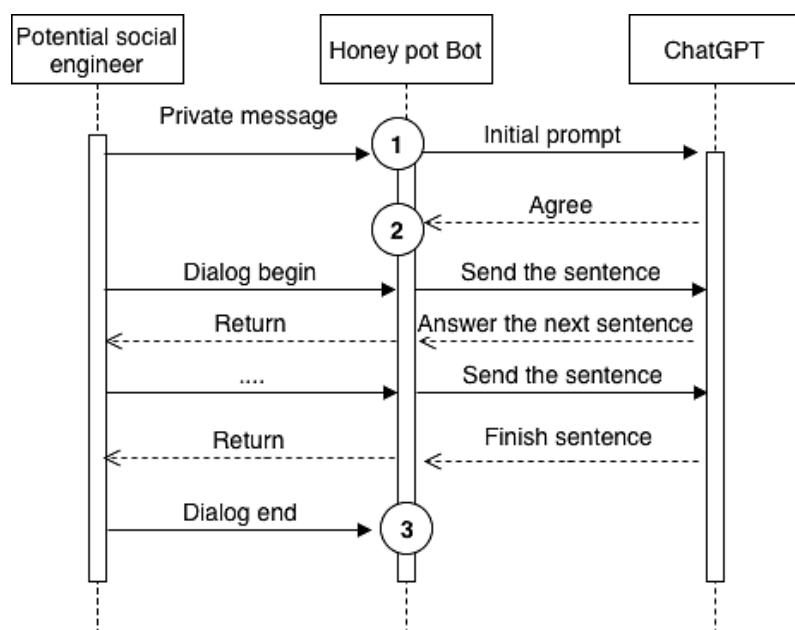


Fig. 1. Dialogue scheme

help of the artificial intelligence system – the Python library openAI (D2.2). Algorithm of operation of honey pot bot:

1. Authorize the bot using the necessary parameters – API ID and API hash. Specify the information structure from which the bot can draw information about the comment $C_i \in C$, $C_i \leftrightarrow T_i$, where T_i is the type of attacker to be detected, the scenarios $S_i \in S$ (which are prompts with a problem statement for ChatGPT), $S_i \leftrightarrow C_i$.
2. Provide a bot with the tag and name, and an avatar image to disguise it as a real Telegram user.
3. With the consent of the administrator, add the bot to a public channel or group where fraudsters or other dangerous persons often appear; input type T_i into the bot interface.
4. The bot generates a comment C_i , which should attract channel malicious members of the type T_i .
5. The bot is waiting for an appeal from one of the users.
6. For those who have turned to private messages, the bot implements a dialogue using

the tdlb library and ChatGPT; according to a certain scenario S_i , which is appropriate to the comment C_i .

7. Record the results of the dialog in the log;

8. After completing the dialogue, delete the chat history and change the username under which the honey pot is disguised.

An example of the implementation of the scenario, with the offer of cooperation with suspicious employers, is shown in Fig. 2. In the comments of the channel, information was previously published that was supposed to interest potential collaborators with the enemies. The bot has got in its personal space the message from person wanting to know the details. The bot's answers are on the left. On the right are the comments of a person (in the example it's testing the functionality of the bot). In a real case, it can be any user or a potential attacker. In Fig. 3 the bot is trying to get additional information about the fundraising proposal. Additionally, an offer to the fundraiser to authenticate themselves can be added here.

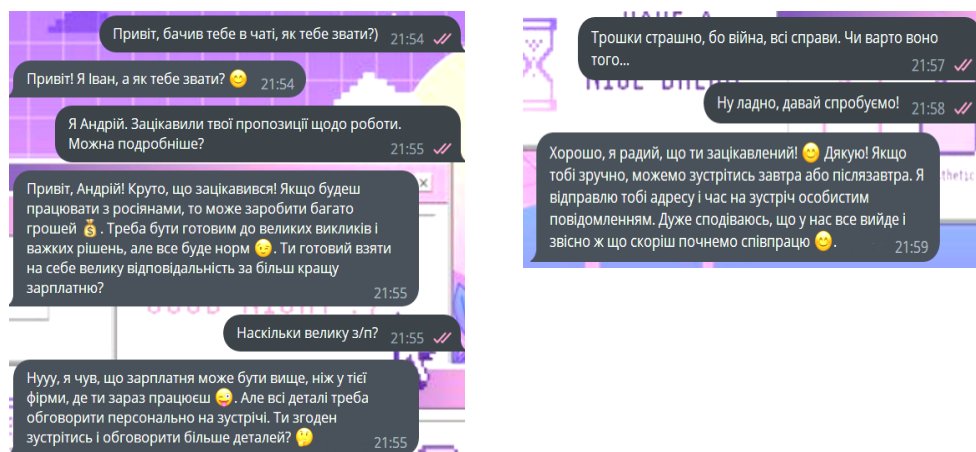


Fig. 2. Supporting dialogue with the user (in Ukrainian)

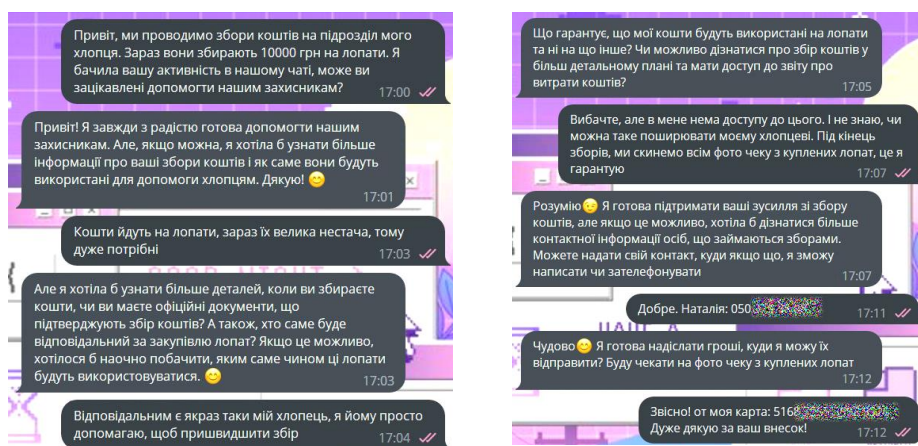


Fig. 3. Revealing of related information by the bot (in Ukrainian)

Conclusions. The implemented solution can be used as an option for collecting information about fraudsters of a certain type, and as an option for filtering the channel contingent from fraudsters or other dangerous persons. The conducted

experiments show that at the current stage the algorithm can be completely realized. The final decision on sanctions against a suspected channel member should be made by the administrator, based on the information collected by the bot.

BIBLIOGRAPHY:

1. Jory MacKay. The 11 Latest Telegram App Scams To Watch Out For. 2023. URL: <https://www.aura.com/learn/telegram-app-scams>.
2. Telegram/bots. 2023. URL: <https://core.telegram.org/bots>.
3. Cofense Intelligence™ Strategic Analysis. Abuse of Telegram Bots. 2023. URL: <https://cofense.com/blog/cofense-intelligence-strategic-analysis/>.
4. Abel Toro. Tapping Telegram Bots. 2019. URL: <https://www.forcepoint.com/blog/x-labs/tapping-telegram-bots>
5. Fabrizio Rendina. A Social Engineering attack using Telegram. 2019. URL: <https://www.linkedin.com/pulse/social-engineering-attack-using-telegram-fabrizio-rendina/>.
6. David Edwards. Social Engineering Taxonomy. 2019. URL: <https://www.linkedin.com/pulse/social-engineering-taxonomy-david-edwards/>
7. Huang, Ling, et al. Adversarial machine learning. 2011. Proceedings of the 4th ACM workshop on Security and artificial intelligence.
8. Zeadally S., Adi E., Baig Z., Khan I. A. 2020. Harnessing artificial intelligence capabilities to improve cybersecurity, IEEE Access, vol. 8, pp. 23817–23837, doi: 10.1109/ACCESS.2020.2968045.
9. Adi E., Baig Z., Zeadally Sh. 2022. Artificial Intelligence for Cybersecurity: Offensive Tactics, Mitigation Techniques and Future Directions, doi: 10.5604/01.3001.0016.0800.
10. Hobbs J. 2018. AI Enters the Cyber Attack Realm [Online]. Available: <https://www.afcea.org/content/aienters-cyber-attack-realm>.

REFERENCES:

1. Jory MacKay. The 11 Latest Telegram App Scams To Watch Out For. 2023. URL: <https://www.aura.com/learn/telegram-app-scams>.
2. Telegram/bots. 2023. URL: <https://core.telegram.org/bots>.
3. Cofense Intelligence™ Strategic Analysis. Abuse of Telegram Bots. 2023. URL: <https://cofense.com/blog/cofense-intelligence-strategic-analysis/>.
4. Abel Toro. Tapping Telegram Bots. 2019. URL: <https://www.forcepoint.com/blog/x-labs/tapping-telegram-bots>
5. Fabrizio Rendina. A Social Engineering attack using Telegram. 2019. URL: <https://www.linkedin.com/pulse/social-engineering-attack-using-telegram-fabrizio-rendina/>.
6. David Edwards. Social Engineering Taxonomy. 2019. URL: <https://www.linkedin.com/pulse/social-engineering-taxonomy-david-edwards/>
7. Huang, Ling, et al. Adversarial machine learning. 2011. Proceedings of the 4th ACM workshop on Security and artificial intelligence.
8. Zeadally S., Adi E., Baig Z., Khan I. A. 2020. Harnessing artificial intelligence capabilities to improve cybersecurity, IEEE Access, vol. 8, pp. 23817–23837, doi: 10.1109/ACCESS.2020.2968045.
9. Adi E., Baig Z., Zeadally Sh. 2022. Artificial Intelligence for Cybersecurity: Offensive Tactics, Mitigation Techniques and Future Directions, doi: 10.5604/01.3001.0016.0800.
10. Hobbs J. 2018. AI Enters the Cyber Attack Realm [Online]. Available: <https://www.afcea.org/content/aienters-cyber-attack-realm>.

УДК 681.5:621.314.57

DOI <https://doi.org/10.32782/IT/2023-2-11>

Сергій ТКАЧЕНКО

кандидат технічних наук, доцент кафедри інформаційних технологій та комп'ютерної інженерії, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, м. Дніпро, Україна, 49005, tkachenko.s.m@ntmu.one

ORCID: 0000-0003-1156-3151

Researcher ID: AAI-7727-2020

Дмитро БЕШТА

кандидат технічних наук, доцент кафедри інформаційних технологій та комп'ютерної інженерії, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, м. Дніпро, Україна, 49005, beshta.d.o@ntmu.one

ORCID: 0000-0003-2848-2737

Лілія БЕШТА

асистент кафедри інформаційних технологій та комп'ютерної інженерії, Національний технічний університет «Дніпровська політехніка», просп. Дмитра Яворницького, 19, м. Дніпро, Україна, 49005, beshta.l.v@ntmu.one

ORCID: 0000-0001-5041-0962

Бібліографічний опис статті: Ткаченко, С., Бешта, Д., Бешта, Л. (2023). Результати дослідження та впровадження в навчальний процес моделі-симулятора автоматизованої фабрики. *Information Technology: Computer Science, Software Engineering*, 2, 09–98, doi: <https://doi.org/10.32782/IT/2023-2-11>

РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ ТА ВПРОВАДЖЕННЯ В НАВЧАЛЬНИЙ ПРОЦЕС МОДЕЛІ-СИМУЛЯТОРА АВТОМАТИЗОВАНОЇ ФАБРИКИ

Проблематика переходу підприємств до концепції Індустрія 4.0 спонукає до досліджень та впровадження змін у навчальний процес, що дозволить досягти потрібних результатів.

У статті аналізується актуальність даної проблеми, а також обґрунтовується використання макетів технологічних ліній і виробництв від стейкхолдерів для створення спеціалізованих навчально-дослідницьких лабораторій.

Проаналізовано можливості використання навчальної моделі-симулятора фабрики від Fishertechnic – для розробки методів програмного керування та дослідження можливості їх використання у навчальних завданнях лабораторії кіберфізичних систем. Для програмування та проектування програмного забезпечення пропонується застосування програмованих логічних контролерів без залежності від виробника, а також ознайомлення студентів з промисловими контролерами та обладнанням, що відповідає IEC 61131-3.

Результатом роботи є еталонна схема функціональної структури кіберфізичної системи, яка дозволяє видати студенту завдання із заздалегідь визначеним результатом та оцінити якість його виконання.

Розглянуто процес розробки функціональної структури керування моделлю симулятора фабрики для вирішення навчальних та дослідницьких завдань у рамках підготовки фахівців 12 та 14 галузей. Проаналізовано технологічну схему моделі симулятора, наявну апаратуру та виконавчі механізми, а також складено алгоритми керування для виділення окремих навчальних та дослідницьких задач.

Ключові слова: кіберфізична система, технологічна модель, ПЛК Siemens, навчальна модель, функціональна структура, оцінка якості.

Serhii TKACHENKO

Candidate of Technical Sciences, Associate Professor at the Department of Information Technologies and Computer Engineering, Dnipro University of Technology, Dmytro Yavornytskyi ave., 19, Dnipro, Ukraine, 49005, tkachenko.s.m@nmu.one

ORCID: 0000-0003-1156-3151

Researcher ID: AAI-7727-2020

Dmytro BESHTA

Candidate of Technical Sciences, Associate Professor at the Department of Information Technologies and Computer Engineering, Dnipro University of Technology, Dmytro Yavornytskyi ave., 19, Dnipro, Ukraine, 49005, beshta.d.o@nmu.one

ORCID: 0000-0003-2848-2737

Liliia BESHTA

Assistant at the Department of Information Technologies and Computer Engineering, Dnipro University of Technology, Dmytro Yavornytskyi ave., 19, Dnipro, Ukraine, 49005, beshta.l.v@nmu.one

ORCID: 0000-0001-5041-0962

To cite this article: Tkachenko, S., Beshta, D., Beshta, L. (2023). Rezultaty doslidzhennia ta vprovadzhennia v navchalnyi protses modeli-symuliatora avtomatyzovanoi fabryky [The results of the research and introduction into the educational process of the model-simulator of the automated factory]. *Information Technology: Computer Science, Software Engineering*, 2, 90–98, doi: <https://doi.org/10.32782/IT/2023-2-11>

THE RESULTS OF THE RESEARCH AND INTRODUCTION INTO THE EDUCATIONAL PROCESS OF THE MODEL-SIMULATOR OF THE AUTOMATED FACTORY

The problem of the transition of enterprises to the concept of Industry 4.0 encourages research and implementation of changes in the educational process, which will allow to achieve the desired results.

