

Міністерство освіти і науки України

Системні технології

System technologies

5 (148) 2023

Регіональний міжвузівський збірник наукових праць

Засновано у січні 1997 року.

У випуску:

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

Системні технології. Регіональний міжвузівський збірник наукових праць. – Випуск 5 (148). - Дніпро, 2023. – 167 с.

ISSN 1562-9945 (Print).

ISSN 2707-7977 (Online).

Редакційна колегія випуску:

Алпатов А.П. - д.т.н., проф. (відп. редактор)

Архипов О.Є. - д.т.н., проф.

Білозьоров В.Є. - д.ф.-м.н., проф.

Бабічев С.А. (Чеська Республіка) - д.т.н., доц.

Єрьомін О.О. - д.т.н., проф.

Прогресивні інформаційні
технології та організація
сучасного виробництва

Гече Ф.Е. - д.т.н., проф., (відп. редактор)

Гуда А.І. - д.т.н., проф.

Гнатушенко Вік.В. - д.т.н., проф.

Скалозуб В.В. - д.т.н., проф.

Математичне
та програмне забезпечення
інтелектуальних систем

Гнатушенко В.В. - д.т.н., проф., (відп. редактор)

Гожий О.П. - д.т.н., проф.

Кіріченко Л.О. - д.т.н., проф.

Светличний Д.С. (Польща) - д.т.н., проф.

Хандецький В.С. - д.т.н., проф.

Системні технології
обробки інформації
та кібербезпека

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

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

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

O.O. Zhulkovskyi, I.I. Zhulkovska, V.V. Kostenko, O.F. Bulhakova

RESEARCH ON SYNCHRONIZATION AND DATA PROTECTION PROBLEMS IN IMPLEMENTING MULTITHREADED PROGRAMS

Abstract. The issue of shared data usage by threads is especially relevant in modern multi core and multiprocessor systems. The main problems of implementing multithreaded programs are race conditions, deadlocks, and thread starvation. The aim of the work is to solve the problem of thread racing in multithreaded calculations of resource intensive tasks with parallel access to shared data using appropriate synchronization mechanisms, such as mutexes. A multithreaded algorithm for implementing a typical task of processing large data arrays with protection of the critical area in concurrent programs running on multiprocessor and multi core systems has been developed and researched.

Keywords: multithreaded computing, parallel programming, thread racing, synchronization mechanisms, mutexes.

Introduction. Urgent problems of modern development of processes and technologies require constant improvement of computer equipment, efficient use of its resources, processing of large volumes of data, and support for the growing requirements of modern information systems.

One of the approaches to address the mentioned problems is the use of multiprocessor and multi-core computing systems, where multiple physical or logical cores can be efficiently utilized, ensuring faster and more productive parallel solutions to specific resource-intensive tasks [1]. Software for computers is adapted and refined to utilize such systems. Parallel techniques, such as parallel algorithms, parallel databases, and parallel programming, are becoming increasingly important to ensure optimal system performance. Big data processing, training, and implementation of complex artificial intelligence algorithms, especially deep learning, have also become integral to many industries. Moreover, multithreaded computing remains a key technology in the powerful video game programming industry with high-quality graphics, the Internet of Things (IoT), and helps solve various tasks in many sectors.

Thus, researching synchronization methods and resource management in multithreaded environments, developing new methods and tools for parallel program-

ming with the creation of efficient multithreaded cross-platform programs, and optimizing algorithms and architectures for scalable computations are relevant tasks.

Literature Review. Multithreading is a property of a platform (e.g., an operating system, virtual machine, etc.) that allows a process created by the operating system to consist of several threads that are executed "in parallel", i.e., without a fixed order in time. During the execution of some tasks, such a distribution can achieve more efficient use of the computing machine's resources [2].

Full parallel execution of tasks is possible only in a multiprocessor (or multi-core) system. In the case of a single-processor multitasking system, processes are actually executed sequentially - here pseudo-parallel execution is used, creating the appearance of parallel work of several processes [2].

From the user's perspective, a process is an instance of a program during execution, and threads are branches of the program that are executed "in parallel". When one program performs many tasks, supporting multiple threads within one process allows distributing responsibility for different tasks among different threads, as well as increasing performance. Furthermore, tasks often need to exchange data, use shared data, or the results of other tasks. Threads within a process provide this capability, as they use the address space of the process to which they belong [2].

A thread can be in one of three states: executing, runnable, and waiting or blocked.

A multithreaded system can be implemented with scalability. For example, in parallel processing, the number of threads created can adapt to the number of processor cores in the system, allowing the program to accelerate within certain limits, fundamentally without changing its code.

When threads need to interact with each other or work with shared data, multithreading problems can arise, most of which are illustrated by the following classic tasks: about dining philosophers, about a sleeping barber, about smokers, about readers-writers, and others [2, 3].

The main problems of implementing multithreaded programs related to thread synchronization are [4–6]:

- race conditions, when two or more threads try to simultaneously access the same data or system resources without proper control, which can lead to unpredictable results;
- deadlocks, when two or more threads mutually block each other, waiting for mutual unlocking, as a result of which the program seems to freeze;
- thread starvation, when one thread is given excessive access to resources,

and other threads suffer from insufficient access to these resources.

Research Objective. The aim of this study is to solve the problem of thread racing in multithreaded calculations of resource-intensive tasks with parallel access to shared data using appropriate synchronization mechanisms, such as mutexes.

To achieve this goal, the following tasks are formulated: development of a multithreaded algorithm for implementing a typical task of processing large data arrays with protection of the critical area using synchronization primitives - mutexes; research of the performance of the developed algorithm with a significant amount of processed data and a variable number of computing threads; development of a concept for further application of effective approaches to data protection in concurrent programs implemented on multiprocessor and multi-core systems.

Research methodology and results. In this work, the problem of race condition is explored - an error in designing a multithreaded system where the program's operation depends on the order of code execution. Since this error is a «heisenbug», it can manifest randomly at different times, making it unpredictable and challenging to analyses, detect, and correct due to its variable behaviour.

The issue of shared data usage by threads is especially relevant in modern multi-core and multiprocessor systems, where many threads compete for access to shared resources, such as memory or files.

Research in this area aims to develop synchronization algorithms that allow threads to interact with shared data without conflicts and ensure the program's correct execution. It's also essential to optimize access to shared resources to ensure efficiency and speed of program execution. This issue is crucial for software developers as it allows efficient use of modern computing systems' capabilities and creates high-performance and reliable applications that can operate under intense resource competition.

When different threads have access to shared data, especially modifying them, synchronization of such access is required.

The thread race problem can be solved using synchronization and appropriate synchronization mechanisms, such as mutexes, conditional variables, atomic operations, and semaphores [2, 3].

Mutexes allow blocking access to shared data to one thread at a particular time. The thread that first locks the mutex gets access to the data, and other threads wait until the mutex is released.

When designing a program to reduce thread conflict, immutable design can avoid the race condition. That is, it's desirable to avoid shared access to mutable da-

ta from multiple threads or use immutable objects and data.

The choice of a specific method depends on the task and context.

Consider the thread race problem by solving a typical task of finding the sum of elements of a massive array in multithreaded mode.

The algorithm for solving this problem is implemented in C++ in the Microsoft Visual Studio 2022 IDE using the `std::vector` class - a dynamic array from the standard STL (Standard Template Library), which provides a convenient and safe way to manage memory and array size.

The program generates a large array of integers and uses a user-specified number of threads for parallel calculation of the sum of array elements. The `std::thread` class is used to create and manage threads in the program [7]. Each thread is responsible for a specific segment of the array, for which it calculates the sum of elements in a specially developed function. After the calculations are completed, the program displays the total sum and execution time.

To avoid thread races when several threads simultaneously try to access a shared variable storing the sum of array elements, the `std::mutex` class is used. The mutex ensures the correct calculation of the sum of the given array elements. In this program, to simulate a more significant mutex capture load, the mutex was captured on each iteration of the sum accumulation loop. Of course, in practice, using such an approach is justified only for experimental purposes.

For the experiments, the following PC infrastructure was used: Intel Core i7-12700H (14 cores, 2.3 GHz / 4.7 GHz); Goodram DDR4 (16 GB)×2 = 32 GB; Microsoft Windows 10; IDE Microsoft Visual Studio C++ 2022.

The size of the output array varied in the range from 100 million to 1 billion elements.

Fig. 1 presents the results of the implementation of the described program without blocking the critical area with a mutex, i.e., with the thread race problem and a deliberately incorrect result of the sum calculation.

As can be seen, with an increase in the size of the output array, the computation time significantly increases, while splitting the program into separate computational threads contributes to a noticeable increase in its performance. For example, with an increase in the array size in the range of 10^8 – 10^9 , the program execution time increases by 6–7 times in multithreaded (2–14 threads) and ten times in single-threaded implementation. The use of a multithreaded algorithm instead of a traditional single-threaded one provides a 1.5–2 times increase in program execution speed, depending on the array size. It should be noted that the results mentioned are

merely informative, as they were obtained without using mechanisms to avoid thread races, and the sum calculation here is most likely incorrect.

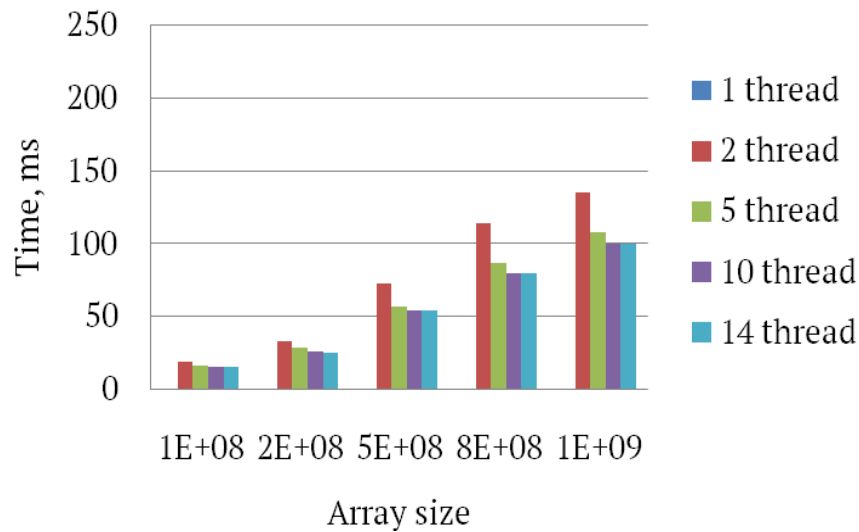


Figure 1 – Dependence of the computation time of the multithreaded program on the size of the output array (without blocking the critical area)

In Fig. 2, the results of the implementation of the described program are presented, addressing the thread race problem by blocking the critical area with a mutex, i.e., obtaining correct calculation results.

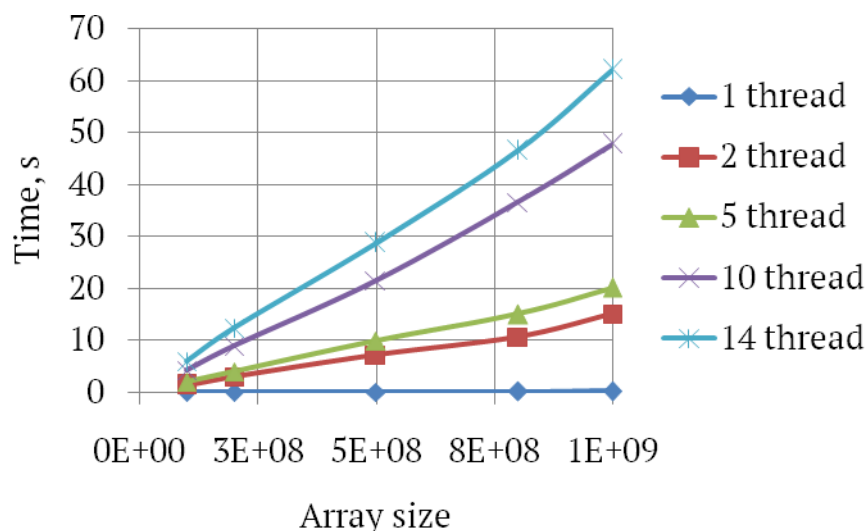


Figure 2 – Dependence of the computation time of the multithreaded program on the size of the output array (with blocking the critical area)

As can be observed, when solving the thread race problem in the considered task, with an increase in the array size in the range of 10^8 – 10^9 , the program execution time increases approximately tenfold. Increasing the number of used threads for parallel calculations with mutex capture on each iteration of the sum accumulation

loop also leads to a noticeable slowdown in program execution. For example, when increasing the number of threads from 2 to 14, the program execution time increases approximately four times, regardless of the output array size.

Each time a thread captures a mutex, other competing threads must wait until the current «owner» of the mutex releases it. The more threads there are, the more situations can arise where threads wait for access to the mutex, affecting the overall program execution time. The more cores in the processor, the more intense the competition for access to the mutex becomes. Additionally, when threads compete for processor time, scheduling operations can cause delays due to context switching between threads.

To improve performance with many threads, it is desirable to consider alternative methods in the future, such as using atomic operations, etc.

Also, promising directions are considered to be the use of special tools, for example, ROMP [8], DataRaceBench [9], and new approaches to detecting thread races based on deep neural network models [10].

Conclusions. As a result of the work, a multithreaded algorithm was developed to implement a typical task of processing extremely large data arrays with protection of the critical area to prevent the thread race problem using mutexes; the performance of the developed algorithm was investigated with a significant amount of processed data (10^8 – 10^9 elements) and a variable (1–14) number of computing threads; a concept was developed for the further application of effective approaches to data protection in concurrent programs implemented on multiprocessor and multi-core systems.

It was established that solving the thread race problem in the considered task on a modern PC with an Intel Core i7-12700H processor, with an increase in the array size in the range mentioned above, slows down the program execution approximately tenfold. Increasing the number of used threads from 2 to 14 during this slows down the application implementation approximately four times, regardless of the amount of processed data.

REFERENCES

1. Almeida S. An Introduction to High Performance Computing / S. Almeida // International Journal of Modern Physics A. – 2013. – vol. 28, no. 22n23, 1340021, p. 1–9. <https://doi.org/10.1142/s0217751x13400216>
2. Tanenbaum A. S. Modern Operating Systems / A. S. Tanenbaum, H. Bos. – New Jersey : Prentice Hall Press, 2014. – 1136 p. ISBN 978-0-13-359162-0
3. Williams A. C++ Concurrency in Action: Practical Multithreading / A. Williams. – 8

Island : Manning Shelter, 2019. – 592 p. ISBN 9781617294693

4. Kim S. Finding Semantic Bugs in File Systems with an Extensible Fuzzing Framework / S. Kim, M. Xu, S. Kashyap, J. Yoon, W. Xu, T. Kim // Proceedings of the 27th ACM Symposium on Operating Systems Principles (SOSP). – 2019. – p. 147–161. <https://doi.org/10.1145/3341301.3359662>
5. Xu W. Fuzzing File Systems Via Two-dimensional Input Space Exploration / W. Xu, H. Moon, S. Kashyap, P-N. Tseng, T. Kim // 2019 IEEE Symposium on Security and Privacy (SP). – 2019. – p. 818–834. <https://doi.org/10.1109/SP.2019.00035>
6. Xu M. Krace: Data Race Fuzzing for Kernel File Systems / M. Xu, S. Kashyap, H. Zhao; T. Kim. // 2020 IEEE Symposium on Security and Privacy (SP). – 2020. – p. 1643–1660. <https://doi.org/10.1109/SP40000.2020.00078>
7. Zhulkovskyi O. O. Evaluation of the Efficiency of the Implementation of Parallel Computational Algorithms Using the <thread> Library in C++ / O. O. Zhulkovskyi, I. I. Zhulkovska, V. V. Shevchenko, H. Ya. Vokhmianin // Computer Systems and Information Technologies. – 2022. – № 3. – p. 49–55. <https://doi.org/10.31891/csit-2022-3-6>
8. Gu Y. Dynamic Data Race Detection for OpenMP Programs / Y. Gu, J. Mellor-Crummey // SC18: International Conference for High Performance Computing, Networking, Storage and Analysis. – 2018. – p. 767–778. <https://doi.org/10.1109/SC.2018.00064>
9. Verma G. Enhancing DataRaceBench for Evaluating Data Race Detection Tools / G. Verma, Y. Shi, C. Liao, B. Chapman, Y. Yan // 2020 IEEE/ACM 4th International Workshop on Software Correctness for HPC Applications (Correctness). – 2020. – p. 20–30. <https://doi.org/10.1109/Correctness51934.2020.00008>
10. TehraniJamsaz A. DeepRace: A Learning-based Data Race Detector / A. TehraniJamsaz, M. Khaleel, R. Akbari, A. Jannesari // 2021 IEEE International Conference on Software Testing, Verification and Validation Workshops (ICSTW). – 2021. – p. 226–233. <https://doi.org/10.1109/ICSTW52544.2021.00046>

Received 01.10.2023.

Accepted 10.10.2023.

Дослідження проблем синхронізації та захисту даних при реалізації багатопоточних програм

Нагальні проблеми сучасного розвитку процесів та технологій потребують постійного підвищення продуктивності комп'ютерної техніки, ефективного використання її ресурсів, обробки великих обсягів даних та підтримки зростаючих вимог сучасних інформаційних систем. Одним із напрямків, що забезпечують ви-

рішення вказаних проблем, є використання багатопроцесорних та обчислювальних систем із багатоядерними процесорами, що забезпечують швидке та продуктивніше паралельне вирішення окремих ресурсномістких завдань.