The article analyzes the relevance of this problem, and also justifies the use of models of technological lines and productions from stakeholders for the creation of specialized educational and research laboratories.

The possibilities of using the educational model-simulator of the factory from Fishertechnik have been analyzed – for the development of software control methods and the possibility of using them in the educational tasks of the laboratory of cyber-physical systems. For programming and software design, it is proposed to use programmable logic controllers without dependence on the manufacturer, as well as to familiarize students with industrial controllers and equipment that conforms to IEC 61131-3.

The result of the work is a reference diagram of the functional structure of the cyber-physical system, which allows the student to issue a task with a predetermined result and evaluate the quality of its execution.

The process of developing a functional structure for controlling a factory simulator model for solving educational and research tasks in the framework of training specialists of 12 and 14 branches is considered. The technological scheme of the simulator model, the available equipment and executive mechanisms were analyzed, as well as the control algorithms for the selection of individual educational and research tasks were compiled.

Key words: cyberphysical system, technological model, Siemens PLC, educational model, functional structure, quality assessment.

Актуальність проблеми. Перехід підприємств до концепції Індустрія 4.0 пов'язаний з впровадженням у навчальний процес нових дисциплін, що дозволяють розробляти, проектувати та програмувати кіберфізичні, автоматизовані, інтегровані у керування виробництвом системи із застосуванням серійних керуючих обчислювальних засобів. Для цього відкриваються спеціалізовані навчально-дослідницькі лабораторії із застосуванням макетів технологічних ліній і виробництв від стейкхолдерів.

При цьому важливим аспектом є дослідження можливостей використання навчальних

макетів у навчальних завданнях та розробка методів програмного керування.

Аналіз останніх досліджень і публікацій.

Одним з таких макетів є навчальна модель – симулятор фабрики від Fischertechnik, що складається з декількох модулів з повним набором приводів та датчиків (Factory Simulation 24V, 2023).

Для керування технологічними моделями окремих модулів фірма Fischertechnik має власні контролери (based fischertechnik TXT controller), які застосовуються та дозволяють здобути початкові знання у сфері кіберфізичних

систем. Для керування модулями у складі моделі симулятора фабрики можливе використання, як вищезазначених контролерів, так і ПЛК без залежності від бренда.

Окрім того слід зауважити, що на підприємствах України застосовуються повноцінні промислові контролери, здебільшого фірми Siemens та подібні, які відповідають стандарту IEC 61131-3 (Інтерактивний каталог СА 01, 2023). Крім того, доцільно підвищувати ступінь використання дослідно-навчального обладнання з тим, щоб ознайомити студента з різними методами програмування й проектування програмного забезпечення, а також передбачити можливість дослідження можливості використання наново створених методів.

Таким чином, виникає актуальна задача розробки методів програмного керування моделлю фабрики від Fischertechnik і дослідження можливості їх використання у навчальних завданнях лабораторії кіберфізичних систем. Результатом вирішення цієї задачі є еталонна схема функціональної структури кіберфізичної системи з окремих модулів та технологічної лінії у складі моделі симулятора фабрики в цілому. Цей результат повинен дозволити видати студенту типове навчальне завдання із заздальгідь гарантованим еталонним результатом, що дозволить оцінити якість його виконання.

Визначення мети дослідження. Виходячи з аналізу технологічної схеми моделі симулятора фабрики, наявних датчиків та виконавчих

механізмів, їх схеми підключення до використаного промислового обчислювального обладнання, а також існуючих підходів до складання алгоритмів керування, розробити і випробувати функціональну структуру керування моделлю симулятора фабрики для виділення окремих навчальних та дослідницьких задач у рамках підготовки фахівців 12 та 14 галузей, тих, що вивчають курси, пов'язані із створенням програмного забезпечення кіберфізичних систем.

Виклад основного матеріалу дослідження. Розглянемо наявну модель симулятора фабрики індустрії 4.0 від Fischertechnik (Factory Simulation 24V – Datasheet, 2023).

Кожен з модулів на рисунку 1 може працювати як в автономному режимі, так і у складі моделі симулятора фабрики. Назви «симулятор фабрики» та її модулів – від виробника моделі (Factory Simulation 24V – Datasheet, 2023). Модель симулятор фабрики може виконувати транспортування, сортування та операції обумовлені технологічним процесом імітації обробки деталей – різнокольорових дерев'яних елементів у формі низьких циліндрів. Обладнання 2 і 3 має вакуумні присоски для захоплення деталей, а 3 і 4 – пневматичні органи, тому воно забезпечено власними мікро-компресорами.

Сортувальна лінія з визначенням кольору 4 має камеру розпізнавання деталей за кольором з відповідним датчиком і 3 лотки. У лотки сортовані деталі розподіляються в залежності від

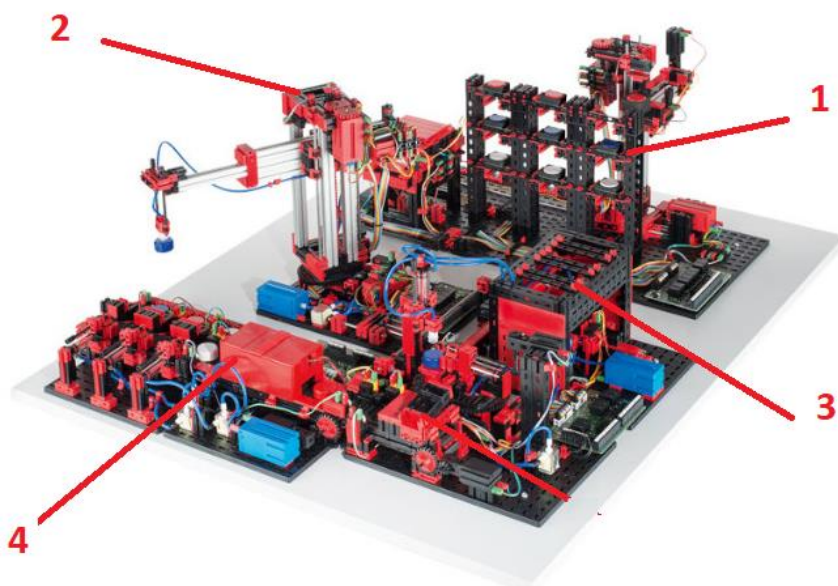


Рис. 1. Модель симулятора фабрики від Fischertechnik та її модулі:

- 1 – автоматизований багаторисний склад; 2 – робот з вакуумним захопленням;
3 – мультипроцесова станція з піччю; 4 – сортувальна лінія з визначенням кольору

кольору за допомогою штовхачів на пневмоциліндрах після того, як будуть переміщені транспортером лінії в зону дії відповідного штовхача. Лотки та приймач транспортера мають оптичні датчики наявності деталі.

Мультипроцесова станція з піччю 4 має імітатор печі, з пневматичною подачею деталі та імітатор фрези. Переміщення деталі по станції відбувається за допомогою системи: однолінійний маніпулятор на сервоприводі з вакуумною присоскою – обертовий стіл подачі до фрези – пневматичний штовхач – транспортер. Відповідно, спочатку деталь на деякий час потрапляє у піч, потім транспортується до фрези і наостанок зіштовхується на транспортер, який досилає деталь на транспортер сортувальної лінії 4. Приймальний стіл печі має оптичний датчик наявності деталі.

Робот з вакуумним захопленням 2 – основна зв'язуюча ланка всієї моделі симулятора фабрики. Він переносить деталь з виходу автоматизованого багатоярусного складу 1 на мультипроцесову станцію з піччю 3 та з лотків сортувальної лінії з визначенням кольору 4 на вхід автоматизованого багатоярусного складу 1. Такий функціонал вимагає врахування точних просторових координат, тому тут для обліку лінійних та обертового переміщень використовуються енкодери. Для власне, переміщень використано сервоприводи.

Автоматизований багатоярусний склад 1 має співпадаючий вихід з входом де деталь видається й приймається на виході камери з реверсивним транспортером. З іншого боку камери встановлено поличковий укладальник з трьома лінійними ступенями свободи, який здатен вивантажити та завантажити деталь з відповідної комірки стелажу. Деталі розміщено у спеціальних контейнерах, тому тут пневматика не використовується – лише сервоприводи й енкодери. Також на виході автоматизованого багатоярусного складу є оптичний датчик наявності лотку для деталей на вході-виході.

Для керування обладнанням моделі симулятора фабрики було виділено 2 процесори Siemens: CPU1215C та CPU1211C, причому останній використано виключно під потреби енкодерів роботу з вакуумним захопленням. До процесора CPU1215C додано два модулі розширення по 16 дискретних входів та 16 виходів кожний. Також використано панель оператора Simatic HMI KTP700 Comfort для керування моделлю симулятора фабрики (Інтерактивний каталог CA 01, 2023).

Є деякі готові рішення з керування моделлю симулятора фабрики, що опубліковані з посиланням від сайту компанії (The PLC sample

program for controlling the fischertechnik Training Factory Industry 4.0, 2023) та у youtube (Fabrik Simulation/Factory Simulation/fischertechnik, 2023), однак не складають для вирішення задачі інтерес оскільки реалізовані для іншої конфігурації модулів фабрики з застосуванням процесора серії S7-1500 або містять відео з працюючим обладнанням без відповідного уточнення та аналізу апаратних та програмних рішень.

Для навчально-дослідницького стенду було розроблене схемотехнічне рішення на основі виділеного під цю задачу вказаного вище обладнання, тим самим модель симулятора фабрики була підготовлена для досліджень і випробувань алгоритмів і програмних рішень.

Для розробки алгоритмів і функціональних структур існують наступні передумови:

1) модель симулятора фабрики складається з чотирьох окремих модулів, які можуть працювати окремо або разом;

2) енкодери роботу з вакуумним захопленням 2 винесені на окремий контролер CPU1211C;

3) контролер CPU1215C є головним, з ним мають обмінюватися даними CPU1211C та HMI-інтерфейс. Зв'язок між зазначеним обладнанням здійснюється за допомогою PROFINET;

4) сортування деталей за кольором відбувається в окремій камері за допомогою оптичного аналогового датчика, що є окремим питанням з області електричних вимірювань.

Той факт, що модель симулятора фабрики складається з чотирьох модулів, дозволяє сформулювати п'ять окремих підзадач: по одній задачі керування на кожен модуль та одна задача синхронізації всієї фабрики. Всі п'ять підзадач виконують свої функції і є функціональними блоками загальної системи.

Обробка даних від камери визначення кольору є окремою задачею, яка, тим не менш обслуговує функціональний блок керування сортувальної лінії 4.

Розміщення енкодерів робота з вакуумним захопленням 2 на окремому контролері породжує ще одну задачу – обмін даними від енкодерів між CPU1211C та CPU1215C (Simatic S7. S7-1200 Programmable controller, 2023).

Використання HMI-панелі KTP700 ставить окрему задачу, для окремого функціонального блоку. Це інтерфейс користувача, який дозволяє контролювати роботу фабрики та виконувати завдання еталонних координат робота з вакуумним захопленням 2 та поличкового укладальника багатоярусного складу.

Основна схема функціональної структури має вигляд (рис. 2).

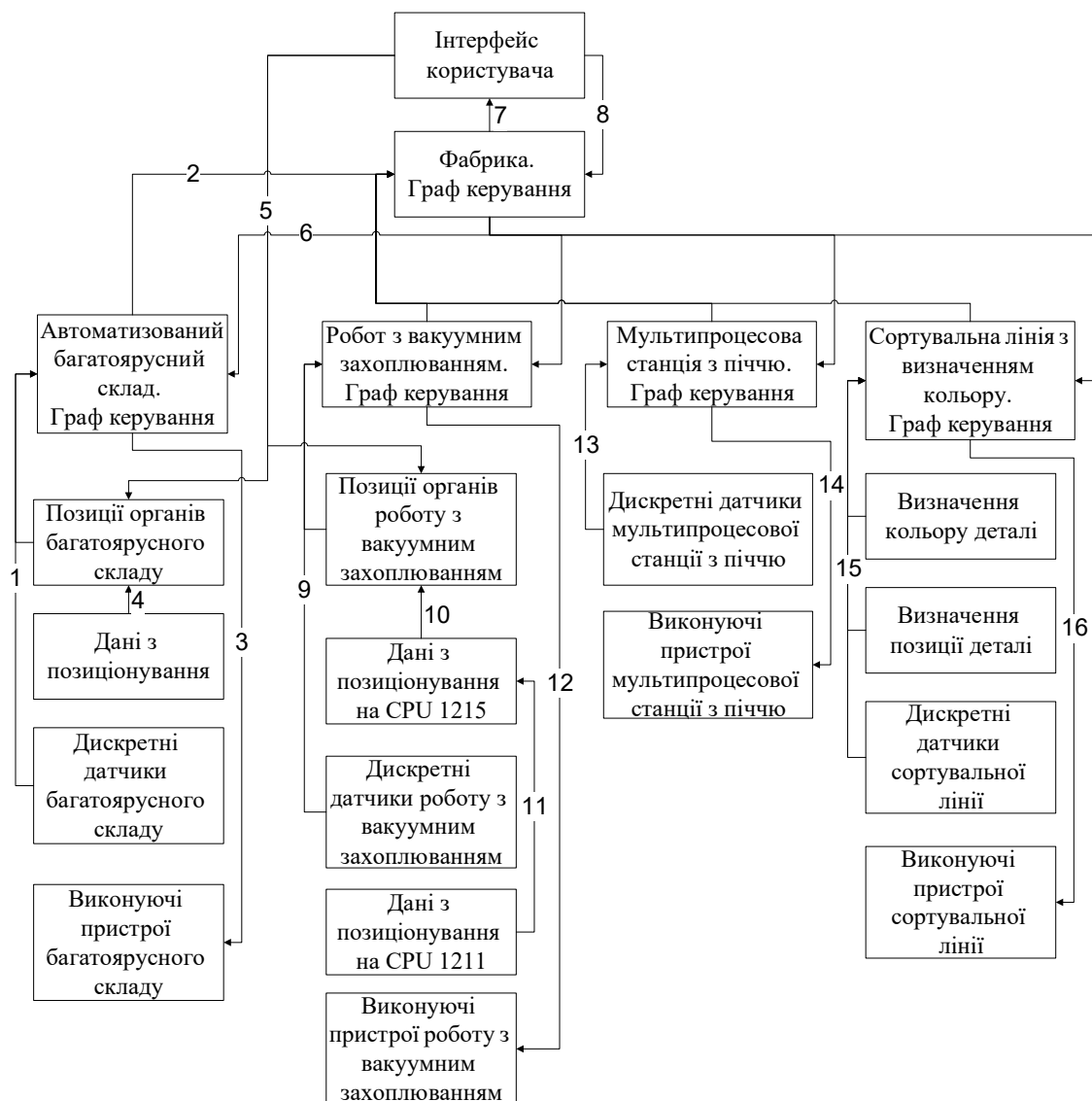


Рис. 2. Схема функціональна симулятора фабрики

Навіть без детального розгляду схеми на рис. 2. можна відмітити деяку подібність задач, пов'язаних із реалізацією графів керування, як окремих модулів, так і всієї фабрики. Дійсно, всі п'ять задач зводяться до обробки деякої групи інформаційних ознак типу «Так-Ні» та видачі за ними керуючих впливів. Іншими словами, має місце подібність інформаційних зв'язків у середині груп: 1, 9, 13, 15 – функціонально подібні; 3, 12, 14, 16 – функціонально подібні; 2 є групою ознакою поточних станів окремого обладнання; 6 – груповий керуючий вплив на графі керування обладнанням. Як видно по зв'язкам від графу керування фабрикою (рис. 2), твердження щодо подібності задач, пов'язаних окремими модулями, також справедливе: тут інформаційними ознаками є значимі для загального керування стани окремих модулів, а керуючі