У зв'язку з цим актуальними задачами є дослідження методів синхронізації та управління ресурсами в багатопоточних середовищах, розроблення нових методів та інструментів для паралельного програмування зі створенням ефективних багатопоточних крос-платформних програм, оптимізація алгоритмів та архітектур для масштабованих обчислень.

Метою даної роботи є вирішення майже важливішої серед проблем багатопоточного програмування – проблеми гонки потоків при багатопоточних обчисленнях ресурсномістких задач із паралельним доступом до спільних даних за допомогою використання м'ютексів.

В роботі вирішені наступні завдання: розроблено багатопоточний алгоритм реалізації типової задачі обробки надвеликих масивів даних із захистом критичної області за допомогою примітивів синхронізації – м'ютексів; досліджено продуктивність виконання розробленого алгоритму при значній (10^8 – 10^9 елементів) кількості оброблюваних даних і змінному (1–14) числі обчислювальних потоків (ядер); вироблено концепцію для подальшого застосування ефективних підходів до захисту даних у конкурентних програмах, що реалізуються на багатопроцесорних та багатоядерних системах.

Встановлено, що вирішення проблеми гонки потоків у розглянутій задачі на сучасному PC із процесором Intel Core i7-12700H (14 cores, 2.3 GHz / 4.7 GHz) при збільшенні розміру масиву у досліджуваному діапазоні уповільнює виконання програми приблизно в 10 разів. Збільшення при цьому числа використаних потоків від 2 до 14 уповільнює реалізацію застосунку приблизно в 4 рази, незалежно від кількості оброблюваних даних.

Жульковський Олег Олександрович – доцент кафедри програмного забезпечення систем, Дніпровський державний технічний університет.

Жульковська Інна Іванівна – доцент кафедри кібербезпеки та інформаційних технологій, Університет митної справи та фінансів.

Костенко Вікторія Вікторівна – старший викладач кафедри комп'ютерних наук та інженерії програмного забезпечення, Університет митної справи та фінансів.

Булгакова Ольга Федорівна – ст. викладач кафедри комп'ютерних наук та інженерії програмного забезпечення, Університет митної справи та фінансів.

Zhulkovskyi Oleg – Associate Professor of Department of Software Systems, Dniprovsky State Technical University.

Zhulkovska Inna – Associate Professor of Department of Cybersecurity and Information Technologies, University of Customs and Finance.

Kostenko Victoria – Senior Lecturer of Department of Computer Science and Software Engineering, University of Customs and Finance.

Bulhakova Olha – Senior Lecturer of Department of Computer Science and Software Engineering, University of Customs and Finance.

В.В. Герасимов, Н.В. Карпенко, І.А. Скуратовський

ДОСТУП ДО ЕЛЕМЕНТІВ СТРУКТУРИ І НЕВИЗНАЧЕНА ПОВЕДІНКА КОДУ НА МОВІ С

Анотація. Досліджено доступ до змінних типу даних структура на мові програмування С. Проаналізовано особливості невизначеної поведінки коду після компілятора Visual Studio 2019 під керуванням Windows 10 та компілятора gcc під керуванням ОС Linux. Показано результати нехтування повідомленнями при компіляції програм, при яких останні працюють некоректно. Надано рекомендації щодо уникнення цієї проблеми.

Ключові слова: невизначена поведінка, структура, компілятор, мова програмування С.

Постановка проблеми. Під час розробки програмного забезпечення розробники-початківці зазвичай отримують багато повідомлень про помилки та просто попереджень різного виду. І якщо є помилки, то код просто не запуститься, але при наявності різноманітних попереджень програма зазвичай запускається. І тут важливо розуміти, до яких наслідків може привести наявність попереджень різного роду.

Аналіз останніх досліджень і публікацій. Дуже часто попередження компілятора дратують програмістів і вони шукають в Інтернеті рішення для усунення цих попереджень через налаштування компілятора. Але досвідчені програмісти скажуть, що попередження ігнорувати не можна ні в якому випадку, оскільки вони можуть бути настільки важливими, що код написаний таким чином буде умовно робочим.

Інформація про те, які попередження існують в IDE Visual Studio, та які з них вимкнені за замовченням, можна знайти на сайті Microsoft [1]. Стаття [2] надає рекомендації щодо налаштування рівня попереджень для середовищ розробки Visual Studio, Code::Blocks та GCC/G++. Також тема про попередження компіляторів піднімається в книзі про покращення коду, написаного на мові C++ автора Scott Meyers [3]. Досить цікавою і корисною є стаття з прикладами усунення помилок сегментації для мов програмування C/C++ [4]. Однак основна маса статей містить лише загальні відомості без прикладів. І коли програ-

міст стикається з чимось незрозумілим, він починає шукати відповідь і не находячи її, залишає своє питання в професійних блогах чи на форумах, але на деякі питання так і немає відповідей.

Метою даної роботи є дослідження невизначеної поведінки коду при роботі зі структурами на мові програмування C при видачі відповідного попередження компілятора про повернення адреси локальної або тимчасової змінної.

Основна частина. Студенти, які тільки почали програмувати, фокусуються на тому, щоб написана програма працювала і перевіряють її на обмеженій кількості комплектів тестових даних. Це приводить до того, що вони можуть навіть і не помітити, коли в програмі є невизначеність поведінки [5].

В процедурній мові програмування C є пращур класу ООП — структура `struct`, яка інкапсулює лише стан сутності. І виникає питання — а чи можна працювати з окремими компонентами-полями такої структури аналогічно мовам ООП?

Розглянемо проблему доступу до елементів структури на мові C на прикладі найпростішої інформаційної системи "Телефонний довідник", де в якості структури даних будемо використовувати наступну структуру [6] (лістинг 1).

Лістинг 1

```
typedef struct {
    int id;
    char firstName[20];
    char lastName[20];
    char number[20];
} person;
```

Для оптимальної реалізації функціоналу пошуку за ім'ям, прізвищем або телефоном потрібно загальну частину оформити у вигляді окремої функції, в якій певна частина структури буде порівнюватись із введеним текстом. В об'єктно-орієнтованих мовах програмування високого рівня (C#, Java тощо) для отримання частини структури (класу) зазвичай використовують так звані гетери — спеціальні функції, призначені для повернення значень полів екземпляру класу (його стану). Спробуємо реалізувати щось подібне на мові C (лістинг 2).

Лістинг 2

```
char* getFirstName(person person)
{
    return person.firstName;
}
```

Ця функція повертає ім'я людини. Аналогічним чином можна реалізувати функції для виокремлення прізвища та номера телефону. Тоді загальна функція для пошуку за будь-якою частиною структури буде мати наступний вигляд (лістинг 3).

Лістинг 3

```
void searchByPartInfo(char* part, char* (*func)(person person))
{
    char arr[FIELD], * info;
    printf("Enter %s\n", part);
    scanf_s("%s", arr, FIELD);
    int counter = 0;

    for (int i = 0; i < N; i++)
    {
        if (massiveStruct[i].id != -1)
        {
            if (strstr(func(massiveStruct[i]), arr))
            {
                printItem(i);
                counter++;
            }
        }
    }

    if (counter == 0)
    {
        puts("Not found");
    }
}
```

де `char* (*func)(person person)` — вказівник на наш псевдогетер, який повертатиме необхідну частину структури даних.

При компіляції вихідного коду в середовищі Visual Studio видаються попередження відносно наших псевдогетерів — `warning C4172: повернення адреси локальної або тимчасової змінної: person`. Але тим не менш, код запускається і працює, що може приспати пильність розробника коду, особливо недосвідченого. А ось розробники зі стажем можуть відразу згадати — `Undefined behaviour`, невизначена поведінка. В чому ж проявляється ця невизначена поведінка?

Тестування коду показало, що пошук за ім'ям і прізвищем в такому випадку працює бездоганно, але за номером телефону (який стоїть на останньому місці під час оголошення структури) — не працює зовсім. Щоб зрозуміти про-

блему необхідно за допомогою налагоджувача (debugger) прослідкувати за ходом виконання коду. Для цього створимо невелику функцію-обгортку (лістинг 4).

Лістинг 4

```
int compare(char str[], char sub[])
{
    return strstr(str, sub);
}
```

Відразу після додавання такої функції-обгортки картина дещо змінилась, а саме в деяких випадках пошук почав виконуватись. Детальний аналіз показав, що відразу після входу в функцію-обгортку символьний рядок, який містить номер телефону залишається без змін (рис. 1,а). Але функція для порівняння символьних рядків strstr повертає рядок, в якому без змін залишаються перші 8 цифр, інші — змінюються випадковим чином (рис. 1,б). Тому якщо область пошуку обмежена тільки першими 8 цифрами — результат успішний.

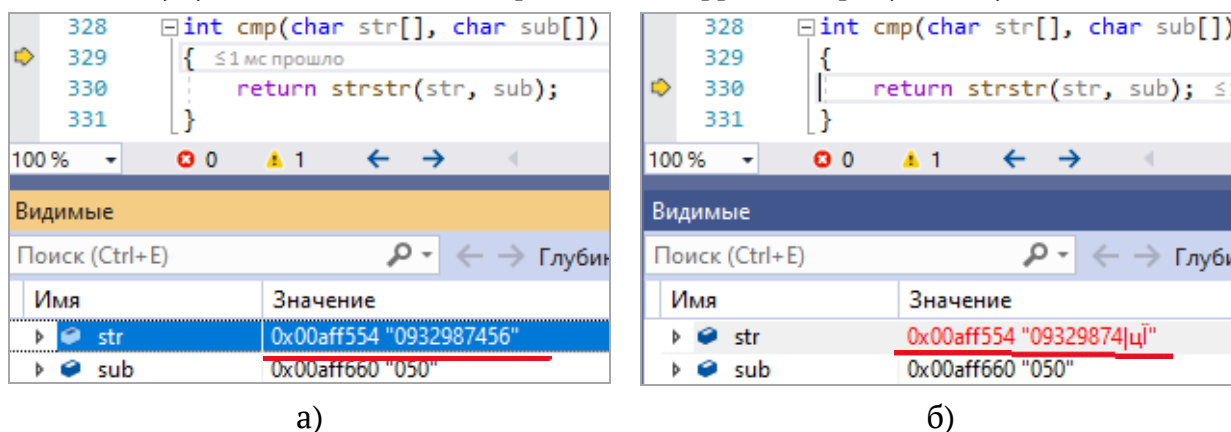


Рисунок 1 – Виявлення моменту заміни символів початкового рядку:
а) вхід в функцію-обгортку; б) результат, який повертає функція strstr

Були спробувані різні варіанти вирішення проблеми. Масив оголошувався і як статичний, і як константний — жодне рішення не допомогло вирішити проблему. Незмінним залишалось лише одне — перші 8 цифр. Решта масиву або змінювалась випадковим чином, або просто відкидалась. Передати на обробку частину структури у вигляді масиву символів-цифр не вдавалось.

Потім були проведені дослідження як змінюється значення останнього члена структури в залежності від довжини масиву (`char number[20]`). Так при наявності 19 цифр (останній байт відводиться під символ закінчення рядку) в цій частині структури, яку ініціалізували цифрами 1234567890123456789, був отриманий результат, представлений на рис. 2, а саме: при вході в функцію-

оболонку масив обрізається до 12 символів (рис. 2,а), а на наступному кроці — обрізається ще раз до 8 символів і доповнюється випадковими символами (рис. 2,б).

Подібні експерименти з іншими членами структури, які стоять не на останньому місці, до таких результатів не призвели. Тобто така поведінка при передачі між функціями спостерігається лише для останнього члена структури, який є символьним масивом.

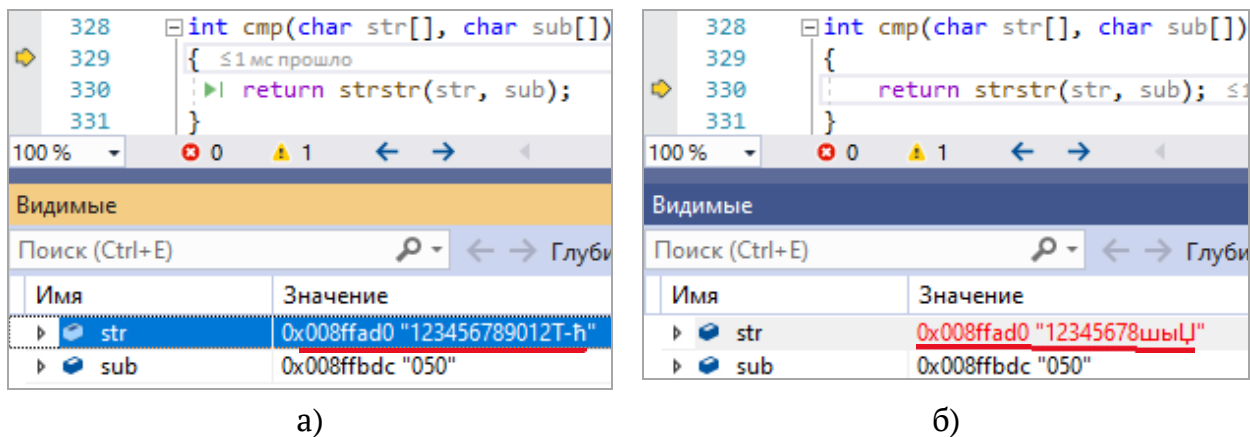


Рисунок 2 – Заповнення останнього члена структури цифрами на максимальну ємність: а) значення на вході в функцію-оболонку; б) результат, який повертає функція `strstr`

Також були проведені дослідження для різної кількості байт пам'яті, яку програміст відводить для останнього члена структури під час оголошення. Так, у випадку розміру масиву в 10 символів (`char number[10]`) при вході у функцію-оболонку залишалось лише 4 символи, а на наступному кроці початкових символів вже не залишалось зовсім, були лише випадкові символи. У випадку масиву в 40 символів (`char number[40]`) при вході у функцію-оболонку залишалось 32 символи, а на наступному кроці їх залишалось лише 28.

Переміщення цілого поля структури в її кінець (поля `id`) в цілому картину не змінило — при передачі в якості аргументу останнього символьного масиву спостерігається втрата інформації, змінюється лише кількість символів, які залишаються без змін.

Дослідження були проведені в середовищі Visual Studio 2019 під керуванням операційної системи Windows 10, коли в налаштуваннях проекту був обраний стандарт мови C за умовчанням — MSVC. Коли стандарт був змінений на більш сучасний стандарт ISO C17 (2018) [7] — в останній частині структури, яка

відповідала за номер телефону, незмінними залишались лише 4 цифри, інші змінювались випадковим чином.

Таким чином ми бачимо дійсно невизначену поведінку, коли результат виклику функції залежить від багатьох чинників: від довжини масиву, від стандарту мови C, від положення масиву в структурі.

А ось спроба провести аналогічні дослідження під керуванням операційної системи Linux Mint з використанням компілятора gcc версії 5.4 не вдалась. При компіляції коду також видавалось попередження про повернення адреси локальної змінної. А ось при спробі здійснити пошук за допомогою вищенаведеного коду за будь якою частиною структури даних (ім'я, прізвище чи телефон) програма просто аварійно закінчувала роботу з повідомленням про помилку сегментування [4]. Відлагодження вихідного коду показало, що у функцію-оболонку передається пустий вказівник, тобто наш "псевдогетер" не повертає взагалі нічого.

Пояснення такої поведінки є досить простим, для цього потрібно лише детально розглянути наш псевдогетер (лістинг 2). В мові програмування C структура передається в метод за значенням, а це означає, що при вході в функцію ця структура буде скопійована в локальну змінну, отже всі поля цієї структури будуть скопійовані в локальній області видимості (в межах методу). При цьому в самій структурі в нас фігурують не вказівники, а саме масиви символів, і тому наші поля структури будуть скопійовані за новими адресами. В псевдогетері ми повертаємо вказівник на перший елемент масиву символів, який після виходу за межі області видимості функції буде помічений як видалений. Тобто система після виходу з методу вважає, що може знову використовувати область пам'яті, яку займав цей масив символів. Отже на наступних ітераціях в цю частину пам'яті можуть бути записані інші дані, причому не обов'язково тієї самої довжини, і збереження значення змінної за цією адресою не гарантується. Ми можемо спостерігати як правильну поведінку, так і повний або частковий перезапис даних.

Всі ці дослідження були проведені для того, щоб з'ясувати — в чому саме полягає невизначена поведінка коду? Ну а для того, щоб все-таки вирішити поставлену задачу доступу до елементів структури на мові C в псевдогетерах та інших подібних випадках необхідно передавати структуру за адресою (посиланням). Або як у випадку найпростішої інформаційної системи можна скористатись глобальним масивом структур і оперувати відповідно тільки з індексами цього масиву [6].

Висновки. У роботі проведено дослідження невизначеної поведінки коду, написаного на мові програмування C, при передачі членів структури в функцію.

Компілятори Visual Studio 2019 і gcc надавали повідомлення про невизначену поведінку коду. Але ця поведінка кардинально відрізнялась для різних операційних систем та різних компіляторів. Якщо після gcc під ОС Linux код просто зовсім не працює і програма зупиняє свою роботу з повідомленням про помилку сегментації, то після Visual Studio під керуванням Windows код запускався, але працював коректно лише на обмежених комплектах даних.

Було показано «підводні камені» різноманітних рішень для виправлення цієї помилки. Таким чином, невизначена поведінка коду усувається передачею в функцію адреси структури, з членом якої потрібно працювати.

ЛІТЕРАТУРА

1. Ошибки и предупреждения компилятора и средств сборки C/C++. URL: <https://learn.microsoft.com/ru-ru/cpp/error-messages/compiler-errors-1/c-cpp-build-errors?view=msvc-170>
2. Конфигурация компилятора: Уровни предупреждений и ошибки. URL: <https://ravesli.com/konfiguratsiya-kompilyatora-urovni-preduprezhdenij-i-oshibki/>
3. Мэйерс С. Эффективное использование C++. 55 верных способов улучшить структуру и код ваших программ. – М.: ДМК Пресс, 2006. – 300 с. – ISBN 5-94074-304-8, 0-321-33487-6.
4. Segmentation Fault in C/C++ - GeeksforGeeks. URL: <https://www.geeksforgeeks.org/segmentation-fault-c-cpp/>
5. Герасимов В. В., Карпенко Н. В. Проблема невизначеної поведінки коду на мові C при роботі зі структурами // XI Міжнародна науково-практична конференція "Сучасні проблеми і досягнення в галузі радіотехніки, телекомунікацій та інформаційних технологій", м. Запоріжжя, НУ "Запорізька політехніка", 12-14 грудня 2022 р., с. 140 – 142.
6. Карпенко Н. В., Герасимов В. В. Сучасний підхід до програмування на мові C від нульового до просунутого рівня: навчальний посібник — Дніпро, ПП «Ліра ЛТД», 2022. — 418 с.
7. The Standard – C. URL: https://www.iso-9899.info/wiki/The_Standard

REFERENCES

1. Oshybky i preduprezhdeniya kompilyatora i sredstv sborky C/C++. URL: <https://learn.microsoft.com/ru-ru/cpp/error-messages/compiler-errors-1/c-cpp-build-errors?view=msvc-170>

2. Konfiguratsiya kompilyatora: urovni preduprezhdenij I oshibki URL: <https://ravesli.com/konfiguratsiya-kompilyatora-urovni-preduprezhdenij-i-oshibki/>
3. Meyers S. Efektivnoe ispolzovanie C++. 55 vernykh sposobov uluchshyt strukturu i kod vashykh programm. – M.: DMK Press, 2006. – 300 s. – ISBN 5-94074-304-8, 0-321-33487-6.
4. Segmentation Fault in C/C++ - GeeksforGeeks. URL: <https://www.geeksforgeeks.org/segmentation-fault-c-cpp/>
5. Herasymov V. V., Karpenko N. V. Problema nevyznachenoї povedinky kodu na movi C pry roboti zi strukturamy // XI Mizhnarodna naukovo-praktychna konferentsiia "Suchasni problemy i dosiahnennia v haluzi radiotekhniky, telekomunikatsii ta informatsiinykh tekhnolohii", m. Zaporizhzhia, NU "Zaporizka politekhnika", 12-14 hrudnia 2022 r., s. 140 – 142.
6. Karpenko N. V., Herasymov V. V. Suchasnyi pidkhid do prohramuvannia na movi C vid nulovoho do prosunutoho rivnia: navchalnyi posibnyk — Dnipro, «Lira LTD», 2022. — 418 p.
7. The Standard – C. URL: https://www.iso-9899.info/wiki/The_Standard

Received 12.10.2023.

Accepted 15.10.2023.

Access to struct members and undefined behavior of C code

During software development, novice developers usually receive a lot of error messages and just warnings of various kinds. And if the code simply won't run when there are errors, then the program usually starts when there are various warnings. And here it is important to understand what consequences the presence of warnings of various kinds can lead to.

This work aims to study the code's undefined behavior when working with structures in the C programming language when issuing a corresponding compiler warning about returning the address of a local or temporary variable.

In the procedural programming language C, there is an ancestor of the OOP class — the structure struct, which encapsulates only the state of the entity. And the question arises — is it possible to work with separate components-fields of such a structure analogously to OOP languages?

For research, a simple structure was taken, which contains information about the person's name, surname, and phone number. To access parts of the structure, pseudogetters were used — functions that returned a pointer to the corresponding part of the structure.

The research was conducted in the Visual Studio 2019 environment under the control of the Windows 10 operating system when the default C language standard - MSVC and the more modern ISO standard C17 (2018) was selected in the project settings.

As a result, a truly undefined behavior of the code was obtained, when the result of the work of the code fragment (function call) depends on many factors: the length of the array, the standard of the C language, the position of a certain part in the structure.

An attempt to conduct similar research under the control of the Linux Mint operating system using the gcc compiler version 5.4 was unsuccessful. When compiling the code, a similar warning about returning the address of a local variable was also issued, as in the case of Visual Studio. But when the program was launched, it simply crashed with a message about a segmentation error.

Thus, both the Visual Studio 2019 compiler and the gcc compiler warned us about undefined code behavior. But this uncertain behavior was radically different for operating systems and compilers. If after gcc under the Linux OS, the code simply does not work at all and the program stops its work with a segmentation error message, then after Visual Studio under Windows, inexperienced developers with improper testing and verification of the code can miss the code that "does not always work", which can lead to unexpected results, not always pleasant, to say the least. And that's why software developers, especially beginners, should pay attention not only to compilation errors but also to warnings, even if the code works.

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

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

Скуратовський Ігор Анатолійович - к.ф.-м.н., доцент кафедри ЕОМ Дніпровського національного університету імені Олеся Гончара.

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

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

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

ADAPTATION OF THE WORLD FRAMEWORK FOR FRAME-BY-FRAME REAL-TIME SPEECH ANALYSIS

Abstract. WORLD is a vocoder based speech synthesis system developed by M. Morise et al. and implemented in C++. It was demonstrated to have improved performance and accuracy when compared to other algorithms. However, it turned out to not perform well in certain scenarios, particularly, when applying the framework to very short waveforms on a frame by frame basis. This paper reviews the issues of the C++ implementation of WORLD and proposes modified versions of its constituting algorithms that attempt to mitigate those issues. The resulting framework is tested on both synthetic signals and on real recorded speech.

Keywords: speech analysis, speech synthesis, real time signal processing, spectral envelope, F0 estimation

Introduction. The speech analysis and synthesis were always popular subjects of interest for many researchers. The techniques and methods of working with speech signals developed over the last decades are now used in a variety of applications such as text-to-speech (TTS) synthesis, speaker identification (SI), voice conversion (VC), speech to text (STT) recognition, and others. One important characteristic of these applications share is that they have to operate with real-time constraints. A TTS system has to produce output from the user-provided text with minimum lag, a SI has to identify speakers quickly, a VC system should be able to operate on a signal and produce continuous output with minimal latency, etc.

There are systems of each type that satisfy these requirements and produce acceptable results. One of such systems is the WORLD framework for high-quality speech analysis and synthesis [1]. WORLD is a vocoder-based speech synthesis system that decomposes the speech signal into a fundamental frequency (F0) using the DIO algorithm [2], a spectral envelope using the CheapTrick method [3], and an excitation signal using PLATINUM algorithm [4]. Each of these components is usually estimated for every millisecond of the speech signal, forming a contour of estimations. The original signal can then be resynthesized by convolving the minimum phase response of the spectral envelope with the excitation signal. The

diagram of WORLD framework is illustrated in Figure 1. Speech resynthesized in this way was shown to be of high quality in terms of waveform similarity with the input signal.

Literature Review. Unlike systems that perform direct sequence-to-sequence conversion on raw audio and usually utilize neural networks and deep learning for modeling non-linear mappings for speech-to-speech [1, 2], text-to-speech [3, 4], or speech-to-text [5, 6] conversions, frameworks like WORLD rely on conventional tools, are highly customizable, and make for a great basis for developing new systems.

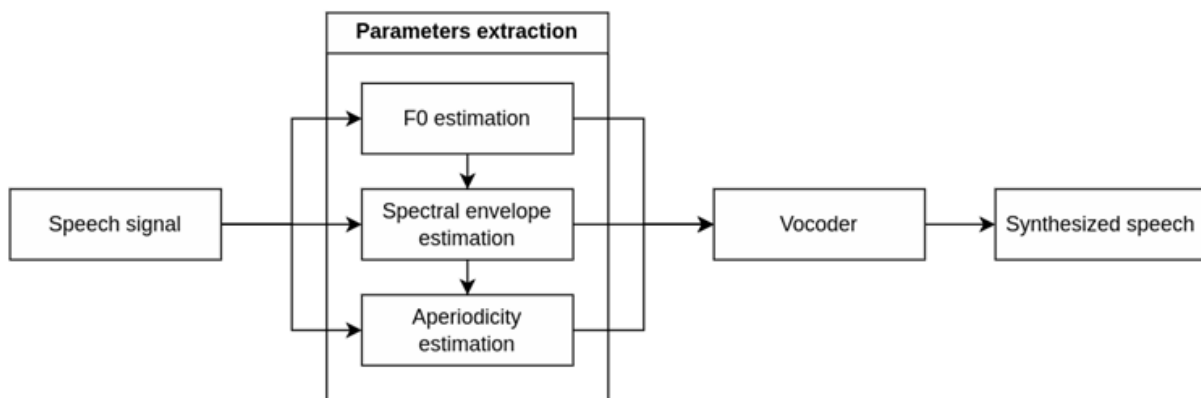


Figure 1 - WORLD framework diagram. The original signal is decomposed into three components that are used to drive the vocoder and produce high-quality speech

Unfortunately, the C++ implementation of WORLD is ill-suited for one specific but very popular use case – frame-by-frame real-time processing of audio signals. This type of processing is common in digital audio systems and used ubiquitously in music production, recording, and sound design. Such systems include, for example, Jack audio server for Linux – a system that is intended for environments in which low latency (up to 5ms) is critical [7]. Jack operates by connecting audio sources to audio sinks and passing pieces of signal between them using very small buffers that usually contain from 64 to 1024 samples depending on the configuration. WORLD requires buffers that are at least three times longer than the period of the lowest signal the system is expected to correctly process. So, if the speech signal is expected to have the lowest pitch of 60Hz or have period of 16ms, then the input should be at least 50ms long or contain at least 800 samples at the 16 kHz sampling rate. Jack usually doesn't provide such long buffers, but this can be solved by allocating a buffer of required size and outputting only silence while it is filling up. This approach will not impact the overall latency but will introduce delay between the

first non-silent input and the first non-silent output. Even with this approach however, the WORLD framework fails at the very first stage – the F0 estimation. It either crashes or provides incorrect estimations after which all the other stages of sound decomposition are impossible. If the buffer sizes are increased significantly (above about 2000 samples), the algorithm succeeds at estimating the F0 but the following stage, the spectral envelope estimation, produces unstable and noisy results that introduce heavy distortion when the sound is resynthesized.

The following sections propose alternative algorithms to those implemented in WORLD in order to enable the framework to work in the environment described earlier.

Research methodology and results. F0 estimation. Fundamental frequency or pitch estimation is a process that determines the minimum frequency at which the analyzed waveform repeats itself. There are multiple and very different approaches to estimating F0 of a given signal. These approaches can be broadly categorized into spectrum- and time domain-based methods. The spectrum-based methods focus on detecting peaks and distances between them in the spectral envelope or the cepstrum of the signal while the time domain-based methods attempt to detect periods between peaks and zero crossings of the raw signal in the time domain [1]. When selecting a F0 detection method for a particular application multiple factors must be taken into account, such as the algorithm's accuracy to performance ratio, the expected noise level of the signal, the time constraints, required signal frame length, etc. It was shown that spectrum-based algorithms perform better than the time domain-based ones [2], but their drawback is that they introduce extra computation cost by requiring the spectrum or cepstrum to be computed for each frame of the signal. For this reason, the WORLD framework uses the time domain-based DIO algorithm for estimating F0 quickly and cheaply. However, the C++ implementation of DIO prevents WORLD from supporting the frame-by-frame analysis and synthesis. This section discusses the adaptation of DIO to enable this support.

As the basis of the reimplementations, the improved version of DIO was selected [3] and modified to make it work with short speech frames. The implementation itself was done in Julia programming language to make it easier to apply in data exploration scenarios. Below the detailed description of the final algorithm is provided. Our implementation of DIO will be referred to as JDIO, the algorithm by Daido and Hisaminato will be referred to as YDIO, and the C++

implementation from the WORLD framework will be referred to as CDIO in the text below.

DIO's basic idea is to narrow down the possible pitch values as much as possible, detect the most stable F0 candidates for a given frame, and then select one candidate that preserves the smoothness of the corresponding voiced region's F0 contour the best. The first step can be thought of as preprocessing, and it includes decimating (downsampling) the input signal so that its sample rate is just enough to represent the highest desirable frequency accurately, removing the DC component from the decimated signal, and applying parallel cascading lowpass filters to attenuate harmonics and make the low pitch detection more reliable. So, the preprocessing produces a collection of signals that have a significantly reduced number of samples due to decimation, have mean amplitude of zero, and differ only by the increasingly aggressive lowpass frequency cutoff. The second step involves applying a period detector to each of these signals and assessing the stability of a period detected in each of them. The final step simply selects one of the detected candidates based on a number of heuristics that minimize the deviation from the F0 estimate for the previous frame. The flowchart representation of the algorithm is given on Figure 2.

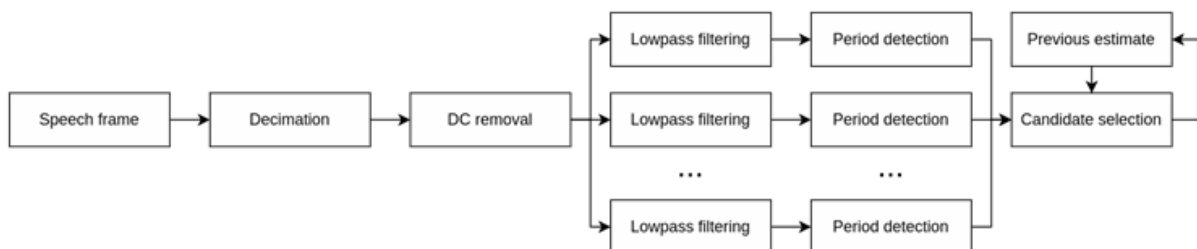


Figure 2 - Top-level representation of the DIO algorithm. Each component can be tweaked to achieve optimal performance for a given application

In JDIO, the decimation, DC component removal, and the lowpass filtering steps were delegated to the DSP.jl package routines [8], while the period detection and candidate selection were implemented using custom algorithms.

Just like in YDIO, the JDIO's period detection is based on marching through the signal and identifying time durations between pairs of events of 4 types: positive zero crossings, positive peaks, negative zero crossings, and negative peaks. To make identification of these events more reliable, each of them is searched for one after the other in the provided order and in circular fashion, i.e. after detecting the negative peak, the algorithm starts searching for positive crossing. Once an event is

detected at time t_j , the new search starts at $t_j + \tau_{\min}$. If the new search fails to find any new events until $t_j + \tau_{\max}$, the search is considered failed. $\tau_{\min}$ and $\tau_{\max}$ are constants that put hard lower and upper bounds on the range of frequencies the algorithm can detect. The search process stops once 8 successful detections have been performed, i.e. when two points of each event type have been found, and a sequence of detected timestamps $[t_1, t_2, \dots, t_8]$ is obtained. From these timestamps, event periods for each event type can be calculated: $p_i = t_{i+4} - t_i, \forall i \in [1..4]$. The process of the period detection is visualized on the Figure 3. The estimated fundamental period of the frame is then given by the average of the periods of each event type:

$$P_0 = \frac{1}{4} \sum_{i=1}^4 p_i \Rightarrow F_0 = \frac{1}{P_0}. \quad (1)$$

Each estimation (1) is assigned a “confidence” value given by how much each event period deviates from their average value:

$$\lambda = 1 - \min \left[1, \frac{\sum_{i=1}^4 |p_i - P_0|}{4P_0} \right] \quad (2)$$

The estimates with confidence lower than a given threshold $\lambda_{\min}$ are rejected, the best F0 candidate is selected among the remaining estimates.

Even though the base algorithm is fairly straightforward, its real world application presents a number of additional challenges.

First, the peak detection can suffer from a high number of false positives when the signal is noisy or contains harmonics. To mitigate this problem to some extent, two thresholds are introduced – noise floor and relative peak amplitude. A peak point candidate s_j is skipped during the peak detection process if its amplitude is lower than the noise floor or does not reach at least α fraction of the previous peak of the opposite sign:

$$|s_j| > \eta \wedge |s_j| > |\alpha s_{prev,peak}|. \quad (3)$$

The noise threshold also automatically filters out any unvoiced regions that otherwise would have to be identified and removed by other means.