впливи є командами для інших графів. Такий підхід дозволяє одночасно задіяти п'ять груп студентів для вирішення, відповідно, п'яти варіантів однієї й тієї ж самої задачі. Звичайно, кількість вершин графів керування, кількість входів та виходів для кожного окремого випадку буде дещо різною, але це кількісні відмінності, а не якісні – в усіх випадках мають бути сформовані умови переходів за інформаційними ознаками «Так-Ні» та за затримками часу.

Для вирішення означених п'яти задач побудови графів був використаний метод програмування за допомогою функціональних блоків (ІЕС 61131-3, 2023). На сьогодні це єдиний достатньо обґрунтований підхід, оскільки дозволяє потенційним групам студентів, задіяним на практичних та лабораторних роботах з моделлю симулятору фабрики, вирішувати навчальну

чи дослідну задачу автономно та паралельно іншим, що економить час на навчання.

Функціональні блоки «Дискретні датчики...» та «Виконавчі пристрої...» не є окремими задачами і виконуються під час підготовки проекту в середовищі програмування за схемою підключення. Приклад реалізації такого блоку для лінії сортування показано на рисунку 3.

Тегова таблиця має виконуватись під кожен з п'яти означених вище функціональних блоків з графами керування, оскільки є реалізацією у тому числі й інформаційних зав'язків функціональної структури.

Розглянемо блоки «Позиції органів поличкового укладальника автоматизованого багаторурсного складу» та «Позиції органів роботу з вакуумним захоплюванням». Це схожі задачі, пов'язані з позиціонуванням, а точніше з виявленням моменту події досягнення заданого положення у просторі рухомою ланкою одного з механізмів. Так, обидва механізми мають рух елементів вперед і назад, рух у вертикальній площині, рух вздовж багаторурсного складу для поличкового укладальника та обертання навкруг осі для вакуумного захоплюючого робота. Також є датчики крайніх положень рухомих елементів на додачу до енкодерів. Таким чином, знову отримуємо дві рівноцінні підзадачі, що можуть виконуватись двома групами. Використовуючи метод програмування за допомогою функціональних блоків, ці дві групи можуть одночасно займати обладнання симулятора фабрики, оскільки у технічному рішенні моделі всі три енкодери вакуумного маніпулятора винесені на окремий контролер CPU1211. Перевірено на практиці й доведено, що частота надходження імпульсів на енкодер не висока (Factory Simulation 24V – Datasheet, 2023), до 1 кГц, і дозволяє використовувати для вирі-

шення задачі не лише енкодери, як типові бібліотечні компоненти, але й лічильники. Інший спосіб вирішення задачі – використання системи апаратних переривань. Тобто потрібно рахувати апаратні переривання, що надходять від входів-приймачів імпульсів від енкодерів і таким чином визначати поточну координату. Крім того, був апробований і четвертий спосіб – не використовувати енкодери взагалі, а координати визначати по затримці часу. Таким чином, студенти можуть відпрацювати різні способи обробки імпульсів та позиціонування маніпуляторів, у тому числі й роботу з перериваннями та структуру програми з кількома організаційними блоками.

У процесі налагодження маніпуляторів, були виявлені «зайві» імпульси від енкодерів. Тобто у початковому положенні маніпулятора може виникнути не нульова координата. До 3–5 зайвих імпульсів спотворення не суттєве, але для позиціонування обертання роботу з вакуумним захоплюванням, наприклад, спотворення сягало до 30 імпульсів. У результаті робота з вакуумним захоплюванням не втрапляв присоскою на поверхню деталі і не міг її підхопити. У зв'язку з цим рекомендується для точного позиціонування під час обертання роботу з вакуумним захоплюванням та для точного позиціонування під час обертання поличкового укладальника багаторурсного складу, активний рівень сигналу скидання екземпляру відповідного функціонального блоку таймера, енкодера чи лічильника затримувати на час не менше 1,5 с. Це було виявлено емпіричним шляхом.

Функціональний блок визначення позиції деталі не суттєвий, реалізується типовим бібліотечним функціональним блоком лічильника або за допомогою обробника апаратного переривання в окремому організаційному блоці.

	Name	Data type	Address	Retain	Acces...	Writa...	Visibl...	Comment
	Pulse_PLC_Dia0	Bool	%I0.0					pulse counter
	LghtInlt_SM1_Dib0	Bool	%I9.0					light-barrier inlet
	LghtBhnd_SM1_Dib1	Bool	%I9.1					light-barrier behind color sensor
	LghtWht_SM1_Dib3	Bool	%I9.3					light-barrier white
	LghtRd_SM1_Dib4	Bool	%I9.4					light-barrier red
	LghtBlu_SM1_Dib5	Bool	%I9.5					light-barrier blue
	Colr_PLC_AIO	Int	%IW64					color sensor
	Conv_SM1_DQb0	Bool	%Q9.0					motor conveyor belt
	Compr_SM1_DQb1	Bool	%Q9.1					Compressor
	FrstEj_SM1_DQb2	Bool	%Q9.2					valve first ejector (white)
	SecEj_SM1_DQb3	Bool	%Q9.3					valve second ejector (red)
	ThrdEj_SM1_DQb4	Bool	%Q9.4					valve third ejector (blue)
	TempOut9	Byte	%QB9					

Рис. 3. Таблиця тегів для модуля сортувальної лінії симулятора фабрики

Більш суттєвим є визначенням кольору деталі у модулі сортувальної лінії. Останнє пов'язане із задачею визначення еталонних проміжків аналогового сигналу 0–10 В для білої, червоної та синьої деталі у камері визначення кольору. Не дивлячись на те, що камера є більш-менш затіненою, точність визначення кольору може залежати від освітленості приміщення, розміщення біля стенду джерела світла, світло-відбиття навколишніх предметів. Крім того, поверхня деталі у зоні дії датчика має округлу форму, тому її світловідбиваючі властивості сильно залежать від положення відносно датчика у момент вимірювання. У зв'язку з цим, група студентів, що вирішує задачу визначення позиції деталі повинна виконати дослідження із визначення еталонних значень сигналів аналогового оптичного датчика для заданих кольорів. Крім того, програмне рішення для сортувальної лінії повинне містити обов'язкову зупинку конвеєра, що транспортує деталь, у певному місці всередині камери визначення кольору. Еталонні координати зупинки конвеєра (у імпульсах), а також координати трьох штовхачів деталей також мають бути визначені студентом емпірично у процесі навчання.

Не менш важливою для підготовки фахівця задачею, що можна вирішувати на навчально-дослідницькій моделі симулятора фабрики, є задача побудови каналу зв'язку PROFINET між контролерами CPU1215C та CPU1211C. Момент досягнення маніпулятором заданої просторової координати можна передати й без PROFINET – через вільні дискретні входи й виходи контролерів, і поточне технічне рішення це дозволяє. Але наявність об'єднаних у мережу контролерів дозволяє обробляти дані централізовано, відображати додаткову інформацію на інтерфейсі кіберфізичної системи та для слухачів курсу набувати навички з організації, налагодження каналу зв'язку, підготовки до передачі та обробки прийнятих даних. Це окрема задача із своїми алгоритмами, як на боці CPU1215C, так і на боці CPU1211C, що вирішуватись, як методом побудови функціональних блоків, так і функцій, чи окремих організаційних блоків. Паралельно можуть бути задіяні 2 підгрупи, котрі, тим не менш, мають співпрацювати одна з одною у питанні організації прикладного інтерфейсу обміну даними.

Інтерфейс користувача має надавати засоби для задання комірки завантаження та вивантаження деталі зі багатоярусного складу. Хоча для цього використано HMI-панель KTP700 Comfort, він може бути реалізований і на персональному комп'ютері засобами WinCC. На

сьогодні тут не передбачено алгоритмів і графів керування і робота має стосуватись виключно питань організації людино-машинного інтерфейсу наявними у Tia Portal V14 засобами.

Звичайно, після вирішення означених вище задач, всі рішення обов'язково мають бути зібрані у одному програмному проекті із подальшими їх взаємоузгодженням та відлагодженням у загальній програмі керування моделлю симулятора автоматизованої фабрики.

Таким чином, отримуємо групу задач, які можуть виконуватись на навчально-дослідницькій моделі симулятора фабрики у рамках навчального курсу, показану у таблиці 1.

Виходячи із міркувань обов'язковості виконання усього комплексу задач у таблиці 1, навчальний курс чи роботи із дослідження складаються із 7 етапів, причому виконання 1 і 2 строго послідовне, а інші можуть виконуватись у будь-якому порядку за виключенням 7-го. Вона у будь-якому випадку є завершальною для всіх підгруп підготовки і з точки зору планування навчального навантаження на модель симулятора фабрики є «вузьким» місцем. Питання легко вирішується, якщо процес відлагодження всього програмного забезпечення організувати розподілено у часі. Для цього кожна окрема підгрупа повинна із самого початку використовувати один і той же програмний проект, крім того підгрупи мають обмінятися одна з одною готовими програмними рішеннями по роботам 1–2 або отримати еталонні рішення у випадку, якщо потрібний варіант завдання не виконувався. Рішення мають бути включені у робочий проект. У результаті завдання 7 буде виконане паралельно відлагодженню рішень за п. 3–6.

Таким чином, запропонована функціональна структура для програмного забезпечення моделі симулятора фабрики дозволяє організувати одночасну підготовку від 12 до 16 підгруп або окремих слухачів з використанням паралельно симуляторів процесорів серії 1200 у самому Tia Portal.

На сьогодні на основі запропонованої функціональної структури впроваджено і випробувано, в якості еталонно-демонстраційної, програма керування обладнанням навчально-дослідницької моделі симулятора фабрики від Fischertechnik. Стенд з цією програмою використовується для реклами профільних спеціальностей у Національно-технічному університеті «Дніпровська політехніка» та для створення методичного забезпечення спеціальних предметів спеціальностей 141 «Електроенергетика, електротехніка та електромеханіка», а також 123 «Комп'ютерна інженерія».

Таблиця 1

Розподіл окремих навчальних та дослідницьких завдань, які можуть виконуватись у курсі навчання із використанням моделі симулятора фабрики

№ з/п	Завдання курсу	Кількість підгруп
1	Підготовка проєкту та таблиця тегів (інформаційних ознак). Варіанти завдання: – автоматизований багатоярусний склад; – робот з вакуумним захопленням; – мультипроцесова станція з піччю; – сортувальна лінія з визначенням кольору; – фабрика із задіянням всіх зазначених модулів.	5 підгруп
2	Написання і відлагодження графу керування. Варіанти завдання: – автоматизований багатоярусний склад; – робот з вакуумним захопленням; – мультипроцесова станція з піччю; – сортувальна лінія з визначенням кольору; – фабрика із задіянням всіх зазначених модулів.	5 підгруп
3	Використання бібліотеки енкодерів. Варіанти завдання: – позиції органів багатоярусного складу; – позиції органів роботи з вакуумним захопленням Використання бібліотеки лічильників. Варіанти додаткового завдання: – позиції органів багатоярусного складу; – позиції органів роботи з вакуумним захопленням. Використання бібліотеки таймерів. Варіанти додаткового завдання: – позиції органів багатоярусного складу; – позиції органів роботи з вакуумним захопленням. Використання обробників апаратних переривань. Варіанти додаткового завдання: – позиції органів багатоярусного складу; – позиції органів роботи з вакуумним захопленням.	2 підгрупи
4	Визначення кольору деталі. Варіанти завдання: – позиціонування деталі з використанням бібліотеки лічильників; – позиціонування деталі з використанням обробників апаратних переривань.	2 підгрупи
5	Передача даних мережею PROFINET. Варіанти завдання: – передача даних від енкодерів CPU1211C; – передача команди очистки буферів даних енкодерів від CPU1215C.	2 підгрупи
6	Створення інтерфейсу користувача: – на панелі HMI; – на SCADA-системі.	2 підгрупи
7	Відлагодження всього програмного рішення для функціональної структури	Залежить від попередньої кількості підгруп

Висновки і перспективи подальших досліджень. Виходячи з аналізу технологічної схеми моделі симулятора фабрики Fischertechnik, його датчиків та виконавчих механізмів, схеми підключення до використаного промислового обчислювального обладнання, існуючих методів та рішень для складання алгоритмів керування, було розроблено, випробувано і впроваджено функціональну структуру керування моделлю симулятора фабрики від Fischertechnik, яка дозволяє структурно визначити окремі навчальні та дослідницькі задачі у рамках підготовки фахівців галузі 12 «Інформаційні технології» та галузі 14 «Електроенергетика, електротехніка та електромеханіка». Запропоноване рішення дозволяє на моделі організувати від 7 до 11 окремих лабораторних робіт, пов'язаних з різноплановими питаннями: підготовка проєкту та даних, складання та програмна реалізація графу керу-

вання, отримання даних з позиціонування органів обладнання різними способами, передача даних мережею PROFINET, створення інтерфейсу користувача, відлагодження колективно створених проєктів. Таким чином, стає можливим більш повно і ефективно використовувати лабораторну модель симулятора фабрики Fischertechnik, як для підготовки фахівців 12-ї та 14-ї галузей, так і для виконання досліджень з проєктування та створення кіберфізичних систем, включаючи методи обробки результатів вимірювань, та передачу даних між розподіленими обчислювальними вузлами.

Подальший розвиток досліджень буде спрямовано на розвиток підходів до створення інтерфейсів користувача: розробці й обґрунтуванню дизайну, протоколюванню й збереженню даних, створенню й збереженню технологічних рецептів, питанням захисту інформації.

ЛІТЕРАТУРА:

1. Factory Simulation 24V – Training models. URL: <http://surl.li/goofq> (дата звернення 24.04.2023 р.)
2. Інтерактивний каталог CA 01. Siemens. URL: <http://surl.li/goosb> (дата звернення 24.04.2023 р.)
3. Factory Simulation 24V – Datasheet. URL: <http://surl.li/googf> (дата звернення 24.04.2023 р.)
4. The PLC sample program for controlling the fischertechnik Training Factory Industry 4.0. URL: <http://surl.li/dcnvj> (дата звернення 24.04.2023 р.)
5. Fabrik Simulation/Factory Simulation/fischertechnik. URL: <http://surl.li/dcnvq> (дата звернення 24.04.2023 р.)
6. Simatic S7. S7-1200 Programmable controller. System Manual. – Siemens, 05/2021. 1418 p. URL: <http://surl.li/gooln> (дата звернення 24.04.2023 р.)
7. IEC 61131-3. Programmable controllers. Part 3: Programming languages. International Electrotechnical Commission. 226 p. URL: <http://surl.li/gooix> (дата звернення 24.04.2023 р.)

REFERENCES:

1. Factory Simulation 24V – Training models. URL: <http://surl.li/goofq> (access date 24/04/2023)
2. Factory Simulation 24V – Datasheet. URL: <http://surl.li/googf> (access date 24/04/2023)
3. Interactive catalog SA 01. Siemens. URL: <http://surl.li/goosb> (access date 24/04/2023)
4. The PLC sample program for controlling the fischertechnik Training Factory Industry 4.0. URL: <http://surl.li/dcnvj> (access date 24/04/2023)
5. Fabrik Simulation/Factory Simulation/fischertechnik. URL: <http://surl.li/dcnvq> (access date 24/04/2023)
6. Simatic S7. S7-1200 Programmable controller. System Manual. Siemens, 05/2021. 1418 p. URL: <http://surl.li/gooln> (access date 24/04/2023)
7. IEC 61131-3. Programmable controllers. Part 3: Programming languages. International Electrotechnical Commission. 226 p. URL: <http://surl.li/gooix> (access date 24/04/2023)

УДК 004.942

DOI <https://doi.org/10.32782/IT/2023-2-12>

Володимир БАБИЧ

старший викладач кафедри комп'ютерних наук, Львівський торговельно-економічний університет, вул Туган-Барановського, 10, м. Львів, Україна, 79005, babych_v@lute.lviv.ua

ORCID: 0000-0003-1996-9332

Анатолій КОСТЕНКО

кандидат фізико-математичних наук, професор кафедри комп'ютерних наук, Львівський торговельно-економічний університет, вул Туган-Барановського, 10, м. Львів, Україна, 79005, avak54@lute.lviv.ua

ORCID: 0000-0002-2162-7852

Василь ПЛЕША

старший викладач кафедри комп'ютерних наук, Львівський торговельно-економічний університет, вул Туган-Барановського, 10, м. Львів, Україна, 79005, plesha_v@i.ua

ORCID: 0000-0001-5321-9602

Михайло ПЛЕША

кандидат фізико-математичних наук, доцент кафедри комп'ютерних наук, Львівський торговельно-економічний університет, вул Туган-Барановського, 10, м. Львів, Україна, 79005, milan@lute.lviv.ua

ORCID: 0009-0002-6496-1102

Леся ХМІЛЯРЧУК

старший викладач кафедри комп'ютерних наук, Львівський торговельно-економічний університет, вул Туган-Барановського, 10, м. Львів, Україна, 79005, _lkh@ukr.net

ORCID: 0000-0002-1753-6472

Бібліографічний опис статті: Бабич, В., Костенко, А., Плеша, В., Плеша М., Хмілярчук, Л. (2023). Задача пошуку найкоротшого шляху: порівняльний аналіз основних алгоритмів. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 99–106, doi: <https://doi.org/10.32782/IT/2023-2-12>

ЗАДАЧА ПОШУКУ НАЙКОРОТШОГО ШЛЯХУ: ПОРІВНЯЛЬНИЙ АНАЛІЗ ОСНОВНИХ АЛГОРИТМІВ

За допомогою графів будуються математичні моделі зв'язків між певними елементами. Наприклад, у вигляді графа можуть бути зображені транспортні, інформаційні, комп'ютерні та інші мережі, карти автомобільних, залізничних, повітряних шляхів, лабіринти і т.п.

Питання про найкоротший шлях і досі є однією з найактуальніших тем у галузі досліджень. Знаходження найкоротших шляхів у графі використовується у різних сферах діяльності, наприклад, для знаходження оптимального маршруту між двома об'єктами на карті місцевості, у логістиці вантажних перевезень, у системах комутації інформаційних пакетів в мережі Internet тощо.

У цій статті представлено основні принципи трьох алгоритмів пошуку найкоротшого шляху та проведено їхнє порівняння шляхом аналізу часової та просторової складності. Крім того, узагальнено область застосування різних алгоритмів. Алгоритм Дейкстри – це класичний алгоритм отримання найкоротшого шляху від конкретної вершини до будь-якої іншої. Його широко використовують в дорожніх мережах. Цей алгоритм можна використовувати лише тоді, коли у графі не існує жодного ребра з від'ємною вагою. Алгоритм Беллмана-Форда можна використовувати на графах з від'ємною вагою ребер, якщо граф не містить негативного циклу, доступного з вихідної вершини. Результат роботи цього алгоритму можна використати для визначення існування циклу від'ємної ваги у графі. Алгоритм Флойда-Уоршелла – це алгоритм динамічного програмування, який може вирішити проблему найкоротшого шляху між будь-якими двома вершинами. Метод використовується на зв'язаних графах, у яких можуть бути як додатні, так і від'ємні ваги ребер, проте у ньому не має бути від'ємних циклів. Таким чином, цей метод загальніший у порівнянні з алгоритмом Дейкстри. Однак у практичному застосуванні ці три алгоритми безпосередньо не застосовуються, а проводиться їхня модифікація та оптимізація для підвищення ефективності.

Ключові слова: задача пошуку найкоротшого шляху, алгоритм Дейкстри, алгоритм Беллмана-Форда, алгоритм Флойда-Уоршелла, граф.

Volodymyr BABYCH

Senior Lecturer at the Department of Computer Sciences, Lviv University of Trade and Economics, Tugan-Baranovskoho str., 10, Lviv, Ukraine, 79005, babych_v@lute.lviv.ua

ORCID: 0000-0003-1996-9332

Anatoly KOSTENKO

Candidate of Physical and Mathematical Sciences, Professor at the Department of Computer Sciences, Lviv University of Trade and Economics, Tugan-Baranovskoho str., 10, Lviv, Ukraine, 79005, avak54@lute.lviv.ua

ORCID: 0000-0002-2162-7852

Vasyl PLESHA

Senior Lecturer at the Department of Computer Sciences, Lviv University of Trade and Economics, Tugan-Baranovskoho str., 10, Lviv, Ukraine, 79005, plesha_v@i.ua

ORCID: 0000-0001-5321-9602

Mykhailo PLESHA

Candidate of Physical and Mathematical Sciences, Associate Professor at the Department of Computer Sciences, Lviv University of Trade and Economics, Tugan-Baranovskoho str., 10, Lviv, Ukraine, 79005, milan@lute.lviv.ua

ORCID: 0009-0002-6496-1102

Lesya KHMILYARCHUK

Senior Lecturer at the Department of Computer Sciences, Lviv University of Trade and Economics, Tugan-Baranovskoho str., 10, Lviv, Ukraine, 79005, _lkh@ukr.net

ORCID: 0000-0002-1753-6472

To cite this article: Babych, V., Kostenko, A., Plesha, V., Plesha, M., Khmilyarchuk, L. (2023). Zadacha poshuku naikоротshoho shliakhu: porivnialnyi analiz osnovnykh alhorytmiv [The problem of finding the shortest path: a comparative analysis of the main algorithms]. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 99–106, doi: <https://doi.org/10.32782/IT/2023-2-12>

THE PROBLEM OF FINDING THE SHORTEST PATH: A COMPARATIVE ANALYSIS OF THE MAIN ALGORITHMS

Mathematical models of connections between certain elements are built with the help of graphs. For example, transport, information, computer and other networks, maps of road, railway, air routes, labyrinths, etc. can be depicted in the form of a graph.

The question of the shortest path is still one of the most relevant topics in the field of research. Finding the shortest paths in a graph is used in various fields of activity, for example, to find the optimal route between two objects on a terrain map, in freight logistics, in information packet switching systems on the Internet, etc.

This article presents the basic principles of three algorithms for finding the shortest path and compares them by analyzing time and space complexity. In addition, the scope of application of various algorithms is summarized. Dijkstra's algorithm is a classic algorithm for obtaining the shortest path from a particular vertex to any other. It is widely used in road networks. This algorithm can be used only when the graph does not have any edge with negative weight. The Bellman-Ford algorithm can be used on graphs with negative edge weights, as long as the graph does not contain a negative cycle accessible from the source vertex. The result of this algorithm can be used to determine the existence of a negative weight cycle in a graph. The Floyd-Warshall algorithm is a dynamic programming algorithm that can solve the shortest path problem between any two vertices. The method is used on weighted graphs, which can have both positive and negative edge weights, but must not contain negative cycles. Thus, this method is more general in comparison with Dijkstra's algorithm. However, in practical application, these three algorithms are not directly used, but their modification and optimization are carried out to improve efficiency.

Key words: problem of finding the shortest path, Dijkstra's algorithm, Bellman-Ford algorithm, Floyd-Warshall algorithm, graph.

Вступ. Проблема знаходження найкоротшого шляху є однією з класичних проблем теорії графів. Її метою є пошук найкоротшого

шляху між двома вузлами графа. Класичним алгоритмом пошуку найкоротшого шляху є знаходження шляху між двома вузлами таким

чином, щоб сума ваг складових його ребер була мінімальною. Задача про знаходження найкоротшого шляху є однією із найбільш фундаментальних і типових задач теорії графів. Інакше вона ще називається задачею про мінімальний шлях. Практичне значення задачі визначається тим, що вона є елементом складніших задач та моделлю для розробки ефективних методів вирішення задач теорії графів.

Значимість розв'язання даної задачі визначається її різними практичними застосуваннями. Наприклад, в GPS-навігаторах здійснюється пошук найкоротшого шляху між двома перехрестями. Вершинами є перехрестя, а дороги – ребрами, які лежать між ними. Якщо сума довжин доріг між перехрестями мінімальна, тоді знайдений шлях найкоротший.

Відзначимо, що на практиці за допомогою цієї задачі можна розв'язувати не лише задачі обчислення довжини шляху, але й фінансові, матеріальні, паливно-енергетичні задачі. Так, наприклад, у економічних за змістом задачах довжиною ребра можуть бути такі величини, як час, вартість, обсяг витрачених ресурсів. При проходженні шляху ці характеристики додаються, відповідно, як і самі довжини при розрахунку загальної довжини. Це дає змогу мінімізувати вищевказані показники.

Теоретично цілком допустиме дослідження графів із ребрами, у яких значення ваг є від'ємними. Однак задачі такого типу, як правило, не мають простого економічного змісту, і складно вирішуються. Задача про найкоротший шлях має рішення для будь-якого зв'язного графа. При цьому рішенням задачі вважається найкоротший шлях і його довжина. Зрозуміло, що довжина найкоротшого шляху визначається однозначно, в той час як найкоротших шляхів може бути декілька. Рішення задачі про найкоротший шлях може бути проведено двома способами: на основі алгоритму повного перебору усіх можливих шляхів або методом динамічного програмування на графі.

Під час вирішення задачі про найкоротший шлях на основі алгоритму повного перебору шляхів будуються всі можливі шляхи між початковою і кінцевою точками, обчислюються довжини усіх отриманих шляхів і вибирається мінімальна з них. Вибраний шлях (один або кілька) і буде розв'язком задачі. При цьому в алгоритм перебору шляхів можна вносити певні поправки, які дозволяють ігнорувати деякі явно невигідні варіанти. Наприклад, не потрібно продовжувати обчислення довжини шляху, якщо його довжина вже перевищує мінімальну довжину знайдених раніше шляхів. Зрозуміло,

що основним недоліком такого алгоритму є його низька швидкість та велика трудомісткість для складних графів.

Основою методу динамічного програмування є розбиття складної задачі на простіші підзадачі. Якщо задачу можна вирішити оптимально, розбиваючи її на підзадачі, а потім рекурсивно шукати оптимальні рішення підзадач, то вона має оптимальну структуру.