Second, the signal being processed is discrete and the points in time at which it's sampled generally speaking do not align with the true peaks and zero crossings. This issue becomes even worse when the signal is decimated and its sampling rate is reduced significantly. An example of this is provided on the Figure 4. To somewhat mitigate this issue, linear and quadratic interpolations can be used. Since zero crossings are detected by searching two consecutive points (t_j, s_j) and (t_{j+1}, s_{j+1}) on

the opposite sides of the X axis, the true zero crossing can be estimated by linearly interpolating between them and solving for zero:

$$t_{zero} = -s_j \frac{t_{j+1} - t_j}{s_{j+1} - s_j} + t_j \quad (4)$$

For detecting peaks, three consecutive points (t_{j-1}, s_{j-1}) , (t_j, s_j) , and (t_{j+1}, s_{j+1}) that satisfy the condition $s_{j-1} < s_j > s_{j+1}$ have to be found. The true peak has to be greater than or equal to the middle point (t_j, s_j) and to approximate it, a quadratic interpolation can be used. In other words, a parabola can be fitted to the three points and its extremum selected as a true peak estimate:

$$\begin{bmatrix} t_{j-1}^2 & t_{j-1} & 1 \\ t_j^2 & t_j & 1 \\ t_{j+1}^2 & t_{j+1} & 1 \end{bmatrix} \begin{bmatrix} a \\ b \\ c \end{bmatrix} = \begin{bmatrix} s_{j-1} \\ s_j \\ s_{j+1} \end{bmatrix} \quad (5)$$

The a , b , and c coefficients for the $f(x) = ax^2 + bx + c$ equation are given by the solution of the equation (5). They can be used to obtain the coordinates of the fitted curve's extremum:

$$\begin{aligned} t_{peak} &= -b/2a \\ s_{peak} &= -b^2/4a + c \end{aligned} \quad (6)$$

It is important to also calculate the amplitude value s_{peak} to use it in the condition (3) check during the next peak detection.

Third, since there are multiple period detectors being applied in parallel to the same signal filtered with different lowpass filters as pictured in Figure 2, multiple F0 candidates $C = \{c_1, c_2, \dots, c_k\}$ are obtained, and the candidate selection procedure has to be put in place. JDIO uses the approach that minimizes the difference between the previous estimate and the new one:

$$F_{n+1} = c: \min_{c \in C} |c - F_n|. \quad (7)$$

The confidence of the new estimates is guaranteed to be acceptable because all candidates with confidence lower than λ_{min} are rejected. This approach helps to prevent frequency halving and frequency doubling errors for harmonic signals.

Four, since the DIO algorithm is usually used for estimating frequency contours for frames with a very short hop time, it is necessary to ensure that the event detection chain starts as close to the start of the frame as possible. To achieve this, JDIO runs each event detector and selects the one that succeeds earlier than all others. The example of different detectors being selected as first ones in the chain can be seen on the Figure 3. In 3a, the first detected event is a positive zero crossing while in 3b, the first detected event is a positive peak that satisfies the condition (3).

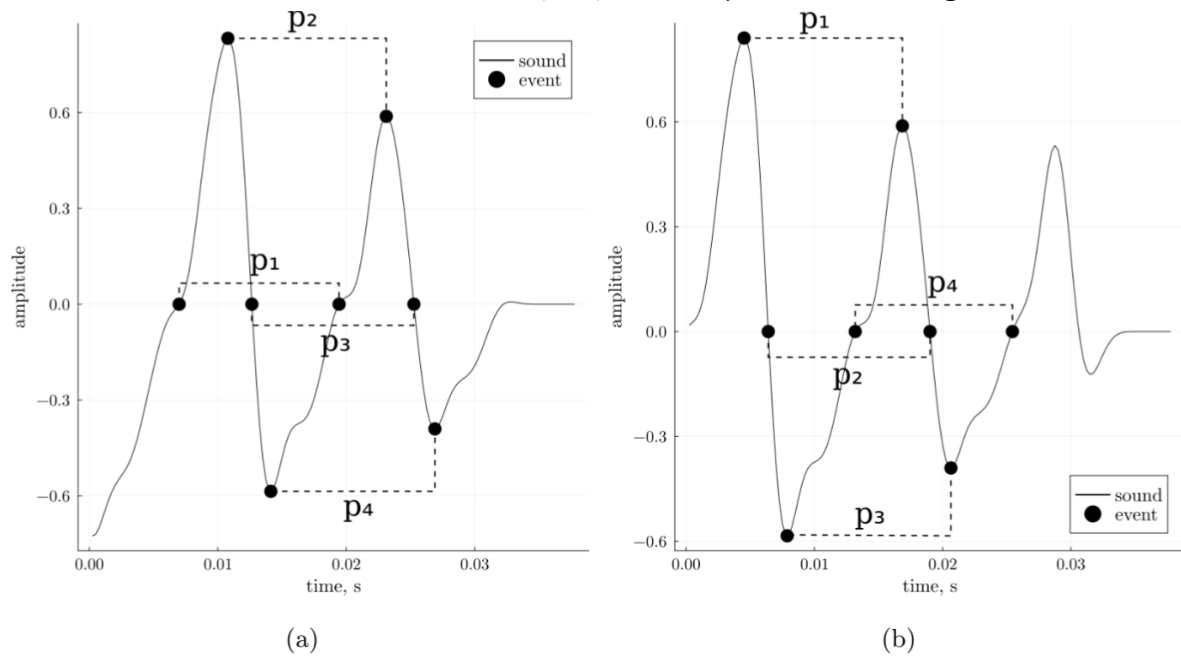


Figure 3 - JDIO detects periods between different kinds of events from a frame of signal

In both cases, the algorithm was able to detect the best chain of events that is closest to the start of the frame and estimate the F0 successfully.

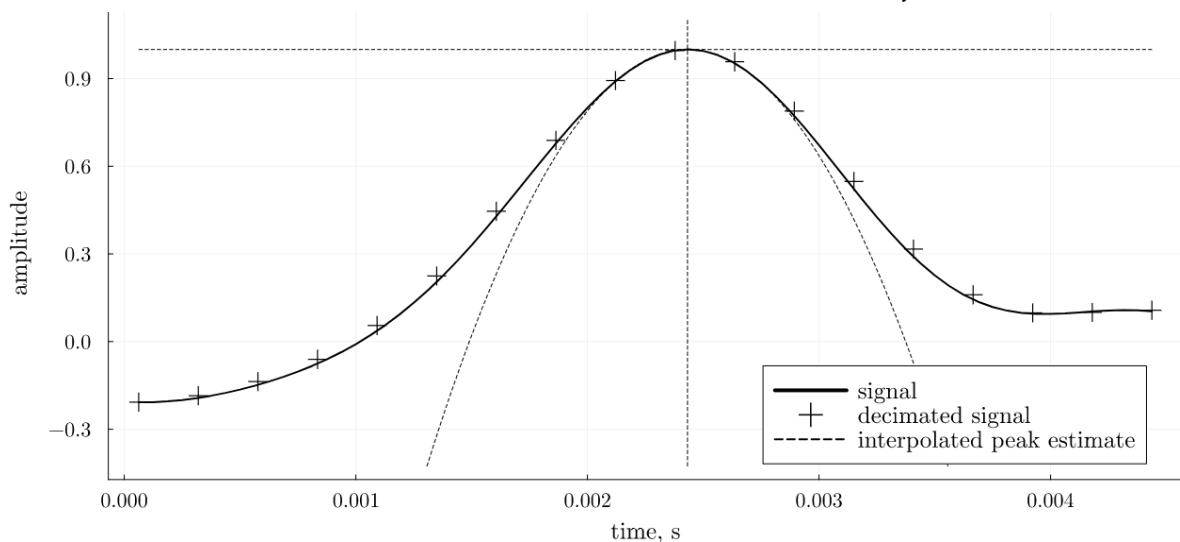


Figure 4 - Positive peak detection. A parabola is fitted to three points that are known to be around the true peak, and the first coordinate of its extremum is selected as the time value at which the true peak occurs

Spectral envelope estimation. A spectral envelope is a curve in the frequency-amplitude plane, derived from a Fourier magnitude spectrum. The spectral envelope should ideally be a smooth line with minimal oscillations that wraps tightly around the magnitude spectrum, linking the peaks [8]. In most signal

processing applications, the spectrum is derived from a short frame of the signal that was windowed and padded for optimal frequency resolution in a procedure known as a short-time Fourier transform (STFT). The power spectrum estimations performed on consecutive overlapping equally sized frames of the signal comprise a spectrogram. A spectrogram represents the change of sound quality over time but because of its high frequency resolution it may contain redundant information and noise. Estimating a spectral envelope for each STFT helps to smooth the spectrum and present a clearer picture of how the large structures in the signal change over time. Spectral envelope is also easier to represent parametrically which can be used for efficient signal compression and reliable feature extraction [8].

WORLD uses the CheapTrick algorithm for estimating the spectral envelope for each frame of the signal. A version of the algorithm is a part of the WORLD's C++ implementation and will be referred to as CCheapTrick in the below text. CCheapTrick works well for pre-recorded sound files, but it is not suited for real-time application because for short audio frames, it produces misaligned or noisy output that results in distortion in the reconstructed signal. Comparison of CCheapTrick's performance on real-time signal and pre-recorded audio file is shown on Figure 5.

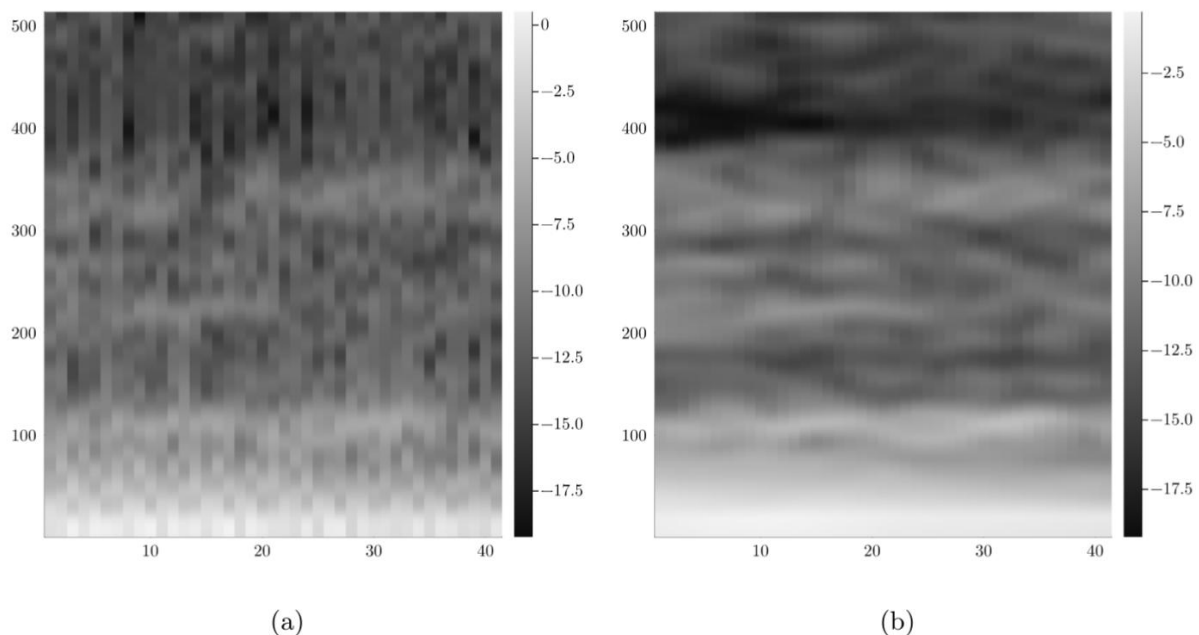


Figure 5 - CCheapTrick is used to estimate the spectral envelope for each frame of real-time audio stream (a) and for pre-recorded audio file (b) of the same speech fragment

A new solution, JCheapTrick, was implemented using the Julia programming language to enable analysis of real-time signals when used in conjunction with JDIO for F0 estimation. Below is the detailed description of the algorithm and motivation behind the choices made during investigation and implementation.

The CheapTrick algorithm is based on adaptive signal windowing that adjusts the window size based on the instantaneous frequency of the current frame. The window itself is a regular Hann window which is given by a squared cosine in range $[-0.5, 0.5]$. The window size is selected to be three times the period of the fundamental frequency $w_n = 3f_s F_0^{-1}$. For example, if a frame that is being analyzed has length of 800 samples at 16 kHz sampling rate and its instantaneous frequency is 220 Hz, the window size will be 218 which means only 218 first samples from the frame will be used for calculating its spectrum. This approach helps to ensure that the fundamental frequency and its harmonics will be the main contributors to the resulting spectrum since the instantaneous frequency is detected as close to the frame's beginning as possible as illustrated in Figure 3. The truncated frame with the Hann window applied to it is then padded with zeros to make its Fourier transform have sufficient resolution. The resolution is selected once and used for all frames and all window sizes to make it possible to combine all spectra into a spectrogram with consistent dimensions.

After the spectrum of a frame has been calculated, its magnitude is adjusted by the sum of the Hann window to accurately represent the energy of the frame. The adjusted spectrum then goes through a series of smoothing transformations.

First transformation is a simple FIR filter with a rectangular window. The length of the window is selected to be the index of the spectrum's energy bin which the estimated F0 falls into. So, a spectrum with lower associated instantaneous frequency will be smoothed with a smaller window than the one with a higher F0. In other words, the smoothing window size is selected so that it is not greater than the distance between the fundamental frequency's harmonics. The window is applied twice in forward and reverse directions to ensure that the result has zero phase distortions, i.e. the original spectrum does not shift in either direction after filtering. The added benefit of this step is that it helps to eliminate zero values from the spectrum which ensures that the cepstrum can be calculated without running into singularity issues.

The second transformation is application of the composite liftering function to the cepstrum of the signal to obtain the resulting spectral envelope.

The cepstrum is a "power spectrum of the power spectrum" that is calculated by applying a Fourier transform to the logarithm of the power spectrum as if it was a regular signal. Just like the spectrum helps to visualize and identify periodicities in time domain signals, cepstrum identifies periodicities in the spectrum itself. For example, a harmonic signal with a fundamental frequency of $F_0 = 220$ Hz will have a power spectrum with evenly spaced peaks that indicate F_0 and its harmonics; a cepstrum calculated for this power spectrum will have a peak at F_0^{-1} indicating the strong periodicity in the power spectrum. The F_0^{-1} value is referred to as quefrency, and it reflects a period of the corresponding frequency in the spectrum. Since the cepstrum is a regular Fourier transform, all techniques that can be applied to spectrum are also valid for cepstrum. For example, it is possible to apply a lowpass or highpass filters to the cepstrum to change the characteristics of the spectrum it was calculated for. Filters applied to the cepstrum are referred to as lifters and the process of applying them is called liftering.

The CheapTrick algorithm uses the following liftering functions:

$$l_s = \begin{cases} \frac{\sin(F_0 \pi \tau)}{F_0 \pi \tau} & \text{if } \tau > 0 \\ 1 & \text{if } \tau = 0 \end{cases} \quad (8)$$

$$l_q = q_0 + 2q_1 \cos(2\pi \tau F_0) \quad (9)$$

Here τ is the quefrency value, q_0 and q_1 are constants that were optimized to have values 1.18 and -0.9 respectively. The function (8) is effectively a lowpass filter for the spectrum that removes all components that correspond to F_0 and its harmonics, but it also removes the temporal instability in the spectrum caused by windowing the waveform and convolutions. The function (9) makes sure that the requirements of the consistent sampling theory are met, namely, that the digital signal must not change after being converted to analog signal and back. Since both of these functions are computationally expensive, the JCheapTrick implementation only computes their values up to the F_0^{-1} quefrency and sets all higher values to zero. So the final liftering function is given by the equation:

$$L(\tau) = \begin{cases} l_s(\tau)l_q(\tau) & \text{if } \tau < F_0^{-1} \\ 0 & \text{if } \tau \geq F_0^{-1} \end{cases} \quad (10)$$

Application of this lifter yields a temporally stable spectral envelope.

Evaluation. JDIO's and JCheapTrick's implementation is meant to provide an alternative to CDIO and CCheapTrick to enable real-time applications without compromising quality and accuracy, but the fact that CDIO and CCheapTrick are unable to work with short audio frames makes the direct comparison impossible.

Instead, the test audio files are processed as is by the C++ versions of the algorithms that produce an F0 contour and a spectrogram respectively while the Julia algorithms operate on the same files frame-by-frame as if they were real-time streams and the outputs are concatenated to enable side-by-side comparison of F0 contours and spectrograms.

For comparing the quality and accuracy of F0 estimation, a number of synthesized signals was generated. Below is the table that lists total mean square errors of CDIO and JDIO on different kinds of signals with length of 1 second and the sampling rate 16 kHz.

Table 1

Mean square deviations of the JDIO and CDIO estimations
from the true F0 of constant and modulated signals

Signal	JDIO	CDIO
Constant pitch signal	0.003	242.0
Amplitude modulated constant pitch signal	0.008	242.0
Linear chirp signal from 60 Hz to 600 Hz	7.8	20.4
Amplitude modulated linear chirp	7.8	20.5
Amplitude modulated linear chirp with harmonics	67.3	20.7
Noisy modulated linear chirp with harmonics	14716.1	25517.0

From the provided numbers it seems like the accuracy of CDIO is increasing as the signals become more complex until noise is introduced and the opposite is happening with JDIO. However, the actual distribution of errors throughout the signal shows that the CDIO's accuracy is consistent in all tests and the main contributors to high error values are the estimates at the beginning and at the end of the F0 contour. Such errors can be explained by the multistage postprocessing that CDIO uses to select valid F0 from all the candidates it estimated for each frame of the signal – candidates at the start and at the end of the signal lack the context for the postprocessing to work properly. Values of F0 that the CDIO estimated in the middle of the signal are highly accurate, but this accuracy slowly decreases towards higher frequencies. JDIO, on the other hand, doesn't employ any form of

postprocessing aside from the consecutive F0 difference minimization step described in (7) and has better consistency in its estimates. The graphical representation of error distributions is provided on Figure 6

F0 contours are easy to test because the ground truth for them is easy to obtain and compare the estimates against. The ground truth for the spectral envelopes, however, is very hard to obtain even for synthetic signals. So, to test the accuracy of JCheapTrick's implementation against CCheapTrick, an indirect approach to comparison was chosen:

The F0 contours obtained by JDIO and CDIO were used to obtain the spectrograms with JCheapTrick and CCheapTrick respectively.

The aperiodicity contours were estimated for both spectrograms.

The F0 contours, the spectrograms, and aperiodicities were used to resynthesize the analyzed waveform using WORLD's vocoder as per Figure 1.

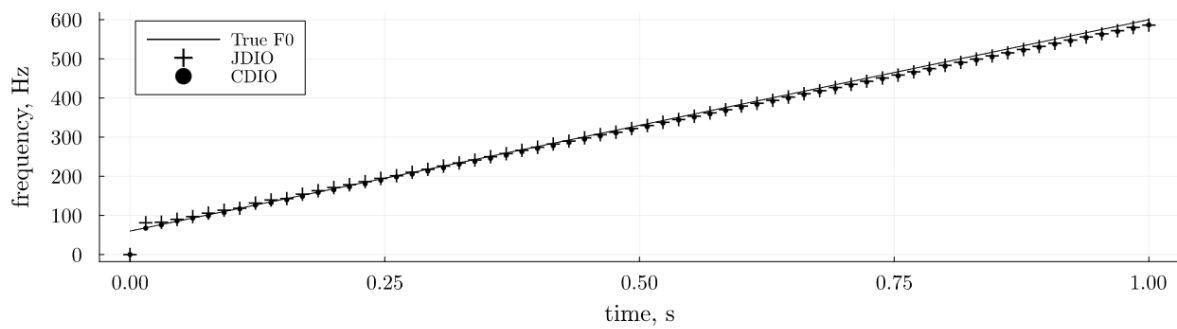
Raw error was obtained by calculating the mean square deviations of the synthesized waveforms from the original sound.

Spectral error was obtained by calculating the mean square deviations of the STFT spectrograms of the synthesized waveforms from the spectrogram of the original sound.

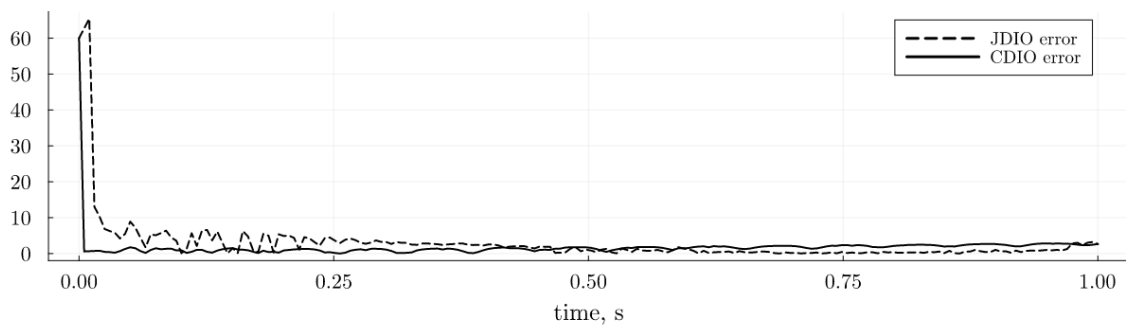
This method also helped to demonstrate that the JDIO and JCheapTrick algorithms were compatible with the WORLD framework implementation.

Instead of the synthesized signals, real speech recordings from the LibriSpeech corpus were used for the tests. The raw errors and spectral errors calculated using the above method are provided on Figure 7.

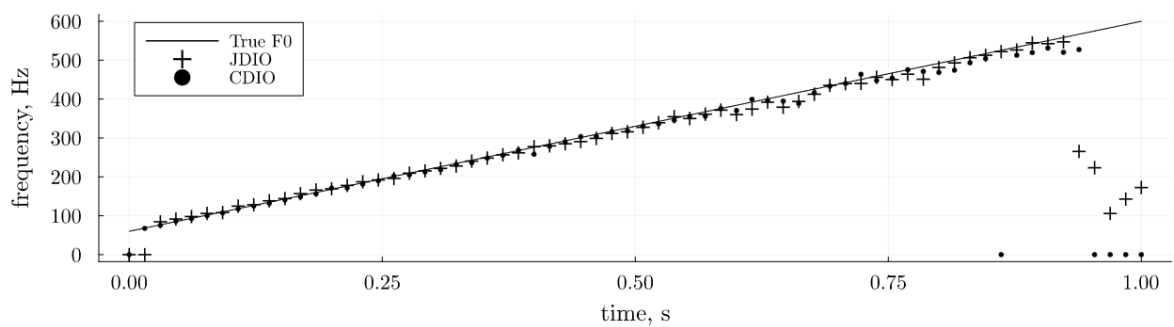
Amplitude modulated linear chirp with harmonics



Square error



Noisy modulated chirp with harmonics



Square error

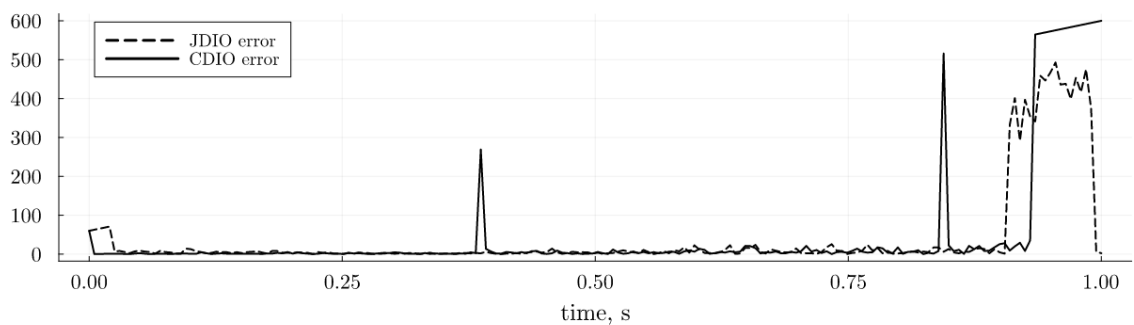


Figure 6 - Accuracy and errors for CDIO and JDIO when estimating the F0 contour for synthesized complex harmonic signals

The errors indicate that JDIO+JCheapTrick yielded more accurate resynthesized signal – the raw error was ~1.27 times better and the spectral

characteristics of the resynthesized sound was up to ~2.51 times better than for the signals resynthesized using CDIO+CCheapTrick.

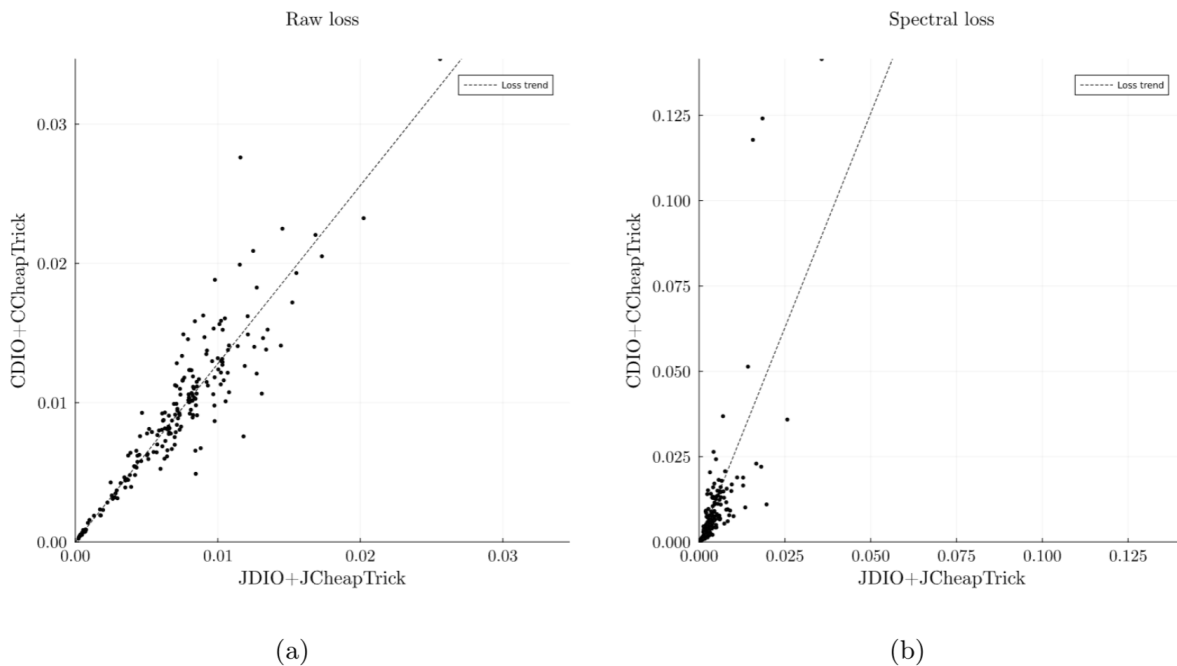


Figure 7 - Raw and spectral errors of resynthesized speech when compared to the original sound. The slopes of the mean lines are 1.27 for raw loss and 2.51 for spectral loss

Conclusions. In this paper, the implementation of the real-time DIO and CheapTrick algorithms was detailed and various accuracy and performance improvements were discussed. The new implementation was shown to have similar or better than the original implementations from the WORLD framework. The new algorithms can serve as a basis for further research and development of new real-time systems for processing speech streams.

REFERENCES

- Arik, S. O., Chrzanowski, M., Coates, A., Diamos, G., Gibiansky, A., Kang, Y., . . . Shoenybi, M. (2017, 06–11 Aug). Deep voice: Real-time neural text-to-speech. In D. Precup & Y. W. Teh (Eds.), Proceedings of the 34th international conference on machine learning (Vol. 70, pp. 195–204). PMLR. Retrieved from <https://proceedings.mlr.press/v70/arik17a.html>
- Daido, R., & Hisaminato, Y. (2016). A fast and accurate fundamental frequency estimator using recursive moving average filters. In Interspeech (pp. 2160–2164).

- Dave, N. (2013). Feature extraction methods lpc, plp and mfcc in speech recognition. International journal for advance research in engineering and technology, 1 (6), 1–4.
- De La Cuadra, P., Master, A. S., & Sapp, C. (2001). Efficient pitch detection techniques for interactive music. In Icmc.
- Hy, N. Q., Lee, S. W., Tian, X., Dong, M., & Chng, E. (2015). High quality voice conversion using prosodic and high-resolution spectral features. CoRR, abs/1512.01809 . Retrieved from <http://arxiv.org/abs/1512.01809>
- Jouviet, D., & Laprie, Y. (2017). Performance analysis of several pitch detection algorithms on simulated and real noisy speech data. In 2017 25th european signal processing conference (eusipco) (p. 1614-1618). doi: 10.23919/EUSIPCO.2017.8081482
- Kornblith, S., Lynch, G., Holters, M., Santos, J. F., Russell, S., Kickliter, J., . . . Smith, J. (2022, December). Juliadsp/dsp.jl: v0.7.8. Zenodo. Retrieved from <https://doi.org/10.5281/zenodo.7406426> doi: 10.5281/zenodo.7406426
- Morise, M. (2012). Platinum: A method to extract excitation signals for voice synthesis system. Acoustical Science and Technology, 33 (2), 123-125. doi: 10.1250/ast.33.123
- Morise, M. (2015). Cheaptrick, a spectral envelope estimator for high-quality speech synthesis. Speech Communication, 67 , 1-7. Retrieved from <https://www.sciencedirect.com/science/article/pii/S0167639314000697> doi: <https://doi.org/10.1016/j.specom.2014.09.003>
- Morise, M., Kawahara, H., & Katayose, H. (2009). Fast and reliable f0 estimation method based on the period extraction of vocal fold vibration of singing voice and speech.
- Morise, M., Yokomori, F., & Ozawa, K. (2016, 07). World: A vocoder-based high-quality speech synthesis system for real-time applications. IEICE Transactions on Information and Systems, E99.D, 1877-1884. doi: 10.1587/transinf.2015EDP7457
- Newmarch, J. (2017). Jack. In Linux sound programming (pp. 143–177). Berkeley, CA: Apress. Retrieved from https://doi.org/10.1007/978-1-4842-2496-0_7 doi: 10.1007/978-1-4842-2496-0_7
- Peng, K., Ping, W., Song, Z., & Zhao, K. (2020, 13–18 Jul). Non-autoregressive neural text-to-speech. In H. D. III & A. Singh (Eds.), Proceedings of the 37th international conference on machine learning (Vol. 119, pp. 7586–7598). PMLR. Retrieved from <https://proceedings.mlr.press/v119/peng20a.html>

- Ping, W., Peng, K., Gibiansky, A., Arik, S. O., Kannan, A., Narang, S., . . . Miller, J. (2017). Deep voice 3: 2000-speaker neural text-to-speech. CoRR, abs/1710.07654 . Retrieved from <http://arxiv.org/abs/1710.07654>
- Schwarz, D. (1998). Spectral Envelopes in Sound Analysis and Synthesis (Diplomarbeit Nr. 1622). Universität Stuttgart, Fakultät Informatik.
- Sun, L., Kang, S., Li, K., & Meng, H. (2015). Voice conversion using deep bidirectional long short-term memory based recurrent neural networks. In 2015 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP) (p. 4869-4873). doi: 10.1109/ICASSP.2015.7178896
- van den Oord, A., Dieleman, S., Zen, H., Simonyan, K., Vinyals, O., Graves, A., . . . Kavukcuoglu, K. (2016). Wavenet: A generative model for raw audio. CoRR, abs/1609.03499 . Retrieved from <http://arxiv.org/abs/1609.03499>

Received 17.10.2023.

Accepted 19.10.2023.

Адаптація фреймворку WORLD

для пофреймового аналізу мовлення в реальному часі

WORLD – це система для синтезу мовлення на основі вокодера, яка була розроблена М. Морісом та ін. і реалізована на C++. Було продемонстровано, дана система має високу ефективність та точність у порівнянні з аналогічними системами. Однак вона виявилася непридатною для використання у певних сценаріях, наприклад, при потоковій обробці аудіо фрейм-за-фреймом. Ця стаття розглядає недоліки C++ імплементації системи WORLD та пропонує модифіковані версії її складових алгоритмів для вирішення виявлених проблем. Результуючий фреймворк було протестовано на синтетичних сигналах та на реальних записах мовлення.

Eugene Koshel - postgraduate student, department of computer technologies, Oles Honchar Dnipro National University.

Кошель Євген - аспірант кафедри комп'ютерних технологій Дніпровського національного університету імені Олеся Гончара.

ВИКОРИСТАННЯ МЕТОДУ НЕЛІНІЙНОГО РЕКУРЕНТНОГО АНАЛІЗУ ДО ПОШУКУ DDOS АНОМАЛІЙ ЧАСОВИХ РЯДІВ МЕРЕЖЕВОГО ТРАФІКУ

Анотація. У статті розглядається питання використання методу нелінійного рекурентного аналізу до проблеми пошуку аномалій у мережевому трафіку, що надана у вигляді даних часових рядів знімку трафіку за певний період часу. Описано методика визначення якісної характеристики для вхідних даних та її використання для побудови відповідної рекурентної діаграми (recurrence plots, RP). В якості яких було запропоновано узяти статистику кількості отриманих байтів у секунду. Для отримання чисельних значень показників RP пропонується використовувати середовище Матлаб та розроблений для цього пакет crptools. Наведені обраховані показники RP дозволили здійснити типізацію отриманих даних та визначити тип якій отримав назву «DDOS RP», що надає змогу виділити деякі види атак типу DoS/DDoS.

Ключові слова: рекурентний аналіз, мережевий трафік, часовий ряд, рекурентна діаграма, параметр затримки, розмірність простору вкладення, RQA аналіз, середовище Матлаб.

Вступ і мета. Виявлення аномалій — це проблема пошуку шкідливих шаблонів у даних, які відрізняються від нормальної поведінки. У комп'ютерних мережах така ненормальна поведінка зазвичай виникає через такі фактори, як відмова в обслуговуванні (DoS), несправність мережевих пристроїв, неправильна конфігурація маршрутизатора, перевантаження трафіку, збій компонентів і вторгнення в мережу.

Відхилення мережевого трафіку від нормальної поведінки відбувається через: зміни шляху та петель маршрутизації, різкі зміни потоку трафіку, варіації затримок трафіку. Ці властивості визначаються та класифікуються шляхом спостереження за потоком трафіку, затримкою пакетів, загальною кількістю надісланих пакетів, співвідношенням байтів, надісланих у кожному напрямку, тривалістю з'єднання, середнім розміром пакетів і часом між надходженнями. Поточне виявлення аномалії спостерігає за статистикою транспортного рівня, і результати порівнюються з пороговим значенням. Проте спостерігається, що характеристики мережевого трафіку, характеристики розподілу ймовірностей

IP-пакетів, щільність трафіку, статистичні характеристики трафіку динамічно змінюються у часовій області.