Аналіз досліджень і публікацій. У 1959 р. Дейкстра запропонував алгоритм пошуку в графі, який можна використати для вирішення задачі знаходження найкоротшого шляху між двома вершинами для будь-якого графа з додатними вагами (Dijkstra, 1959). Цей алгоритм пізніше був модифікований Лі в 2006 році і застосований до системи наведення автомобіля. Він має дві модифікації, а саме: найкоротший шлях та найшвидший алгоритм (Chen et al., 2009). Алгоритм визначення найкоротшого шляху обчислює найкоротший маршрут між кожною парою вершин графа, в той час як алгоритм найшвидшого шляху визначає шлях з мінімальним часом його проходження. Майбутній час проходження шляху можна передбачити на основі моделей прогнозування, використовуючи дані про час проходження маршруту (Chen; Makki; Pissinou, 2009).

Meghanathan (2012) розглянув алгоритм Дейкстри та алгоритм Беллмана-Форда для пошуку найкоротшого шляху в графі. Він дійшов висновку, що часова складність алгоритму Дейкстри дорівнює $O(|E| \cdot \log |V|)$, в той час як складність алгоритму Беллмана-Форда – $O(|V| \cdot |E|)$ (Meghanathan, 2012) (тут і далі $|V|$ та $|E|$ – позначення, відповідно, кількості вершин та кількості ребер графа (V, E)).

Lili Cao та співавт. (2005) виявили, що пошук найкоротшого шляху є важливим для численних програм на основі графів, особливо розроблених для соціальних мереж. Прикладом може бути платформа LinkedIn, де користувачі виконують запити, щоб знайти найкоротший шлях до «соціальних посилань», яка пов'язує їх із певним користувачем. Цей тип запитів є складним завданням для графів середнього розміру, і стає обчислювально нездійсненним для графів, які моделюють сучасні соціальні мережі, більшість з яких містять мільйони вузлів і мільярди ребер. Автори пропонують програму Atlas – новий підхід до масштабованих наближених найкоротших шляхів між вузлами графа, використовуючи колекцію охоплюючих дерев. Охоплюючі дерева легко генерувати, вони компактні щодо графів, і їх можна розподілити між комп'ютерами для паралелізації

запитів. У роботі показано ефективність такого підходу на прикладі шести великих соціальних графів від Facebook, Orkut та Renren, найбільший з яких включає 43 мільйони вузлів та 1 мільярд ребер. Описано методи поступового оновлення, оскільки соціальні графи з часом змінюються. Було зафіксовано динаміку графів, використовуючи 35 щоденних знімків мережі Facebook, і показано, що Atlas може з часом амортизувати вартість оновлення дерев. Нарешті, Atlas застосовано до кількох графічних додатків і показано, що результати є наближені до ідеальних (Cao L., Zhao X., Zheng H., Zhao B., 2011).

Wadhwa (2000) заявив, що дослідники вирішили проблему проектування мережі (проблема кабелю та траншеї), яка передбачає компроміс між витратами на використання та капітальними витратами на будівництво мережі. Більша мережа (дерево найкоротшого шляху) може коштувати дорожче, але може зменшити витрати на використання, включаючи привабливіші шляхи відправлення-призначення. І навпаки, менша мережа (мінімальне охоплююче дерево) може збільшити витрати на використання. Надано евристику, яка дає оптимальні або майже оптимальні рішення (Wadhwa, 2000).

Pallottino і Scutella (1997) описали алгоритм найкоротшого шляху в транспортних моделях: класичні та інноваційні аспекти. Вони розглянули алгоритми найкоротшого шляху при перевезеннях у двох аспектах. Перший – це класичні алгоритми, які є найбільш цікавими у транспортних задачах або щодо теоретичних міркувань, або щодо їхньої ефективності, а також з огляду на їхнє практичне використання в транспортних моделях. Другий аспект вказує на динамічні проблеми щодо найкоротшого шляху, які часто виникають у транспортній галузі. Автори проаналізували основні особливості проблем, які існують у відповідних умовах щодо часу транспортування та функцій витрат, загальної «хронологічної» алгоритмічної парадигми, яка називається Chrono-SPT (Pallottino; Scutella, 1998).

Ahmat (2005) докладно вивчав комунікаційні мережі. У дослідженні описані основні поняття теорії графів та їхній зв'язок із комунікаційними мережами, представлено деякі проблеми оптимізації, пов'язані з протоколами маршрутизації та моніторингом мережі, і показано, що багато проблем оптимізації є NP-повними або NP-складними. Також описано деякі загальні інструменти, які використовуються для створення топологій мережі на основі теорії графів. Xiao вирішив проблему найшвидших онлайн-

відповідей на запити, використовуючи насичену симетрію в графах. Алгоритм Дейкстри є найвідомішим та широко використовуваним для вирішення проблеми знаходження найкоротшого шляху, оскільки він швидкий і використовує структуру даних «купа» для пріоритетних черг запитів найкоротшого шляху, необхідних у багатьох додатках (Ahmat, 2009).

Steinhardt (2006) дійшов до висновку, що алгоритм Дейкстри спеціалізується на пошуку найкоротших шляхів між вершинами графа. Andrew V. Goldberg (2008) вивчав алгоритми найкоротшого шляху «вершина–вершина» (P2P). Останнім часом набули розвитку алгоритми пошуку найкоротшого шляху від точки до точки з попередньою обробкою. На практиці алгоритми виявилися ефективними на дорожніх мережах та деяких інших видах графів. Є кілька питань, особливо теоретичних, які залишаються відкритими. Отже, необхідне теоретичне обґрунтування цих алгоритмів. Обчислювальні охоплення – це ще один набір відкритих питань. Можна змінити стандартний алгоритм знаходження найкоротшого шляху для всіх пар, щоб досягти результату за однаковий час, тобто $O^*(n^2)$ для розріджених графів. Оскільки результат для задачі про всі пари дорівнює n^2 , можливості для удосконалення обмежені.

Shirinivas та співавт. (2010) представили важливість теоретичних ідей графів у різних областях комп'ютерних додатків, таких як алгоритм найкоротшого шляху в мережі, пошук мінімального охоплюючого дерева, пошук планарності графа, алгоритми пошуку матриць суміжності, алгоритми пошуку зв'язності, алгоритми знаходження циклу в графі, алгоритми пошуку елемента в структурі даних (DFS, BFS) (Shirinivas; Vetrivel; Elango, 2010).

Sommer (2010) досліджував обробку запитів із найкоротшим шляхом у мережах як теоретично, так і практично. Експериментальне дослідження проводилося з використанням дорожньо-транспортної мережі. Дослідження показало простий і загальний метод для ефективно підтримки найкоротших запитів у неорієнтованих графах з дуже низькими накладними витратами на попередню обробку та конкурентним часом запитів за рахунок точності. Ефективність цього методу, який є розумною альтернативою існуючим точним методам, спеціально розробленим для транспортних мереж, доведена для різних типів графів (Sommer, 2010).

Abbasi та ін. (2011) розглядали динамічну проблему знаходження найкоротшого шляху, яку застосовують у динамічних потоках мінімальних витрат для задачі перетворення.

Дослідження показало, що ця проблема еквівалентна класичній задачі пошуку найкоротшого шляху у так званій розширеній у часі мережі (Abbasi; Ebrahimnejad, 2011).

Hung (2003) аналізує обернену проблему знаходження найкоротшої довжини шляху (ISPL) у задачах вдосконалення транспортної мережі та цінового потоку (Hung Cheng-Huang, 2003).

Li та співавт. (2008) запропонували ефективний алгоритм під назвою Li-Qi (LQ) для пошуку простого шляху з найменшою сумарною вагою від конкретної початкової або вихідної вершини до кожної іншої вершини графа (Li T.; Qi L.; Ruan, 2008).

Виклад основного матеріалу. До найефективніших алгоритмів знаходження найкоротшого шляху на графах належать алгоритм Дейкстри (використовується для знаходження оптимального маршруту між двома вершинами), алгоритм Флойда-Уоршелла (для знаходження найкоротшого шляху між парами вершин) і алгоритм Беллмана-Форда (для знаходження оптимального маршруту між усіма парами вершин). Зазначені алгоритми легко виконуються при малій кількості вершин у графі. При збільшенні їхньої кількості задача пошуку найкоротшого шляху ускладнюється, і тому без застосування сучасної комп'ютерної техніки практично не обійтись.

Алгоритм Дейкстри. Алгоритм голландського вченого Едсгера Дейкстри знаходить усі найкоротші шляхи з однієї наперед заданої вершини графа до всіх інших (Dijkstra's algorithm). За його допомогою, при наявності всієї необхідної інформації, можна, наприклад, дізнатися, яку послідовність доріг краще використовувати, щоб дістатися з одного міста до кожного з інших, або в які країни вигідніше експортувати нафту тощо.

Недолік цього методу – неможливість обробки графів, у яких є ребра з від'ємними вагами. Наприклад, якщо для деякої системи передбачено збиткові маршрути, то для роботи з нею алгоритм Дейкстри не підходить.

Щоб цей алгоритм реалізувати, потрібно два масиви: логічний *visited* – для зберігання інформації про відвідані вершини, і чисельний *distance*, в який будуть заноситися знайдені довжини найкоротших шляхів. Отже, є граф $G = (V, E)$. Кожна з вершин, які входять у множину V , спочатку відзначена як невідвіdana, тобто елементам масиву *visited* присвоєно значення *false*.

Оскільки найвигідніші шляхи тільки належить знайти, на початку значенням кожного

елемента вектора *distance* буде таке число, яке завідома є більшим від будь-якого потенційного шляху (зазвичай це є нескінченність, але в програмі використовують, наприклад, максимальне значення конкретного типу даних). Як вихідний пункт вибирається вершина *s*, для неї записується нульовий шлях: $\text{distance}[s] = 0$, оскільки немає ребра від *s* до *s* (метод не передбачає петель).

Далі знаходяться всі сусідні вершини (до яких існує ребро від *s*) (наприклад, *t* і *u*) та досліджуються по черзі, а саме обчислюється вартість маршруту від *s* до кожної з них:

$\text{distance}[t] = \text{distance}[s] + \text{вага інцидентного до } s \text{ і } t \text{ ребра};$

$\text{distance}[u] = \text{distance}[s] + \text{вага інцидентного до } s \text{ і } u \text{ ребра}.$

Цілоком ймовірно, що до тієї чи іншої вершини від *s* існує кілька шляхів, тому вартість шляху до такої вершини в масиві *distance* доведеться переглядати, тоді найбільше (неоптимальне) значення ігнорується, а найменше ставиться у відповідність до вершини.

Після обробки суміжних із *s* вершин вона позначається як відвіdana: $\text{visited}[s] = \text{true}$, і активною стає та вершина, шлях до якої від вершини *s* є мінімальним. Припустимо, шлях від *s* до *u* коротший, ніж від *s* до *t*, отже, вершина *u* стає активною. Далі таким самим чином досліджуються її сусіди, за винятком вершини *s*. Наступним кроком *u* позначається як пройдена: $\text{visited}[u] = \text{true}$, активною стає вершина *t*, і вся процедура повторюється для неї. Алгоритм Дейкстри триває доти, доки всі доступні від *s* вершини не будуть досліджені.

Алгоритм Флойда-Уоршелла. Цей метод названо на честь двох американських дослідників Роберта Флойда і Стівена Уоршелла, які одночасно відкрили його у 1962 році.

Алгоритм Флойда-Уоршелла – динамічний алгоритм обчислення значень найкоротших шляхів для кожної з вершин графа. Метод працює на зважених графах, з додатними та від'ємними вагами ребер, але без від'ємних циклів. Таким чином, цей метод загальніший у порівнянні з алгоритмом Дейкстри, оскільки останній не працює з від'ємною вагою ребер, до того ж класична його реалізація має на меті визначення оптимальних відстаней від однієї вершини до всіх інших.

Для реалізації алгоритму Флойда-Уоршелла сформуємо матрицю суміжності $D[][]$ графа $G = (V, E)$, в якому кожна вершина понумерована від 1 до $|V|$. Ця матриця має розмір $|V| \cdot |V|$, і кожному її елементові $D[i][j]$ присвоєно вагу ребра від вершини *i* до вершини *j*. Під час вико-

нання алгоритму ця матриця буде перезаписуватися: в кожну з її комірок буде заноситися значення, яке визначає оптимальну довжину шляху від вершини i до вершини j (відмова від виділення спеціального масиву для цього заощадить пам'ять і час).

Тепер розглянемо вміст матриці найкоротших шляхів. Оскільки кожен її елемент $D[i][j]$ повинен вміщувати найменший із наявних маршрутів, то відразу можна сказати, що для одиначної вершини він дорівнює нулю, навіть якщо вона має петлю (від'ємні цикли не розглядаються), отже, всім елементам головної діагоналі ($D[i][i]$) потрібно надати нульове значення.

А щоб нульові недіагональні елементи (матриця суміжності могла мати нульові значення там, де немає безпосереднього ребра між вершинами i та j) змінили по можливості своє значення, позначимо їх рівними нескінченності (в програмі це може бути, наприклад, максимально можлива довжина шляху в графі, або просто – велике число).

Основна частина алгоритму складається із трьох циклів:

```
for k від 1 до |V| do
  for i від 1 до |V| do
    for j від 1 до |V| do
      if  $D[i][k] + D[k][j] < D[i][j]$  then  $D[i][j] \leftarrow D[i][k] + D[k][j]$ 
```

Найкоротший шлях від вершини i до вершини j може проходити як через них самих, так і через множину інших вершин $k \in (1, \dots, |V|)$. Оптимальним від i до j буде шлях, який або проходить через k , або не проходить. Якщо шлях проходить через вершину k , потрібно замінити значення найкоротшого шляху $D[i][j]$ сумою $D[i][k] + D[k][j]$. Іншими словами, як результат виконання алгоритму матриця суміжності буде замінена матрицею найкоротших шляхів.

Алгоритм Флойда–Уоршелла ефективний при обчисленні усіх найкоротших шляхів у щільних графах, коли є багато пар ребер між парами вершин. Його складність $O(|V|^3)$. У випадку розріджених графів ефективнішим є алгоритм Дейкстри.

Алгоритм Беллмана–Форда. Історія алгоритму пов'язана з трьома незалежними математиками Лестером Фордом, Річардом Беллманом та Едвардом Муром, іноді його називають алгоритмом Беллмана–Форда–Мура. Цей метод використовується в деяких протоколах дистанційно-векторної маршрутизації, наприклад у RIP (Routing Information Protocol).

Алгоритм обчислює найкоротші шляхи від однієї вершини до всіх інших у зваженому графі (як і алгоритм Дейкстри). Він може працювати

на графах, які мають ребра з від'ємними вагами. Проте застосувати його можна не до всіх таких графів, оскільки кожен черговий прохід шляхом, який складається з ребер з від'ємною сумою ваг, лише покращує результат. Нескінченна кількість покращень робить неможливим визначення одного конектного оптимального значення. Тому алгоритм Беллмана–Форда не можна застосувати до графів з від'ємними циклами, проте він дозволяє визначити такі графи.

Щоб знайти всі найкоротші шляхи від вершини s до всіх інших за допомогою алгоритму Беллмана–Форда, потрібно скористатися методом динамічного програмування: розбити задачу на підзадачі та знайти їхній розв'язок. Розв'язком кожної такої підзадачі є визначення найкоротшого шляху від однієї окремо взятої вершини до якоїсь іншої.

Для зберігання результатів роботи алгоритму створимо одновимірний масив $d[]$. У i -му елементі цього масиву буде зберігатися значення найкоротшого шляху від вершини s до вершини i (якщо таке існує). Спочатку присвоїмо елементам масиву $d[]$ значення, які дорівнюють умовній нескінченності (наприклад, число, яке більше за суму всіх ваг), а елементові $d[s]$ присвоїмо значення нуль, оскільки відомо, що шлях від вершини s до неї самої рівний 0. Необхідно припустити, що інші вершини з вершини s недоступні. Під час виконання алгоритму для деяких з них це припущення виявиться хибним і будуть обчислені оптимальні значення довжини шляху до цих вершин з вершини s .

Отже, задано граф $G=(V, E)$, $n=|V|$, а $m=|E|$. Позначимо суміжні вершини цього графа символами v і u , а вагу ребра (v, u) символом w . Іншими словами, вага ребра з вершини v до вершини u буде дорівнювати w .

Основна частина алгоритму Беллмана–Форда:

```
for i від 1 до n-1 do
  for j от 1 до m do
    if  $d[v] + w(v, u) < d[u]$  then  $d[u] \leftarrow d[v] + w(v, u)$ 
```

На кожному n -му кроці здійснюється спроба покращити значення елементів масиву $d[]$: якщо сума ваги ребра $w(v, u)$ і ваги, яка зберігається в $d[v]$, менша ваги $d[u]$, то вона присвоюється останньому.

Складність алгоритму при представленні графа списком ребер становить $O(|V| \cdot |E|)$. Якщо ж граф задати матрицею суміжності, алгоритм буде виконуватися за час $O(|E|^3)$.

Висновки. У підсумку, проаналізувавши основні алгоритми знаходження найкоротших шляхів у графі, зазначимо. Якщо необхідно знайти відстань від однієї вершини до іншої або до всіх вершин графа і ваги всіх ребер графа

є додатними або дорівнюють нулю, то най-ефективнішим є алгоритм Дейкстри із часом роботи $O(|E| \cdot \lg|V|)$. Якщо ж ваги ребер можуть бути від'ємними, то необхідно застосовувати алгоритм Беллмана–Форда, час роботи якого $O(|V| \cdot |E|)$. Якщо необхідно знайти відстані між усіма парами вершин графа, граф є розрідженим і всі ребра мають невід'ємні ваги, то можна виконати $|V|$ разів алгоритм Дейкстри. Якщо

необхідно знайти відстані між усіма парами вершин, ваги ребер можуть бути від'ємними і граф не є розрідженим ($|E|$ прямує до $|V|^2$), то доречно використовувати алгоритм Флойда–Уоршелла. Жоден з наведених алгоритмів не може бути застосований для графів, які містять негативні цикли. Проте алгоритм Беллмана–Форда, як і алгоритм Флойда–Уоршелла, допоможуть такі цикли виявити.

ЛІТЕРАТУРА:

1. Dijkstra's algorithm. Wikipedia: The Free Encyclopedia. URL: https://en.wikipedia.org/wiki/Dijkstra%27s_algorithm (application date: 20.06.2023)
2. Abbasi S., Ebrahimnejad S. Finding the Shortest Path in Dynamic Network using Labeling Algorithm. *International Journal of Business and Social Science*. 2011. Vol. 2. № 20. P. 239–243.
3. Ahmat K. A. (n.d.). Graph Theory and Optimization Problems for Very Large Networks. City University of New York. New York, Information Technology. 2009. 5 p.
4. Chen K., Makki K., Pissinou N. A real-time wireless route guidance system for urban traffic management and its performance evaluation. 2009 IEEE 70th Vehicular Technology Conference Fall. 2009. 5 p.
5. Dijkstra E.W. A note on two problems in connexion with graphs. *Numerische Mathematik*. 1959. Vol. 1. P. 269–271.
6. Hung Cheng-Huang. Inverse Shortest Path Length Problem. School of Industrial and systems engineering Georgia Institute of Technology. 2003. 24 p.
7. Li T., Qi L., Ruan D. An efficient Algorithm for the single source Shortest Path Problem in Graph Theory. International Conference on Intelligent System and Knowledge Engineering. 2008. P. 152–157.
8. Cao L., Zhao X., Zheng H., Zhao B. ZhaoApproximating Shortest Path in Social Graph. UC Santa Barbara: Computer Science Department. 2011. 20 p.
9. Meghanathan D. N. (n.d.). Review of Graph Theory. MS: Department of Computer Science, Jackson State University. 2012.
10. Pallottino S., Scutellà M.G. Shortest Path Algorithms in Transportation models: classical and innovative aspects. *Equilibrium and Advanced Transportation Modelling* / Patrice Marcotte, Sang Nguyen. Springer New York, NY, 1998. P.245–281.
11. Shirinivas S.G., Vetrivel S., Elango N.M. Application of Graph theory in Computer Science an overview. International Journal of Engineering Science and Technology. 2010. Vol. 2 (9). P. 4610–4621.
12. Sommer C. Approximate Shortest Path and Distance Queries in Networks. *Graduate School of Information Science and Technology*: Department of Computer Science, 2010. Tokyo, 2010.
13. Wadhwa Sh. Analysis of a network design problem. Lehigh University. 2000.

REFERENCES:

1. Dijkstra's algorithm. Wikipedia: The Free Encyclopedia. URL: https://en.wikipedia.org/wiki/Dijkstra%27s_algorithm (application date: 20.06.2023)
2. Abbasi, S., Ebrahimnejad, S. (2011). Finding the Shortest Path in Dynamic Network using Labeling Algorithm. *International Journal of Business and Social Science*, 1, 20, 239–243.
3. Ahmat, K. A. (n.d.). (2009). Graph Theory and Optimization Problems for Very Large Networks. *City University of New York. New York, Information Technology*.
4. Chen, K., Makki, K., Pissinou, N. (2009). A real-time wireless route guidance system for urban traffic management and its performance evaluation. *2009 IEEE 70th Vehicular Technology Conference Fall*.
5. Dijkstra, E.W. (1989). A note on two problems in connexion with graphs. *Numerische Mathematik*, 1, 269–271.
6. Hung, Cheng-Huang. (2003). Inverse Shortest Path Length Problem. *School of Industrial and systems engineering Georgia Institute of Technology*.
7. Li, T., Qi, L., Ruan, D. (2008). An efficient Algorithm for the single source Shortest Path Problem in Graph Theory. *International Conference on Intelligent System and Knowledge Engineering*, 152–157.
8. Cao, L., Zhao, X., Zheng, H., Zhao, B. (2011). ZhaoApproximating Shortest Path in Social Graph. *UC Santa Barbara: Computer Science Department*.

9. Meghanathan, D. N. (n.d.). (2012). Review of Graph Theory. *MS: Department of Computer Science, Jackson State University*.
10. Pallottino, S., Scutellà, M.G. (1998). Shortest Path Algorithms in Transportation models: classical and innovative aspects. *Equilibrium and Advanced Transportation Modelling* (pp. 245–281). NY: Springer New York.
11. Shirinivas, S.G., Vetrivel, S., Elango, N.M. (2010). Application of Graph theory in Computer Science an overview. *International Journal of Engineering Science and Technology*, 2(9), 4610–4621.
12. Sommer, C. (2010). Approximate Shortest Path and Distance Queries in Networks. *Graduate School of Information Science and Technology*. Tokyo, 2010.
13. Wadhwa, Sh. (2000). Analysis of a network design problem. Lehigh University.

УДК 004.41, 004.05

DOI <https://doi.org/10.32782/IT/2023-2-13>

Оксана КИРИЧЕНКО

асистент кафедри математичних проблем управління і кібернетики, Чернівецький національний університет імені Юрія Федьковича, вул. Коцюбинського 2, м. Чернівці, Україна, 58002

ORCID: 0000-0003-0282-9958

Бібліографічний опис статті: Кириченко, О. (2023). Особливості архітектури програмного забезпечення для збору та аналізу статистичної інформації в глобальній мережі. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 107–112, doi: <https://doi.org/10.32782/IT/2023-2-13>

ОСОБЛИВОСТІ АРХІТЕКТУРИ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ ДЛЯ ЗБОРУ ТА АНАЛІЗУ СТАТИСТИЧНОЇ ІНФОРМАЦІЇ В ГЛОБАЛЬНІЙ МЕРЕЖІ

Для вивчення структури та характеристики web-простору необхідно використовувати спеціалізоване програмне забезпечення. В даній роботі розглянуто особливості архітектури подібного програмного забезпечення та приділена особлива увага важливим аспектам його роботи, таким як: механізми навігації, стратегії сканування та обробка отриманих даних. У статті також розглянуто проблеми ефективності та масштабованості подібних рішень при обробці великої кількості веб-ресурсів. Розроблено програмне забезпечення, яке складається з кроулера та аналітичного модуля. Функціональним призначенням аналітичного модуля є проведення кластерного та статистичного аналізу великих об'ємів даних за допомогою різноманітних статистичних методів. Архітектура розробленого додатку відповідає останнім тенденціям у сфері програмного забезпечення, враховуючи сучасні вимоги та стандарти. За допомогою даного програмного продукту з використанням розробленого аналітичного модуля досліджені статистичні та кластерні характеристики різних сегментів веб-простору (українського – edu.ua, net.ua; польського сегменту – edu.pl та ізраїльського – ac.il).

Ключові слова: програмне забезпечення, тестування програмного забезпечення, рівні тестування програмного забезпечення, вимоги до програмного забезпечення: функціональні характеристики (вимоги), нефункціональні характеристики (вимоги), специфікація, статистичні методи.

Oksana KYRYCHENKO

Assistant Professor, Department of Mathematical Problems of Control and Cybernetics, Yuriy Fedkovych Chernivtsi National University, Chernivtsi, Ukraine, 58002, o.kyrychenko@chnu.edu.ua

ORCID: 0000-0003-0282-9958

To cite this article: Kyrychenko O. (2023). Osoblyvosti arkhitektury prohramnoho zabezpechennia dlia zboru ta analizu statystychnoi informatsiii v hlobalnii merezhi. [Features of software architecture for collecting and analyzing statistical information on the global network]. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 2, 107–112, doi: <https://doi.org/10.32782/IT/2023-2-13>

FEATURES OF SOFTWARE ARCHITECTURE FOR COLLECTING AND ANALYZING STATISTICAL INFORMATION ON THE GLOBAL NETWORK

To study the structure and characteristics of World Wide Web, specialized software is essential. This paper examines the peculiarities of the such software architecture and pays high attention to important aspects of its performing, such as navigation mechanisms, scanning strategies, and the processing of acquired data. The article also addresses the challenges of efficiency and scalability in similar solutions when handling a large number of web resources. The developed software consists of a crawler and an analytical module. The functional purpose of the analytical module is to conduct cluster and statistical analysis of large volumes of data by means of statistical methods. The architecture of the developed application aligns with the latest trends in software development, considering modern requirements and standards. Using this software with integrated analytical module, statistical and cluster characteristics of various segments of WWW have been researched, including the Ukrainian segments (edu.ua, net.ua), the Polish segment (edu.pl), and the Israeli segment (ac.il).

Key words: software, software testing, levels of software testing, software requirements: functional requirements, non-functional requirements, requirements, statistical methods.

Основною задачею програмного забезпечення для збору, аналізу та використання інформації з глобальної мережі є навігація по всесвітній павутині та завантаження контенту веб-сторінок. Подібне програмне забезпечення повинно взаємодіяти з веб-ресурсами для знаходження та збору інформації, використовуючи ефективні алгоритми навігації та сканування, видобувати та аналізувати дані з різних джерел, включаючи структуровані та неструктуровані дані.

Важливим компонентом будь-яких аналітичних програм, які обробляють велику кількість веб-сторінок є веб-кроулер.

Веб-кроулер – це програма, яка за допомогою однієї або кількох початкових URL-адрес завантажує веб-сторінки, пов'язані з цими URL-адресами, витягує будь-які гіперпосилання, що містяться в них, і рекурсивно продовжує завантажувати веб-сторінки за цими гіперпосиланнями (Najork, Heydon, 2002; Najork, 2017).

Загалом розробка високопродуктивних веб-кроулерів є нетривіальною задачею та включає ряд складнощів і викликів. Зокрема, веб-кроулерам доводиться стикатися з обмеженнями, накладеними веб-серверами. Деякі сервери можуть обмежувати швидкість доступу до сторінок для одного IP, що ускладнює підтримку високої швидкості сканування. Також, сканування великих обсягів даних може призводити до проблем із зберіганням, обробкою та передачею цих даних мережею. Кроулери повинні ефективно масштабуватися та використовувати паралельну обробку даних. Робота із великою кількістю задач одночасно потребує правильної архітектури та керування ресурсами.

Крім того, однією з найскладніших задач є обробка динамічного контенту сайтів, згенерованого з використанням технологій, таких як AJAX або JavaScript.

Для розробки якісного програмного забезпечення потрібно визначити функціональні та нефункціональні вимоги, тобто визначити характеристики, якими повинен володіти програмний продукт, щоб відповідати потребам поставленої задачі.

Функціональні вимоги представляють основні функціональні можливості або особливості, якими повинна володіти система програмного забезпечення, щоб відповідати призначеній меті. Створюючи функціональні вимоги, важливо мати на увазі, що вони повинні бути конкретними, вимірними, досяжними, релевантними та обмеженими у часі.