Нещодавно було запропоновано повторний нелінійний рекурентний аналіз (RQA) для спостереження та вивчення цих нестаціонарних властивостей [1]. З часом, з'явилися інші роботи, які виконували аналіз на нестандартизованих даних, та надавали різну до знаходження нових підходів до пошуку аномалії [2,3]. Деякі роботи використовують у якості вхідних даних згорткової нейронної мережі [4]. Проблема досі є досить складною і актуальною.

Мета даної роботи – запропонувати підхід реалізації пошуку аномалій для атаки сервісу типу DDoS/DoS даних мережевого трафіку за допомогою чисельних показників нелінійного рекурентного аналізу.

Огляд проблеми та її аналіз. Основна проблема аналізу часового ряду полягає в тому, що для побудови рекурентної діаграми (RP) необхідно чіткого визначення прихованих параметрів m і τ . m - це розмірність простору вкладення часового ряду, для якого виконується аналіз, а τ - параметр його затримки. До визначення даних параметрів деякі дослідники [2-4] поставилися суто формально, що є помилковим з точки зору синергетики, а тим більше нелінійної динаміки.

У дослідженні [6] представили хороше опитування щодо виявлення аномалій. Опитування щодо методів виявлення аномалій у сферах машинного навчання та статистики можна знайти в [7]. Для виявлення аномалій мережі доступно декілька інструментів. Деякі з цих інструментів використовують чітко визначені правила для аналізу моделей трафіку. Кілька інших методів включають експоненціальне згладжування та прогноз Холта Вінтерса, адаптивну порогову кумулятивну суму, оцінку максимальної ентропії.

Багато з цих детекторів працюють на основі статистики транспортного ривня, включаючи співвідношення байтів, надісланих у кожному напрямку, середній розмір і середній час між надходженнями пакетів. Далі результати порівнюються з пороговими значеннями. Ці значення не залежать від використання мережі та кількості користувачів. Мережеву аномалію можна виявити за допомогою інших моделей, включаючи байєсівську, опорну векторну машину, нейронні мережі тощо. У літературі для вирішення проблеми класифікації мережевого трафіку застосовано кілька методів аналізу даних. [8] запропонував виявлення аномалій мережевого трафіку на основі байтів пакетів. У даній роботі ми запропонуванували інші види візуалізації отриманої ними інформації: рекуре-

нті діаграми. Тому має сенс спрямувати зусилля на більш детальне дослідження параметрів m і τ , а також чисельних показників RP.

У роботі використовувалося програмне забезпечення, розроблене і реалізоване у середовищі типу Матлаб (Октавіа) – `crptool` [9] під ОС Windows. Для дослідження часових рядів було використано набір даних оцінки методології пошуку аномалій CIC-IDS2017 [5].

Опис даних мережевого трафіку. CIC-IDS2017 набір даних було захоплено з 3 червня 2017 року по 7 червня 2017. Впродовж 5 днів у лабораторних умовах було відтворено ряд базових втручань у внутрішню систему зокрема: веб-атаки, проникнення, ботнет і DDoS/DOS.

У експерименті приймали участь 25 комп'ютерів які симулювали стандартну роботу офісну (пересилання документів, поштові операції, пошук інформації у Google, тощо) і два сервери які працювали як стандартні файлообмінники. На сервери середовища було здійснено атаку DoS/DDoS знято статистику у наступних часових проміжках:

- 05.07.2017 з 9:47 по 10:10 було здійснено атаку DoS slowloris з адреси 205.174.165.73 на сервер 192.168.10.50
- 05.07.2017 з 10:14 по 10:35 було здійснено атаку DoS Slowhttpstest з адреси 205.174.165.73 на сервер 192.168.10.50
- 05.07.2017 з 10:43 по 11:00 було здійснено атаку DoS Hulk з адреси 205.174.165.73 на сервер 192.168.10.50
- 05.07.2017 з 11:10 по 11:23 було здійснено атаку DoS GoldenEye з адреси 205.174.165.73 на сервер 192.168.10.50
- 07.07.2017 з 15:56 по 16:16 було здійснено атаку DDoS LOIT з адрес 205.174.165.69, 205.174.165.70 205.174.165.71 на сервер 192.168.10.50

Окрім даних, які містять аномально активність, було також зібрано 5 проміжків по 10 хвилин з даних зібраних у понеділок 03.07.2017, які містять схожу за динамікою активність, але не містять втручання. Така динаміка може бути викликана на початку робочого дня, під час роботи програм моніторингу активності або масштабної передачі даних.

Інформація була захоплена за допомогою програмного забезпечення Wireshark. За допомогою даної програми також було зібрано статистику по кількості пакетів отриманих сервером з інтервалом у одну секунду.

Проблеми пошуку параметрів. Задача пошуку аномалій у часових рядах полягає у виявленні незвичайних та несподіваних змін у поведінці системи, яка описується часовим рядом. Аномалії можуть вказувати на виникнення

проблем або надзвичайних подій, що потребують уваги та відповідного реагування. У випадку з відхиленнями у даних мережевого трафіку технічний спеціаліст може зупинити атаку шляхом фільтрації запитів з певних адрес, або проводити системні відключення від системи заражених користувачів з подальшим усуненням проблеми.

Одним з підходів до пошуку аномалій є порівняння між очікуваними значеннями та фактичними значеннями часового ряду. Якщо спостерігається відхилення, то це може вказувати на наявність аномалії. Інші підходи включають виявлення змін точок перелому та застосування різноманітних методів машинного навчання, таких як ансамблеві методи, навчання з учителем та навчання без учителя[].

Нажаль, зі збільшенням систем пошук аномалій може займати все більше інформаційного ресурсу. Моделі засновані на принципах математичної статистики в реальних умовах обмежені у часових, обчислювальних і матеріальних ресурсів. Отже необхідні методології з використанням технологій інтелектуального аналізу даних.

Отже, маємо наступний ряд задач: визначити зі знятих реальних даних мережевого трафіку відхилення у числових показниках, які вказуватимуть на наявність втручання та аномальних процесів у системі; проаналізувати дані показники для того, щоб відокремити для пошуку атаки класу DoS/DDoS.

У якості математичного апарату було обрано нелінійний рекурентний аналіз. Сучасні методології пошуку аномалій будуються на нейронних диференціальних рівняннях (NeuralODE)[].

Застосування RQA аналізу може бути використаним є цікавим і перспективним саме для таких систем. Запропонована наступна методика обробки реальної інформації мережевих потоків по кожному з інцидентів втручання, яка складалася з наступних кроків:

1. Надання графіку статистики кількості отриманих даних сервером за одну секунду (*Rate*) у нормальних та аномальних .
2. Обрахування значень параметра затримки τ .
3. Визначення значення розмірності простору вкладення m (для $\min$ та $\max$ -норми) з урахуванням вже визначених у пункті 2 значень параметра затримки τ .
4. Побудова RP діаграми для визначених параметрів τ та m . Значення ϵ обирається не більше 20%.

5. Розрахунок для визначених значень τ та m чисельних показників RP діаграми - RQA аналіз.

6. Отримати усі значення чисельних показників RP діаграми, визначити якісні характеристики за якими можна виявляти втручання та аномальні процеси типу атак DoS/DDoS.

Визначення чисельних показників RP діаграми. Графік рекурентності - це квадратна матриця, яка є симетричною відносно головної діагоналі та містить час виникнення станів як у стовпцях, так і в рядках. Кожен елемент матриці відповідає певній парі часів та містить 1, якщо стан повторюється, та 0, якщо ні.

Іншим визначенням графіку повторення є $N \times N$ матриця, що складається лише з чорних та білих точок з позначенням, що чорна точка відображає повторення, разом з двома осями часу. Математично, графік повторності можна виразити наступним рівнянням.:

$$R_{i,j} = \theta(\varepsilon_i - \|\vec{x}_i - \vec{x}_j\|), \quad \vec{x}_j \in \mathbb{R}^m \quad i, j = 1, \dots, N$$

де розглядаються стани $\mathbf{X}_i$ мають кількість N ; ε_i - порогове значення відстані (околиця); $\|\cdot\|$ - норма; $\theta(\cdot)$ - функція Хевісайду. Коли відстань між двома станами, тобто $\mathbf{X}_i$ та $\mathbf{X}_j$, менша за порогове значення ε , визначається повторення.

Для показників, які використовуємо, представимо їх визначення, що на дано у рекурентному аналізі. Для обчислення характеристик, які використовуються у RP, значення порога ε є фіксованими.

Міра рекурентності (recurrence rate, RR)

$$RR = \frac{1}{N^2} \sum_{i,j=1}^N R_{i,j}^{m,\varepsilon}$$

показує щільність рекурентних точок, просто підраховуючи їх. N - кількість точок на траєкторії у фазовому просторі. Дана міра показує можливість знаходження рекурентної точки в RP (ймовірність повторення стану).

Наступна міра розглядає діагональні лінії. Частотний розподіл довжин l діагональних ліній в RP $P^\varepsilon(l) = \{l_i; i = 1, \dots, N_l\}$, де N_l - абсолютна кількість діагональних ліній (кожна лінія надається тільки один раз). Процеси зі стохастичною поведінкою можуть породжувати дуже короткі діагоналі чи взагалі не породжувати їх, тоді як детерміністські процеси дають довгі діагоналі і малу кількість окремих рекурентних точок.

Отже, відношення рекурентних точок

$$DET = \frac{\sum_{l=l_{min}}^N l P^{\varepsilon}(l)}{\sum_{i,j=1}^N R_{i,j}^{m,\varepsilon}},$$

називається *мірою детермінізму (determinism, DET)* або передбачуваності системи. Слід зазначити, що цей захід не має значення реального детермінізму процесу. Порогове значення l_{min} виключає діагональні лінії, утворені тангенціальним рухом траєкторії у фазовому просторі. Якщо $l_{min} = 1$, $DET=RR$. Більше значення детермінізму вказує більш діагональну лінію в RP і, отже, більш сильну передбачуваність системи.

Діагональні структури показують час, протягом якого ділянка траєкторії підходить досить близько до іншої ділянки траєкторії. Таким чином, ці лінії дозволяють судити про розбіжність елементів траєкторій.

Середня довжина діагональних ліній

$$L = \frac{\sum_{l=l_{min}}^N l P^{\varepsilon}(l)}{\sum_{l=l_{min}}^N P^{\varepsilon}(l)},$$

це середній час, протягом якого дві ділянки траєкторії проходять близько одна до одної, і може розглядатися як середній час передбачуваності. Чим більше значення L, тим менша випадковість, тобто легше визначити поведінку ознаки системи.

Міра ентропії (entropy, ENTR) співвідноситься з ентропією Шеннона (Shannon) частотного розподілу довжин діагональних ліній

$$ENTR = - \sum_{l=l_{min}}^N p(l) \ln p(l), \quad p(l) = \frac{P^{\varepsilon}(l)}{\sum_{l=l_{min}}^N P^{\varepsilon}(l)},$$

та відображає складність детерміністської складової у системі. Велике значення ентропії має на увазі періодичність системи, а низьке - хаотичність. Інакше кажучи, велика ентропія слідує за складнішою системою.

Міра завмирання (laminarity, LAM)

$$LAM = \frac{\sum_{v=v_{min}}^N v P^{\varepsilon}(v)}{\sum_{i,j=1}^N R_{i,j}^{m,\varepsilon}},$$

визначається відношенням кількості рекурентних точок, що утворюють вертикальні лінії, до загальної кількості рекурентних точок. LAM характеризує наявність станів завмирання системи (тобто коли рух системи фазової траєкторії зупиняється або просувається дуже повільно). Ламінарність обчислює ймовірність того, що стан залишиться для наступного тимчасового кроку.

Середня довжина вертикальних структур

$$TT = \frac{\sum_{v=v_{min}}^N v P^{\varepsilon}(v)}{\sum_{i,j=1}^N P^{\varepsilon}(v)},$$

називається мірою часу зупинки (*trapping time*, TT) і характеризує середній час, який система може провести у певному стані або тривалість часу, протягом якого кожен стан перебуває у пастці.

Проведемо пошаговий розрахунок для нормального та аномального часового ряду. Як було зазначено в методології розрахунки велися для міри отриманої кількості даних у байтах у 1 секунду. Назвемо данні зібранні з даних процесів *Rate1* та *Rate2*.

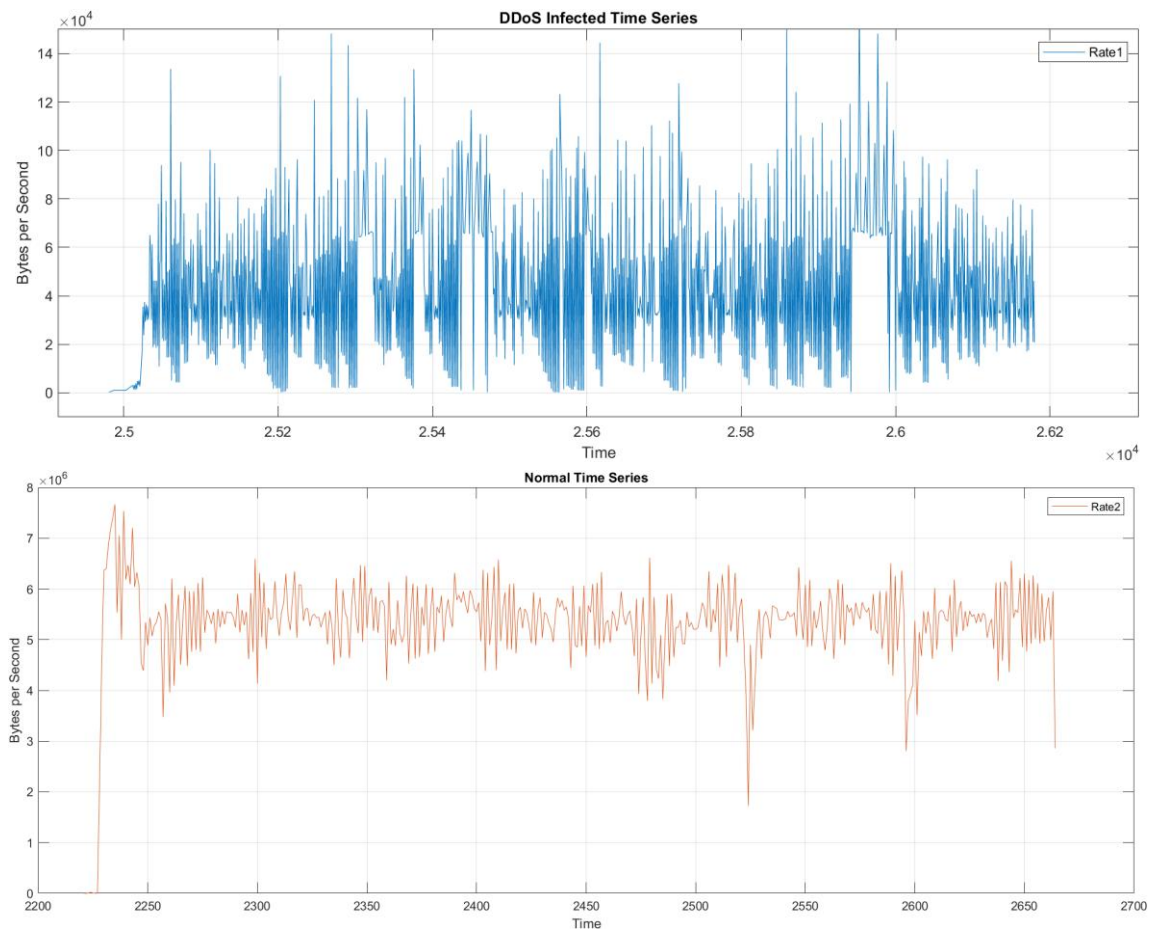


Рисунок 1 - Траєкторія поведінки сигналів у т. *Rate1*(синя), *Rate2*(червона)

1) Розрахунки для точки *Rate1*.