Нефункціональні вимоги доповнюють функціональні вимоги, вказуючи, як програмна система

повинна виконувати певні функції. Вони задають атрибути якості системи, а не її специфічні особливості, встановлюють стандарти продуктивності, безпеки та зручності використання системи, пояснюють обмеження системи, що проектується. До характеристик, якими визначаються нефункціональні вимоги відносять, продуктивність, масштабованість, портативність, сумісність, надійність, доступність, безпека, локалізація, зручність використання тощо.

Саме тому, визначення вимог є ключовим етапом у розробці подібного програмного забезпечення (Говорущенко, Боднар, Кушнір, 2019; Novorushchenko et al., 2019).

Опишемо функціональні та нефункціональні вимоги до системи, яка розроблена (Paech et al., 2002; Alashqar et al., 2015).

Функціональні вимоги:

- задавати вхідні URL-адреси та глибину сканування;
- генерувати список посилань для завантаження;
- фільтрувати на основі URL-адрес, у якому шаблони визначають імена хостів та/або шляхи, що цікавлять;
- паралельно завантажувати веб-сторінки за гіперпосиланнями;
- парсити завантажені веб-сторінки та видобувати гіперпосилання;
- перетворювати отримані посилання відповідно до правил нормалізації посилань;
- забезпечувати збереження відповідної структури посилань;
- здійснювати проведення статистичного аналізу з використанням статистичних методів;
- здійснювати проведення кластерного аналізу різних зон веб-простору.

Нефункціональні вимоги:

- планування (визначає порядок завантаження URL-адрес);
- політика відвідування сторінок (обмежує навантаження на веб-сервер при скануванні веб-простору);
- обробка веб-сторінок (видобуток гіперпосилань);
- визначення дублікатів (розпізнавати під час сканування дублікатів);
- інтерактивність (підтримувати файли cookie та прості схеми паролів);
- стабільність (бути стійким до можливих перерв у роботі, забезпечити ведення логування для моніторингу своєї поведінки та діагностики будь-яких проблем);
- протоколи (додавати нові протоколи завантаження до кроулеру);
- висока пропускну здатність;

- портативність (здатний працювати на Linux, виконувати невеликі та тестові запуски на Windows).

Розроблений кроулер складається з декількох модулів: Downloader, Content parser, Link extractor, Url resolver, Data access module, Url manager, Logging, Analytical module, Url queue, Download queue, Data store (Кириченко, Kanovsky, Остапов, 2013, с. 102). Модифіковану архітектуру розробленого кроулера зображено на рис. 1.

Гнучкість конфігурування та стабільна робота під навантаженням забезпечується модульною архітектурою кроулера. Наведемо короткий опис призначення модулів розробленого програмного забезпечення (Кириченко, Kanovsky, Остапов, 2013, с. 102).

Download queue – забезпечує збереження гіперпосилань для подальшого завантаження.

Downloader – завантажує контент сторінки за вказаним URL, обробляє заголовки запиту за сторінкою, дозволяє здійснювати завантаження контенту у багатому поточному режимі.

Content parser – розбирає контент веб-сторінки, виділяючи структуру html-тегів та інформацію, яка знаходиться всередині тегів.

Link extractor – знаходить та добуває з контенту завантаженої сторінки посилання, враховуючи зовнішні та внутрішні а також посилання Flash-об'єктів.

Url resolver – перетворює отримані посилання відповідно до правил нормалізації посилань. Такі перетворення запобігають завантаженню однакових ресурсів більш, ніж один раз.

Data access module – забезпечує обробку та збереження відповідної структури посилань в базі даних.

Url manager – генерує список посилань, для завантаження. Цей модуль реалізує політику планування щодо порядку завантаження URL-адрес, відповідає за усунення дублікатів та визначає пріоритети при завантаженні сторінок.

Data store – використовується для збереження веб-графу, тобто сторінок, їх структури посилань, різноманітної статистики.

Analytical module – надає інтерфейс для проведення статистичного та кластерного аналізу веб-графу.

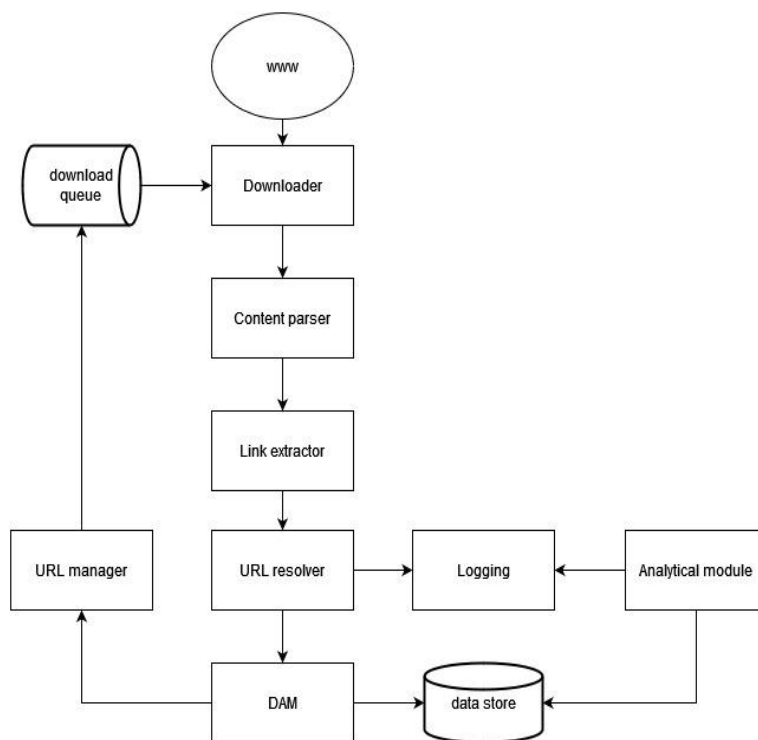


Рис. 1. Архітектура розробленого кроулера

Logging module – агрегує всі лог журнали, створені різними компонентами, щоб уніфікувати всі події в системі.

Головною особливістю розробленого нами програмного забезпечення є наявність аналітичного модуля, детальний опис якого наведено нижче.

Аналітичний модуль надає інтерфейс для проведення статистичного та кластерного аналізу веб-графу. А саме:

- побудови розподілу ймовірності вузлів за ступенями по вхідних зв'язках (in degree) та розподілу ймовірності вузлів за ступенями по вихідних зв'язках (out degree);
- обчислення коефіцієнтів кластерності веб-графу;
- визначення середніх значень ступеня вузла для неорієнтованих графів;
- визначення оптимальної кількості кластерів методом k-Core decomposition;
- знаходження центрів кластерів;
- розбиття досліджуваних мереж на кластери за допомогою алгоритму PIC (Power iteration clustering).

Архітектуру аналітичного модуля зображено на рис. 2.

Наведемо короткий опис компонентів аналітичного модуля.

Service interface – надає API для користувача, яке дозволяє ініціалізувати процес проведення статистичного та кластерного аналізу,

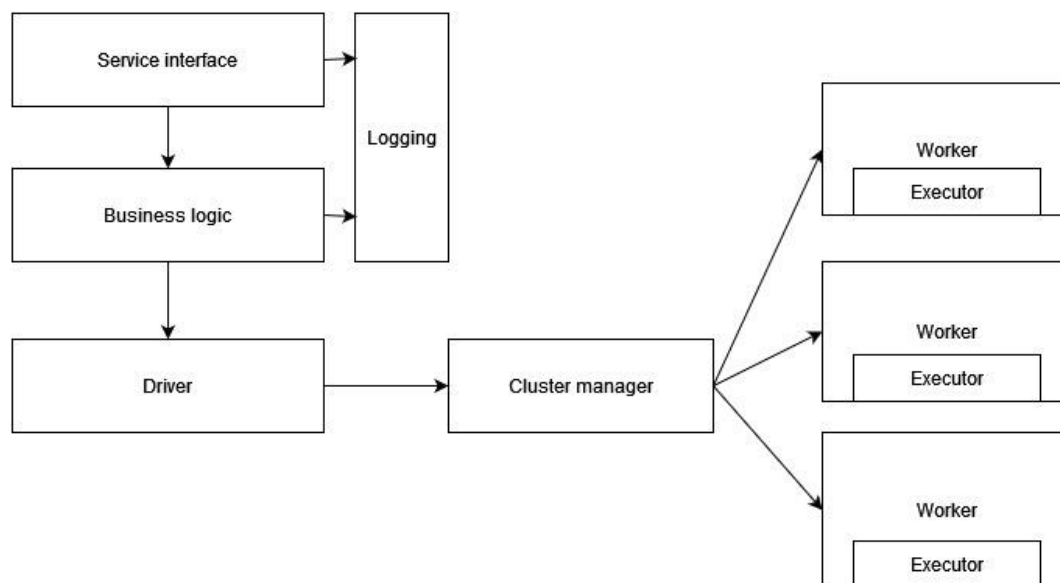


Рис. 2. Архітектура аналітичного модуля

отримує вхідні параметри та перенаправляє їх на рівень бізнес-логіки.

Business logic – відповідно до вхідних параметрів витягує необхідні дані зі сховища даних, здійснює їх перетворення та передає на вхід драйверу для здійснення обчислень.

Driver – координує воркери і контролює виконання завдань. Він складається з коду для виконання і відповідного контексту. Контекст отримує загальну задачу та ділить її на невеликі завдання, які обробляються виконавцями.

Workers – отримують від драйвера завдання та виконують їх. Виконавці утворюють кластер.

Cluster manager – відповідає за взаємодію з драйвером та виконавцями, здійснюючи такі завдання:

- керування виділенням ресурсів;
- керування поділом програми;
- керування виконанням програми.

Logging – це не фактичний логер, а наскрізний компонент (тобто – він доступний для всіх логерів). Цей компонент містить бібліотеку логування, яка використовується компонентом, і генерує події, які обробляються модулем логування.

Для проведення статистичного та кластерного аналізу веб-графу модуль надає REST API, що дозволяє користувачу вказати початкові URL-адреси.

Таким чином, використання багаторівневої архітектури дозволяє забезпечувати незалежну розробку і розвиток окремих частин додатку, ще більше збалансувати навантаження на різні вузли і мережу. Додаток складається з логічних та функціональних компонент, які можна використовувати повторно. Такий підхід значно

спрощує розробку та підтримку розподілених систем (Jaiswal, 2019, p. 2454).

З іншого боку, подібна архітектура дозволяє ізолювати додаток, повністю інкапсулюючи середовище виконання, і запускати його на віртуальній машині замість фізичного сервера. З метою оптимізації використання ресурсів віртуального або фізичного сервера була запроваджена технологія контейнеризації.

Контейнери – це абстракція на рівні програми, яка пакує код і залежності разом. Кілька контейнерів можуть працювати на одній машині та спільно використовувати ядро операційної системи з іншими контейнерами, кожен із яких працює як ізольований процес у просторі користувача. Контейнери займають менше місця, ніж віртуальні машини, можуть працювати з більшою кількістю програм і потребують менше віртуальних машин і операційних систем (Cook, 2017, p. 31). Найпопулярнішим інструментом для контейнеризації є Docker. Docker – це відкрита платформа для розробки та експлуатації додатків.

Для розробки веб-кроулера, який би працював під операційною системою Windows та операційними системами сімейства Unix, було обрано мову програмування Java, яка є потужним інструментом для розробки розподілених та багатопоточних систем. Для створення додатків на Java використали Java Platform, Enterprise Edition (JEE), яка підтримує:

- багаторівневу архітектуру;
- незалежність від платформи;
- роботу в багатопоточному режимі;
- масштабованість;
- гнучкість;

- високу здатність до розширення;
- контейнеризацію;
- швидке розгортання.

Було проведено тестування розробленого програмного забезпечення. Основною метою тестування програмного забезпечення є оцінка відповідності програмного забезпечення зазначеним потребам.

Тестування програмного забезпечення зазвичай здійснюється з метою виявлення помилок для гарантування якості програмного забезпечення та виконання остаточної перевірки відповідності специфікаціям, дизайну. Тестування програмного забезпечення є необхідним етапом, який зазвичай виконується відповідно до установлених вимог (специфікацій) (Novorushchenko et al., 2019). Тестування поділяється на різні рівні, такі як Unit (тестування окремих модулів або компонентів програми), Integration (тестування взаємодії між модулями або компонентами програми), System (тестування всієї системи як єдиного цілого) та Acceptance testing (тестування для підтвердження готовності системи до використання та відповідності вимогам) (Umar, 2019). Для здійснення ефективного тестування необхідно ретельно обирати засоби тестування.

Для оцінки якості розробленого програмного забезпечення використовувалися конкретні тестові сценарії та метрики, а саме коректність обробки структури веб-сайту, швидкості сканування, масштабованості. Функціонал компонентів програмного продукту покривався Unit тестами та інтеграційними тестами. Особливу увагу приділялася приймальному тестуванню, яке засвідчило, що розроблене програмне

забезпечення повністю відповідає зазначеним вище (функціональним та нефункціональним) вимогам.

За допомогою розробленого програмного забезпечення методом зондування досліджувались зони Web-простору: українська (edu.ua та net.ua), ізраїльська (ac.il) та польська (edu.pl). Для кожного сегменту було проскановано більше 400 тисяч веб-сторінок. За допомогою аналітичного модуля проведена статистична обробка зібраних даних, розраховано: ступінь вузла (вхідний та вихідний), розподіл ступенів вузлів по вхідних зв'язках (indegree) та вихідних зв'язках (outdegree), коефіцієнт кластеризації та ін. Дослідження кластерної структури зон веб-простору проводилось за допомогою алгоритму спектральної кластеризації PIC (Power iteration clustering), обчислено оптимальну кількість кластерів методами «ліктя» та k-Core decomposition.

Висновки. За результатами проведених досліджень можна зробити наступні висновки:

1. Розроблено програмне забезпечення (кроулер та аналітичний модуль) для збору інформації з мережі WWW і проведення статистичного та кластерного аналізу.

2. Обґрунтовано вибір архітектурного стилю, елементів та обмежень при розробці програмного забезпечення.

3. Розроблений додаток підтримує багаторівневу архітектуру, забезпечує незалежність від платформи, роботу в багатопоточному режимі, масштабованість, гнучкість, високу здатність до розширення, контейнеризацію, швидке розгортання. Такий підхід відповідає сучасним тенденціям у сфері розробки програмного забезпечення.

ЛІТЕРАТУРА:

1. Najork M., Heydon A. High-Performance Web Crawling. In: (eds) *Handbook of Massive Data Sets. Massive Computing* / J. Abello, P.M. Pardalos, M.G.C. Resende, Boston : Springer, 2002. Vol. 4. P. 25-45.
2. Najork M. Web Crawler Architecture. *Encyclopedia of Database Systems*, 2017. P. 1–4.
3. Говорущенко Т.О., Боднар М.А., Кушнір В.О. Сучасні проблеми формування та аналізу вимог до програмного забезпечення. *Вимірювальна та обчислювальна техніка в технологічних процесах*. 2019. № 1. С. 45–53.
4. T. Novorushchenko, O. Pavlova, M. Bodnar. Development of an Intelligent Agent for Analysis of Nonfunctional Characteristics in Specifications of Software Requirements. *Eastern-European Journal of Enterprise Technologies*. 2019. Vol. 1. No. 2 (97). P. 6–17.
5. Paech, Barbara & Dutoit, Allen & Kerkow, Daniel & Knethen, Antje. (2002). Functional requirements, non-functional requirements, and architecture should not be separated. A position paper. 2002.
6. Alashqar A., Elfetouh A., El-Bakry H. Requirement Engineering for Non-Functional Requirements. *International Journal of Information and Communication Technology Research*. 2015. № 5. P. 21–27.
7. Кириченко О.Л., Каповский И., Остапов С.Е. Програмне забезпечення для дослідження статистичних характеристик глобальної мережі WWW. *Системи обробки інформації*. 2013. Вип. 3 (110). Т. 2. С. 99–104.
8. Jaiswal M. Software Architecture and Software Design. *International Research Journal of Engineering and Technology*. 2019. Vol. 6. P. 2452–2454.

9. Cook J. Docker for Data Science: Building Scalable and Extensible Data Infrastructure Around the Jupyter Notebook Server. Berkeley : Apress, 2017. 257 p.

10. Umar M. Comprehensive study of software testing: Categories, levels, techniques, and types. *International Journal of Advance Research, Ideas and Innovations in Technology*. 2019. № 5. P. 32–40.

REFERENCES:

1. Najork, M., Heydon, A. (2002). High-Performance Web Crawling. In: Abello, J., Pardalos, P.M., Resende, M.G.C. (eds) *Handbook of Massive Data Sets. Massive Computing*, vol. 4. Springer, Boston, MA. DOI: 10.1007/978-1-4615-0005-6_2

2. Najork, Marc. (2017). Web Crawler Architecture. *Encyclopedia of Database Systems*. DOI: 10.1007/978-1-4899-7993-3_457-3

3. Hovorushchenko T.O., Bodnar M.A., Kushnir V.O. (2019). Suchasni problemy formuvannia ta analizu vymoh do prohramnoho zabezpechennia. [Modern problems of the formation and analysis of software requirements. Measuring and computing equipment in technological processes]. *Vymiriuvalna ta obchysliuvalna tekhnika v tekhnolohichnykh protsesakh – Measuring and computing equipment in technological processes*, 1, 45–53 [in Ukrainian].

4. T. Hovorushchenko, O. Pavlova, M. Bodnar (2019). Development of an Intelligent Agent for Analysis of Nonfunctional Characteristics in Specifications of Software Requirements. *Eastern-European Journal of Enterprise Technologies*, 1, 2 (97), 6–17.

5. Paech, Barbara & Dutoit, Allen & Kerkow, Daniel & Knethen, Antje. (2002). Functional requirements, non-functional requirements, and architecture should not be separated. A position paper.

6. Alashqar, Abdelkareem & Elfetouh, Ahmad & El-Bakry, Hazem. (2015). Requirement Engineering for Non-Functional Requirements. *International Journal of Information and Communication Technology Research*, 5, 2127.

7. Kyrychenko O.L., Kanovsky I., Ostapov S.E. (2013). Prohramne zabezpechennia dlia doslidzhennia statystychnykh kharakterystyk hlobalnoi merezhi WWW [Software for researching the statistical characteristics of the WWW global network]. *Systemy obrobky informatsii. – Information processing systems*, 2: 3 (110), 99–104 [in Ukrainian].

8. Jaiswal, Manishaben, Software Architecture and Software Design (November 5, 2019). *International Research Journal of Engineering and Technology (IRJET)* e-ISSN: 2395-0056, p-ISSN: 2395-0072, Vol. 06, Is. 11, s. no-303, pp. 2452–2454. Available at: <https://www.irjet.net/archives/V6/i11/IRJET-V6I11303.pdf>, Available at SSRN: <https://ssrn.com/abstract=3772387> or <http://dx.doi.org/10.2139/ssrn.3772387>

9. Cook, Joshua, (2017). Docker for Data Science: Building Scalable and Extensible Data Infrastructure Around the Jupyter Notebook Server. Apress, Berkeley, 257 p.

10. Umar, Mubarak Albarka. (2020). Comprehensive study of software testing: Categories, levels, techniques, and types. DOI: 10.36227/techrxiv.12578714.v2.

ЗМІСТ

Анна КОРЧЕНКО, Сергій МАЦІЮК, Олександр ЧОБАЛЬ, Олександр КРУЧІНІН, Тарас ПАРАЩУК МЕТОД ФОРМУВАННЯ ПАРАМЕТРІВ ФУНКЦІОНАЛЬНИХ ОБОВ'ЯЗКІВ ДЛЯ ОЦІНКИ ЗАГРОЗ В СОЦІОТЕХНІЧНИХ СИСТЕМАХ.....	3
Лариса КОРЯШКІНА, Марина САЗОНОВА, Микола ОДНОВОЛ АЛГОРИТМИ ТЕРИТОРІАЛЬНОЇ СЕГМЕНТАЦІЇ ДЛЯ МЕРЕЖІ СЕРВІСНИХ ЦЕНТРІВ ІЗ ПЕРЕКРИТТЯМ ЗОН ОБСЛУГОВУВАННЯ.....	12
Володимир КУВАЄВ, Леонід МЕЩЕРЯКОВ, Альона ХАРЬ, Артем ПОЛІТОВ ІНТЕРФЕЙС ТЕХНІЧНОГО СУПРОВОДЖЕННЯ СКЛАДНИХ ІНФОРМАЦІЙНО-КЕРУЮЧИХ СИСТЕМ АВТОМАТИЗАЦІЇ.....	26
Борис МОРОЗ, Леонід КАБАК, Нонна ВАРЕХ, Дмитро МОРОЗ СИСТЕМА КЛАСИФІКАЦІЇ ТЕКСТОВИХ ДОКУМЕНТІВ ІЗ ВИКОРИСТАННЯМ ТЕХНОЛОГІЙ BIG DATA.....	34
Костянтин ОНИЩЕНКО, Анатолій АФАНАСЬЄВ ОЦІНКА ЕФЕКТИВНОСТІ ВИКОРИСТАННЯ ЗГОРТКОВИХ ТА РЕКУРЕНТНИХ НЕЙРОННИХ МЕРЕЖ У ЗАВДАННІ ОБРОБКИ ТЕКСТОВИХ ДАНИХ.....	41
Сергій ПРИХОДЬКО, Артем ТРУХОВ ПОБУДОВА ПРАВИЛ ПРИЙНЯТТЯ РІШЕНЬ ДЛЯ РОЗПІЗНАВАННЯ ОБЛИЧ НА ОСНОВІ КВАДРАТА ВІДСТАНИ МАХАЛАНОВІСА ДЛЯ НОРМАЛІЗОВАНИХ ДАНИХ.....	50
Олексій СЕРГЄЄВ, Світлана УС АНАЛІЗ СУЧАСНИХ ПІДХОДІВ ДО РОЗВ'ЯЗАННЯ ДИСКРЕТНИХ ТА НЕПЕРЕРВНИХ БАГАТОЕТАПНИХ ЗАДАЧ РОЗМІЩЕННЯ.....	59
Світлана СИВАШ, Галина СОКОЛОВСЬКА НАБЛИЖЕНЕ РОЗВ'ЯЗАННЯ ЗАДАЧІ СТАЦІОНАРНОЇ ТЕПЛОПРОВІДНОСТІ З ЕКСТРЕМАЛЬНОЮ ГРАНИЧНОЮ УМОВОЮ.....	71
Владислав СІДАНЧЕНКО, Олена НІКОЛЬСЬКА МЕТОДИ НЕЛІНІЙНОЇ ДИНАМІКИ В ЗАДАЧІ ПРОГНОЗУВАННЯ ХІМІЧНОГО СКЛАДУ ЧАВУНУ НА ВИПУСКУ.....	76
Iryna STOROSHKINA, Mykola ILYIN, Oleksandra PONOMARENKO SOCIAL ENGINEERING IN MODERN MESSENGERS: APPLICATIONS FOR OFFENSIVE SECURITY.....	84
Сергій ТКАЧЕНКО, Дмитро БЕШТА, Лілія БЕШТА РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ ТА ВПРОВАДЖЕННЯ В НАВЧАЛЬНИЙ ПРОЦЕС МОДЕЛІ-СИМУЛЯТОРА АВТОМАТИЗОВАНОЇ ФАБРИКИ.....	90
Володимир БАБИЧ, Анатолій КОСТЕНКО, Василь ПЛЕША, Михайло ПЛЕША, Леся ХМІЛЯРЧУК ЗАДАЧА ПОШУКУ НАЙКОРОТШОГО ШЛЯХУ: ПОРІВНЯЛЬНИЙ АНАЛІЗ ОСНОВНИХ АЛГОРИТМІВ.....	99
Оксана КИРИЧЕНКО ОСОБЛИВОСТІ АРХІТЕКТУРИ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ ДЛЯ ЗБОРУ ТА АНАЛІЗУ СТАТИСТИЧНОЇ ІНФОРМАЦІЇ В ГЛОБАЛЬНІЙ МЕРЕЖІ.....	107

CONTENTS

Anna KORCHENKO, Sergii MATSIUK, Oleksandr CHOBAL, Oleksandr KRUCHININ, Taras PARASCHUK THE METHOD OF FORMING THE PARAMETERS OF FUNCTIONAL RESPONSIBILITIES FOR ASSESSING THREATS IN SOCIOTECHNICAL SYSTEMS.....	3
Larysa KORIASHKINA, Maryna SAZONOVA, Mykola ODNOVOL ALGORITHMS OF TERRITORIAL SEGMENTATION FOR FACILITY NETWORK WITH OVERLAPPING SERVICE ZONES.....	12
Volodymyr KUVAIIEV, Leonid MESHCHERIAKOV, Alona KHAR, Artem POLITOV INTERFACE FOR TECHNICAL SUPPORT OF COMPLEX AUTOMATION INFORMATION CONTROL SYSTEMS.....	26
Borys MOROZ, Leonid KABAK, Nonna VAREKH, Dmitriy MOROZ TEXT DOCUMENT CLASSIFICATION SYSTEM WITH BIG DATA TECHNOLOGIES USAGE.....	34
Kostiantyn ONYSHCHENKO, Anatolii AFANASIEV EVALUATION OF THE EFFECTIVENESS OF USING CONVOLUTIONAL AND RECURRENT NEURAL NETWORKS IN THE TASK OF TEXT DATA PROCESSING.....	41
Sergiy PRYKHODKO, Artem TRUKHOV CONSTRUCTION OF DECISION RULES FOR FACE RECOGNITION BASED ON THE SQUARED MAHALANOBIS DISTANCE FOR NORMALIZED DATA.....	50
Oleksii SERHIEIEV, Svitlana US ANALYSIS OF MODERN APPROACHES TO SOLVING DISCRETE AND CONTINUOUS MULTI-STAGE ALLOCATION PROBLEMS.....	59
Svitlana SYVASH, Galya SOKOLOVSKA APPROXIMATE SOLUTION OF STATIONARY HEAT CONDUCTIVITY PROBLEM WITH EXTREME BOUNDARY CONDITION.....	71
Vladyslav SIDANCHENKO, Olena NIKOLSKA METHODS OF NON-LINEAR DYNAMICS IN THE PROBLEM OF FORECASTING THE CHEMICAL COMPOSITION OF CAST IRON AT THE OUTPUT.....	76
Iryna STOPOCHKINA, Mykola ILYIN, Oleksandra PONOMARENKO SOCIAL ENGINEERING IN MODERN MESSENGERS: APPLICATIONS FOR OFFENSIVE SECURITY.....	84
Serhii TKACHENKO, Dmytro BESHTA, Liliia BESHTA THE RESULTS OF THE RESEARCH AND INTRODUCTION INTO THE EDUCATIONAL PROCESS OF THE MODEL-SIMULATOR OF THE AUTOMATED FACTORY.....	90
Volodymyr BABYCH, Anatoly KOSTENKO, Vasyl PLESHA, Mykhailo PLESHA, Lesya KHMILYARCHUK THE PROBLEM OF FINDING THE SHORTEST PATH: A COMPARATIVE ANALYSIS OF THE MAIN ALGORITHMS.....	99
Oksana KYRYCHENKO FEATURES OF SOFTWARE ARCHITECTURE FOR COLLECTING AND ANALYZING STATISTICAL INFORMATION ON THE GLOBAL NETWORK.....	107

НОТАТКИ

INFORMATION TECHNOLOGY: COMPUTER SCIENCE, SOFTWARE ENGINEERING AND CYBER SECURITY

Випуск 2

Коректура • Ірина Миколаївна Чудеснова

Комп'ютерна верстка • Юлія Станіславівна Семенченко

Формат 60x84/8. Гарнітура Arial.

Папір офсет. Цифровий друк. Ум. друк. арк. 12,79. Замов. № 0823/543. Наклад 300 прим.

Видавництво і друкарня – Видавничий дім «Гельветика»

65101, Україна, м. Одеса, вул. Інглєзі, 6/1

Телефон +38 (095) 934 48 28, +38 (097) 723 06 08

E-mail: mailbox@helvetica.ua

Свідоцтво суб'єкта видавничої справи

ДК № 7623 від 22.06.2022 р.