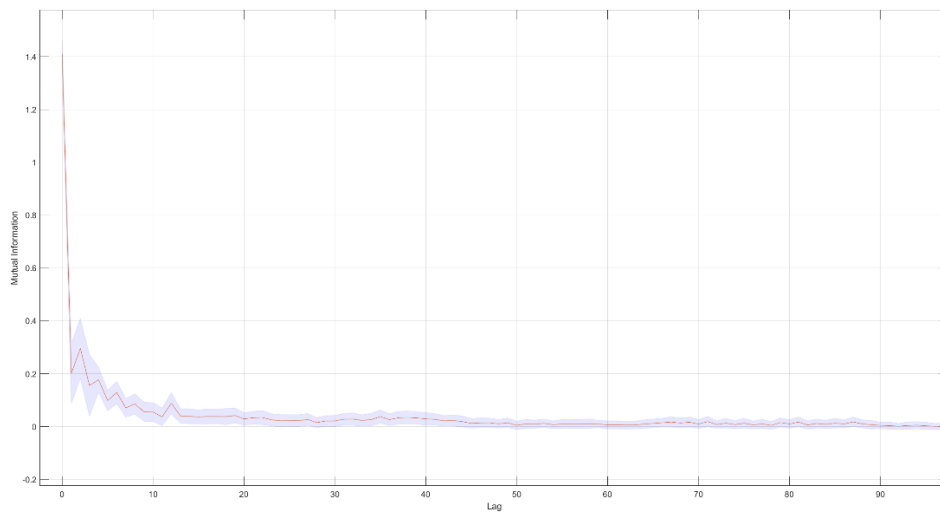


Рисунок 2 - Визначення параметра затримки τ для точки *Rate1*
(значення 28, 50, 78)

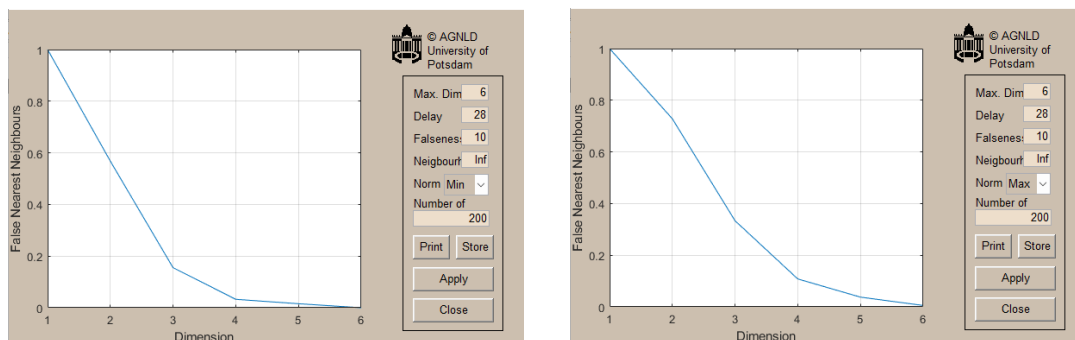


Рисунок 3 - Обрахування m для точки *Rate1*, $\tau=28$ (перше значення ціле нижче 0.1) $m=4$ зліва та $m=5$ справа (надано для обох норм)

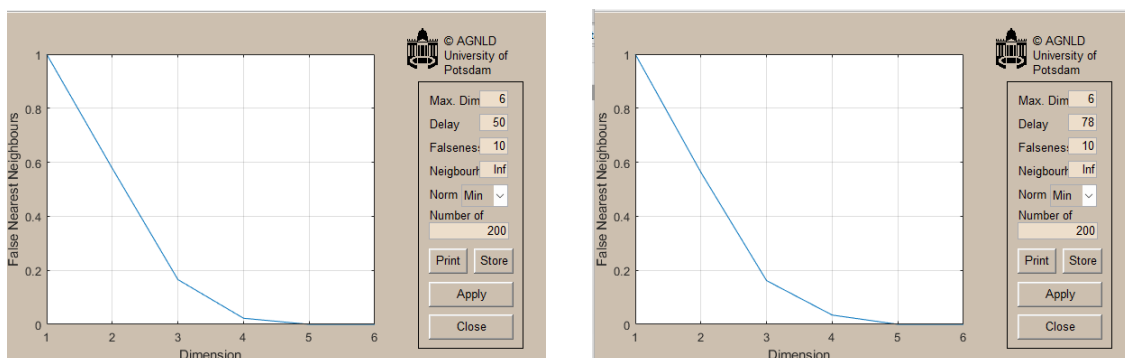


Рисунок 4 - Обрахування m для точки *Rate1*, $\tau=50$ (зліва), 78 (справа)
 $m=4$ зліва та $m=4$ справа (надано для норми min)

Результати щодо норми max такі, як на рис.3 для норми min $m=5$.

Для розрахунку рекурентної діаграми було викроєстано інший інструмент MatLab:

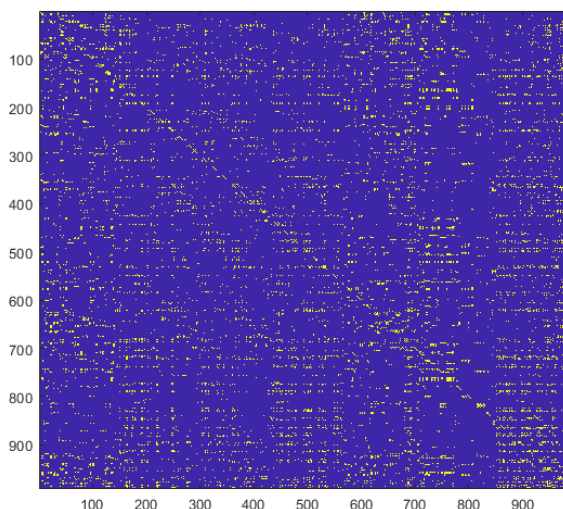


Рисунок 5 - Рекурентна діаграма сигналу O1 ($m=5$, $\tau=28$, $\varepsilon=0.2$)

Програма видала наступні значення показників RQA:

Таблиця 1

DDoS LoIC	RR	DET	ENTR	L	LAM	TT
<i>Rate1</i>	0.046105	0.24954	0.06717	2.2379	0.019056	2.29866

2) Для точки *Rate2* значення параметра затримки τ аналогічно рис.2. було визначено та мало наступні значення: $\tau=21, 43, 87$.

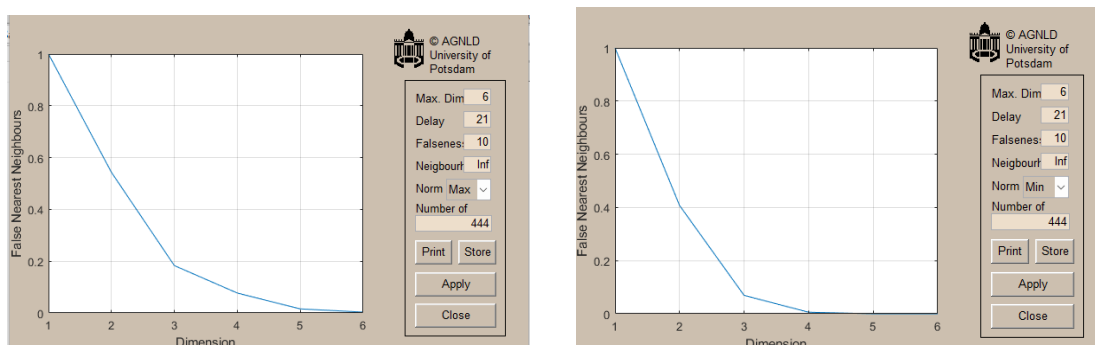


Рисунок 6 - Обрахування m для точки *Rate2*, $\tau=21$ (перше значення ціле нижче 0.1) $m=5$ зліва та $m=4$ справа (надано для обох норм)

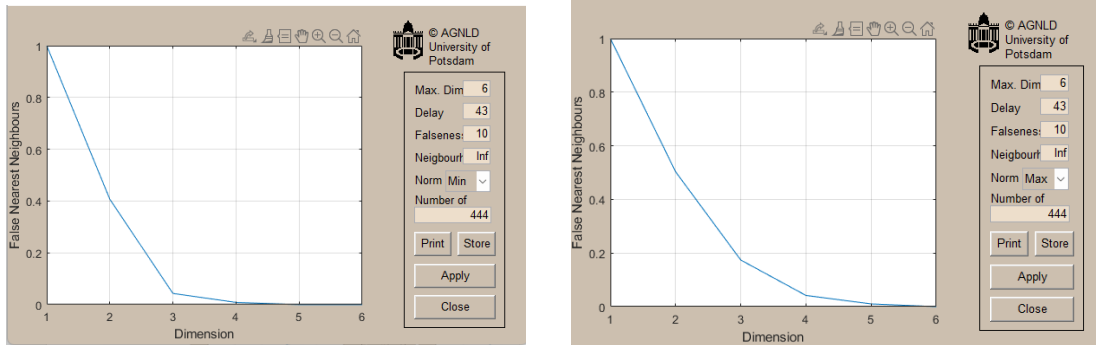


Рисунок 7 - Обрахування m для точки *Rate2*, $\tau = 43$ (перше значення ціле нижче 0.1) $m=4$ зліва та $m=5$ справа (надано для обох норм)

Результати щодо точки O2 для $\tau = 87$, $m=5$ для max норми та $m=4$ для min норми. Такім чином вони аналогічні для усіх значень параметру затримки.

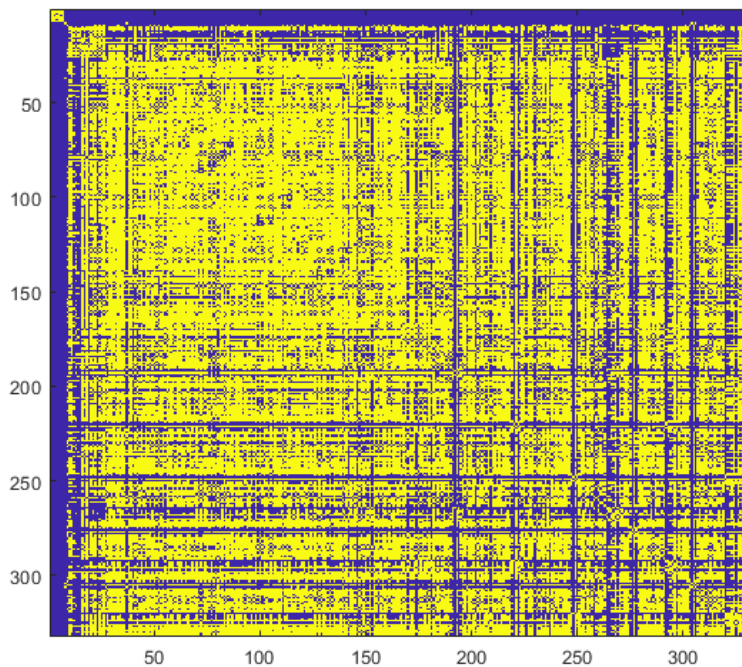


Рисунок 8 - Рекурентна діаграма RP сигналу *Rate2* ($m=5$, $\tau=21$, $\epsilon=0.2$)

Якщо дивитися на рис. 5 та рис. 8, то можна помітити лише деяку незначну структуру схожості. Але зробити якісь висновки досить складно адже процеси відбуваються з різною інтенсивністю.

Таблиця 2

DDoS LoIC	RR	DET	ENTR	L	LAM	TT
<i>Rate2</i>	0.62454	0.87469	0.33398	4.52934	0.86912	5.64760

Після першого аналізу отриманої інформації можна зробити висновок, що використання RP діаграми для будь-якого візуального аналізу є досить важким завданням. Наявними є також деякі проблеми з визначення значень параметру затримки з представленого графіку (Рис.2).

Використовуючи надану методику, розраховуємо чиселні показники для усіх нормальних та аномальних даних мережевого трафіку. Усі данні будуть зведені у єдину таблицю групувану по маркуванню. Результати обробки мережевих даних, що описані у першому розділі, представлено у підсумковій таблиці 3.

Таблиця 3

Показники RQA у аномальних процесах						
	RR	DET	ENTR	L	LAM	TT
DDoS LoIC	0.04611	0.24954	0.06717	2.23791	0.019056	2.29866
DoS slowloris	0.00134	0,38128	0.05415	2.12443	0.991322	2.10536
DoS Slowhttptest	0.00913	0.56233	0.00124	12.1245	0.017162	6.41234
DoS Hulk	0.04694	0.27281	0.04524	2.40560	0.426120	2.75914
DoS GoldenEye	0.04797	0.29281	0.06428	3.81642	0.049356	4.1282
Показники RQA у нормальних процесах						
	RR	DET	ENTR	L	LAM	TT
Normal1	0.624542	0.87469	0.33398	4.52934	0.86912	5.64760
Normal2	0.00412	0.00831	1.08273	85.12433	0.12543	47.12932
Normal3	0.22784	0.51283	1.69012	25.09128	0.71293	10.09120
Normal4	0	0.95190	4.68395	112.10353	0.99153	215.12093
Normal5	0.00084	0.72612	0.91243	2.02089	0.02132	2.33333

Безпосередній порівняльний аналіз значень, який знайшов відображення у таблиці 3, дозволив виявити деякі нечіткі межі щодо ряду обрахованих чисельних показників. Ці показники дозволяють виконати виділити аномальну поведінку під час атаки DDoS у переважній більшості випадків. У подальшому дані числові характеристики можуть використовуватися для уточнення методології пошуку аномалій у якості додаткової фільтрації.

Під час атаки DoS/DDoS данні не є однорідними. Вдалось виділити один стандартний тип, який має суттєву різницю від нормальної поведінки системи за параметрами: міра рекурентності, міра детермінізму та міра ентропії. Ці характеристики є наступними:

- показники RR мають усі значення біля нуля [0.044-0.048]
- показники DET, ENTR значення різні, але з одного невеликого діапазону (наприклад, [0.27-0.3] або [0.056-0.071]);

У результаті обчислювального експерименту, який полягав у проведенні процедури – розрахунків чисельних показників рекурентних діаграм (RP)– авторам вдалося виділити один тип – «DDoS-RP», якій відрізняє аномальний процес від нормального та визначає саме атаки типу DoS/DDoS.

У наступному пропонується розглянути задачу налаштування та навчання нейронної мережі з використанням числових характеристик RQA. З розвитком Neural-ODE починають розвиватись методології додаткового налаштування нейронних мереж заснованих на методах аналізу нелінійних систем. Також запропоновано розглянути інші види атак зокрема атаки типу botnet та різні види втручання.

Висновки. 1. У роботі наведено методику розрахунків чисельних показників рекурентної діаграми (RP) та визначення нейрофізіологічного значення її параметрів — міра рекурентності (RR ймовірність повторення стану), міра детермінізму (DET більше значення детермінізму вказує більш сильну передбачуваність системи), середня довжина діагональних ліній(чим більше значення L, тим менша випадковість, тобто легше визначити поведінку ознаки системи), міра ентропії (ENTR відображає складність детерміністської складової у системі. Велике значення ентропії має на увазі періодичність системи, а низьке - хаотичність. Інакше кажучи, велика ентропія слідує за складнішою системою), міра завмирання (LAM ламінарність обчислює ймовірність того, що стан залишиться для наступного тимчасового кроку), міра часу зупинки (TT характеризує середній час, який система може провести у певному стані або

тривалість часу, протягом якого кожен стан перебуває у пастці) для аномальних даних мережевого трафіку і нормальних.

2. Аналіз отриманої інформації дозволив визначити деякі види аномальної активності типу DoS/DDoS, які мають наступні ознаки:

- показники RR мають усі значення біля нуля [0.044-0.048]
- показники DET, ENTR значення різні, але з одного невеликого діапазону (наприклад, [0.27-0.3] або [0.056-0.071]);

3. В результаті обчислювального експерименту, вдалося виділити один тип рекурентних діаграм – «DDoS-RP», якій відрізняє аномальний процес від нормального та визначає саме атаки типу DoS/DDoS.

ЛІТЕРАТУРА / LITERATURE

1. Palmeieri, F. & Fiore, U. Network anomaly detection through non-linear analysis. Computers Security, 2010, 29(7), 737-55.
2. Somenath Mukherjeea, Rajdeep Ray, Rajkumar Samantac, Mofazzal H. Khondekar ,Goutam Sanyal Nonlinearity and chaos in wireless network traffic Chaos, Solitons and Fractals 96 (2017) 23–29
3. N. Jeyanthi, R. Thandeeswaran, J. Vinithra RQA Based Approach to Detect and Prevent DDoS Attacks in VoIP Networks CYBERNETICS AND INFORMATION TECHNOLOGIE Volume 14, No 1 11-24 DOI: 10.2478/cait-2014-0002
4. Yun Chen Shijie Sum, Hui Yang Convolutional Neural Network Analysis of Recurrence Plots for Anomaly Detection International Journal of Bifurcation and Chaos, Vol. 30, No. 1 (2020) 2050002 (13 pages)
5. Iman Sharafaldin, Arash Habibi Lashkari, and Ali A. Ghorbani, “Toward Generating a New Intrusion Detection Dataset and Intrusion Traffic Characterization”, 4th International Conference on Information Systems Security and Privacy (ICISSP), Portugal, January 2018
6. Chandola, V.; Banerjee, A. & Kumar, V. Anomaly detection: A survey. ACM Computing Surveys, 2009, 41(3), 1-58.
7. Hodge, V. & Austin, J. A survey of outlier detection methodologies. Art. Intel. Revi., 2004, 22(2), 85-126
8. Mahoney, M. Network traffic anomaly detection based on packet bytes. In Proceedings of ACM Symposium on Applied Computing, 2003, pp. 346-50.
9. Mekler A.A. Application of the Apparatus for Nonlinear Analysis of Dynamic Systems for EEG Signal Processing // Actual Problems of Modern Mathematics: Scientific Notes. p / ed. prof. Kalashnikova E.V., ed. LGU them. A.S. Pushkin, St. Petersburg, 2004. T. 13 (issue 2), p. 112-140.

REFERENCES

- Palmeieri, F. & Fiore, U. Network anomaly detection through non-linear analysis. Computers Security, 2010, 29(7), 737-55.
- Somenath Mukherjeea, Rajdeep Ray, Rajkumar Samantac, Mofazzal H. Khondekar ,Goutam Sanyal Nonlinearity and chaos in wireless network traffic Chaos, Solitons and Fractals 96 (2017) 23–29
- N. Jeyanthi, R. Thandeeswaran, J. Vinithra RQA Based Approach to Detect and Prevent DDoS Attacks in VoIP Networks CYBERNETICS AND INFORMATION TECHNOLOGIE Volume 14, No 1 11-24 DOI: 10.2478/cait-2014-0002
- Yun Chen Shijie Sum, Hui Yang Convolutional Neural Network Analysis of Recurrence Plots for Anomaly Detection International Journal of Bifurcation and Chaos, Vol. 30, No. 1 (2020) 2050002 (13 pages)
- Iman Sharafaldin, Arash Habibi Lashkari, and Ali A. Ghorbani, “Toward Generating a New Intrusion Detection Dataset and Intrusion Traffic Characterization”, 4th International Conference on Information Systems Security and Privacy (ICISSP), Portugal, January 2018
- Chandola, V.; Banerjee, A. & Kumar, V. Anomaly detection: A survey. ACM Computing Surveys, 2009,41(3), 1-58.
- Hodge, V. & Austin, J. A survey of outlier detection methodologies. Art. Intel. Revi., 2004, 22(2), 85-126
- Mahoney, M. Network traffic anomaly detection based on packet bytes. In Proceedings of ACM Symposium on Applied Computing, 2003, pp. 346-50.
- Mekler A.A. Application of the Apparatus for Nonlinear Analysis of Dynamic Systems for EEG Signal Processing // Actual Problems of Modern Mathematics: Scientific Notes. p / ed. prof. Kalashnikova E.V., ed. LGU them. A.S. Pushkin, St. Petersburg, 2004. T. 13 (issue 2), p. 112-140.

Received 17.10.2023.

Accepted 21.10.2023.

Using the method of nonlinear recursive analysis for detecting DDoS anomalies in time series data

This research endeavors to address this gap by determining a qualitative characteristic for server network traffic and use it to construct the corresponding recurrence plot (RP). The goal of this study is to develop and assess a novel technique based on nonlinear recursive analysis to detect Distributed Denial of Service (DDoS) anomalies in network traffic time series data. With the increasing frequency of DDoS attacks on modern digital

infrastructures, there is a pressing need for more efficient and accurate detection methods.

There has been some attempts to apply nonlinear analysis to network traffic [2-4], but those studies lack critical steps in determining parameters for embedding space dimension m and delay time τ . More recent studies have explored machine learning and deep learning approaches [7], which offer improved accuracy but can be computationally intensive and require extensive training data. Despite the advancements, there remains a need for a method that is both accurate and efficient, especially in real-time detection scenarios.

The researchers employed nonlinear recursive analysis by estimating RQA parameters and determining a qualitative characteristic of data points of DDoS attack contained in CIC-IDS2017 dataset. A technique for determining hidden information for this series and its use for constructing the corresponding recurrence diagram (RP) at the points of information retrieval are described. It is shown that the use of RP has significant drawbacks associated with the visualization of information on a computer monitor screen, so another way of research is proposed - the calculation of numerical indicators of RP

The given calculated RP indicators made it possible to typify the received data and determine the type, which was named "DDOS-RP", which makes it possible to distinguish some types of DoS/DDoS type attacks. The study concludes by recommending further exploration of this method in diverse network environments and against more complex DDoS attack patterns.

Key words: recurrent analysis, network traffic, time series, recurrent diagram, delay parameter, dimension of the nesting space, RQA analysis, Matlab environment

Гулий Тарас Олександрович – аспірант, кафедри комп'ютерних технологій Дніпровського національного університету імені Олеся Гончара.

Білозьоров Василій Євгенійович – доктор фізико-математичних наук, професор кафедри комп'ютерних технологій Дніпровського національного університету імені Олеся Гончара.

Hulyi Taras - postgraduate student, department of computer technologies, Oles Honchar Dnipro National University.

Belozyorov Vasily - Doctor of Physical and Mathematical Sciences, Professor, Department of computer Technologies, Oles Honchar Dnipro National University.

REAL-TIME DATA VISUALIZATION FOR IOT NETWORK SYSTEMS: CHALLENGES AND STRATEGIES FOR PERFORMANCE OPTIMIZATION

Abstract. Real time data visualization has become an essential tool for decision making systems in various industries, including finance, healthcare, IoT, and manufacturing. Real time data visualization enables organizations to monitor and analyze data as it is generated, providing real time insights into critical business operations. However, real time data visualization poses several challenges, including performance, data quality, and visualization complexity. This paper will explore the importance of real time data visualization in IoT network systems, and the challenges associated with it. Specifically, the paper will discuss the challenges of real time data visualization and ideas to increase performance. The paper will also provide a comprehensive analysis of the impact of real time data visualization on IoT network and decision making systems, highlighting its benefits and potential drawbacks. The paper will begin by discussing the importance of real time data visualization in IoT network systems, highlighting its role in providing timely insights into critical operations. It will then delve into the challenges associated with real time data visualization, including data quality, visualization complexity, and performance. The paper will provide a detailed analysis of each challenge, outlining the potential impact on real time data visualization systems and decision making processes. The paper will also explore ideas to increase performance in real time data visualization, including implementing high performance computing infrastructure, optimizing data processing and analysis, using caching techniques, using visualization techniques optimized for performance, implementing data compression, and using real time analytics. The paper will provide a comprehensive analysis of each idea, outlining its potential impact on real time data visualization systems' performance and overall effectiveness. Finally, the paper will conclude by highlighting the importance of real time data visualization in IoT network systems and the need to address the challenges associated with it. The paper will also provide recommendations about how to implement real time data visualization systems, outlining key considerations and best practices to ensure successful implementation and optimal performance.

Keywords: real time data visualization, IoT, decision making, performance optimization.

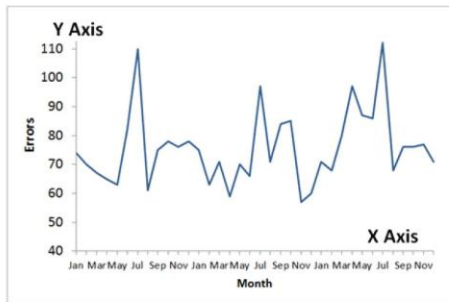
Problem Statement. Data visualization has become an essential tool for IoT network systems. With the advent of real-time data, the ability to visualize this data in real-time has become increasingly important. Real-time data visualization pro-

vides decision-makers with the ability to quickly analyze and interpret data as it is generated, allowing them to make timely and informed decisions. The amount of data generated by modern IoT network systems is constantly increasing, and it is becoming increasingly challenging to process and make sense of this data in a timely manner [1]. Real-time data visualization provides a solution to this problem by allowing decision-makers to identify trends and patterns quickly and easily in large datasets. Visual analytics “combines automated data analysis with interactive visualizations for an effective understanding, reasoning, and decision making on the basis of very large and complex datasets” [2]. However, real-time data visualization also presents several challenges, including performance issues that can slow down the system and affect the accuracy of the visualizations.

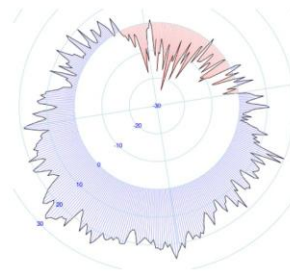
Analysis of Recent Research and Publications. Real-time data visualization is a process of generating and displaying data in real-time. It provides decision-makers with the ability to view data as it is generated, allowing them to make informed decisions based on current information. Real-time data visualization is particularly useful for IoT network systems that require timely responses, such as financial trading, transportation management, and emergency response systems. Real-time data visualization can be achieved using a variety of techniques, including dashboards, charts, graphs, and maps. Recent research in visual analytics specify different methods of visual representation and performance optimizations [3-8]. Various chart types are shown on Fig. 1.

Visualization techniques may be described in three dimensions: layout, scale, and representation. It is suitable for storytelling chart classification [5].

These visualizations are typically generated using data from real-time data sources, such as sensors, databases, and social media feeds. Real-time data visualization tools often provide users with the ability to filter and drill down into data sets, allowing them to quickly identify trends and patterns.



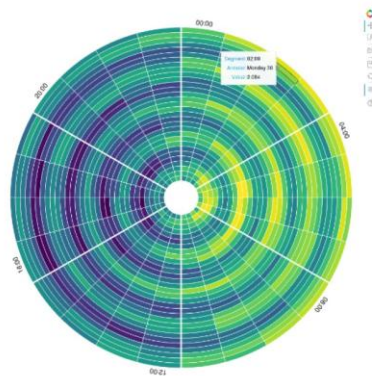
a) Line graph



b) Radial graph



c) Heatmap graph



d) Radial heatmap graph

Figure 1 - Time series visualisations in different techniques. Line graph (a), radial graph (b), heatmap graph (c), radial heatmap graph (d)

Real-time data visualization is particularly useful for IoT network systems that require collaborative decision-making. Collaborative decision-making involves multiple decision-makers working together. Real-time data visualization tools allow decision-makers to share data and collaborate in real-time, allowing them to make informed decisions more quickly and effectively. Electromagnetic situation awareness is an example of such decision making [8]. It utilises different visual analytics methods and charts, some of them are shown on Fig. 2.

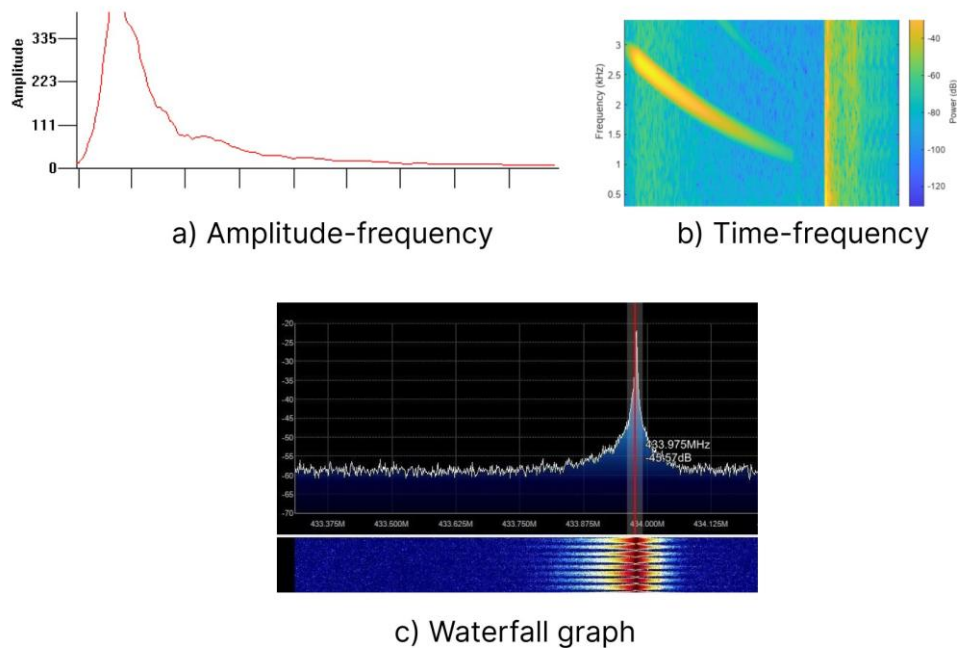


Figure 2 - Traditional spectrum diagrams: (a) spectrum amplitude-frequency diagram, (b) spectrum time-frequency diagram, (c) spectrum waterfall diagram

Formulation of Research Objective. The objective of this paper is to explore the challenges of real-time data visualization and propose strategies to optimize performance. Specifically, the paper aims to identify the performance issues that affect real-time data visualization and discuss potential solutions to these issues.

Presentation of the Main Research Material. A visualization system typically consists of several components including collection, data flow, processing, storage, and delivery. The collection component is responsible for gathering information while the data flow component facilitates the transfer of information from the collection components to the processing components. The processing components analyze the data and convert it into a suitable format for storage and delivery. Typically, the delivery is done through a visual interface such as a web or desktop program [9]. Let us formulate the challenges of real-time data visualization.

The field of visual analytics faces many challenges, especially when it comes to real-time data visualization. These challenges arise from specific applications of visual analytics, each with its own practical requirements in its problem domain. However, some challenges are common to more than one domain and application [1].

One of the most significant challenges is scalability. It is difficult to develop scalable visual analytics solutions that can handle both visual representations and automatic analysis. The solution needs to be able to scale in size, dimensionality,

data types, and levels of quality. Efficient methods are required to deal with noisy high-resolution input data as well as continuous input data streams of high bandwidth. Relevant data patterns and relationships need to be visualized on different levels of detail and with appropriate levels of data and visual abstraction.

Another major challenge is dealing with uncertainty. This is especially difficult in visual analytics due to the large amount of noise and missing values that originate from heterogeneous data sources and bias introduced by automatic analysis methods as well as human perception. To face this problem, the notion of data quality and the confidence of the analysis algorithm need to be appropriately represented. Analysts need to be aware of the uncertainty and be able to analyze quality properties at any stage of the data analysis process.

Hardware is also a significant challenge for real-time data visualization. More efficient computational methods and powerful hardware are needed to support near real-time data processing and visualization for large data streams. In addition to high-resolution desktop displays, advanced display devices such as large-scale power walls and small portable personal assistants need to be supported. Visual analytics systems should adapt to the characteristics of the available output devices, supporting the visual analytics workflow on all levels of operation.

Interaction is also a significant challenge in real-time data visualization. Novel interaction techniques are needed to fully support seamless intuitive visual communication with the system. The analyst should be able to fully focus on the task at hand and not be distracted by overly technical or complex user interfaces and interactions. User feedback should be taken as intelligently as possible, requiring as little user input as possible. Such interactions guarantee the full support of the user in navigating and analyzing the data, memorizing insights, and making informed decisions.

Evaluation is another challenge in visual analytics. Due to the interdisciplinary nature and complex visual analytics process, it is difficult to assess the quality of visual analytics solutions. A theoretically founded evaluation framework needs to be developed to assess the effectiveness, efficiency, and user acceptance of new visual analytics techniques, methods, and models. Such a framework will lead to a better understanding of the field and more successful and efficient development of innovative methods and techniques.

Finally, infrastructure is a challenge that needs to be addressed in real-time data visualization. Most current visual analytics solutions develop their infrastructures for solving their specific problems. Although some systems can be connected

through various communication mechanisms such as direct library linking and web services, there is still a mismatch between the level of service provided and the real need for visual analytics in terms of fast and precise answers with progressive refinement, incremental re-computation, and steering the computation towards data regions that are of higher interest to the user. More research is needed to develop a high-level infrastructure to bind together all the processes, functions, and services supplied by various disciplines. There is also a need to build repositories of available visual analytics solutions to ensure that common components are reusable.

Real-time data visualization presents several challenges, including performance issues that can slow down the system and affect the accuracy of the visualizations. These issues can arise due to the large volume of data being processed, the complexity of the visualizations, and the limitations of the hardware being used.

To optimize performance, several strategies can be used to reduce the amount of data being processed. One such strategy is data aggregation, where data is summarized at a higher level to reduce the volume of data being processed. For example, instead of plotting every data point on a graph, the data can be aggregated into smaller time intervals or geographic regions, reducing the number of data points being plotted. This not only reduces the amount of data being processed but also makes the visualization easier to read. Such method as k-means and DBSCAN may be used for clustering. However, traditional clustering methods experience difficulties in achieving accurate and efficient signal clustering, so custom clusterization algorithms may be required for efficient work in certain applications [8].

Another strategy is data pre-processing, where the data is cleaned and transformed before being visualized. This can involve removing duplicates, filling in missing values, and transforming the data to a format that is more suitable for visualization. By pre-processing the data, the visualization can be made more accurate and meaningful.

Compression is another technique that can be used to reduce the amount of data being processed. Compression involves compressing the data into a smaller format, making it easier to store and process. This can involve using algorithms such as gzip or deflate to compress the data before visualization or clustering on IoT device side before sending.

In addition to reducing the amount of data being processed, several techniques can be used to improve the speed and accuracy of the visualizations. Caching is one such technique, where the results of previous visualizations are stored in memory for

future use. This can significantly improve the speed of the visualizations, as the data does not need to be reprocessed each time.

Indexing is another technique that can be used to improve the speed of visualizations. Indexing involves creating an index of the data, allowing the system to quickly access specific data points without having to scan the entire dataset. This can significantly improve the speed of visualizations, especially when dealing with large datasets. This method is suitable for visualizations that use both real-time and stored data, for example in signal classification applications.

Incoming data grouping is also an essential way of optimizing performance. Developing a new algorithm that can group many data chunks from numerous devices in a fast and memory-efficient manner may increase performance.

Finally, parallel processing can be used to improve the speed and accuracy of the visualizations. Parallel processing involves breaking down the data into smaller chunks and processing them simultaneously on multiple processors. This can significantly reduce the processing time, allowing the visualizations to be generated more quickly.

By implementing these strategies, decision-makers can make more informed decisions in real-time, improving the overall efficiency and effectiveness of the decision-making process. However, it is important to note that these strategies may not be suitable for all types of data and visualizations, and it is important to carefully evaluate each strategy before implementation. Additionally, it may be necessary to balance performance with accuracy, as overly aggressive optimization may lead to inaccurate or misleading visualizations.

Conclusions. Real-time data visualization faces several challenges that require further research and development to overcome. These challenges include scalability, uncertainty, hardware, interaction, evaluation, and infrastructure. Addressing these challenges will enable the development of effective and efficient real-time data visualization solutions that can provide valuable insights for IoT network systems.

Future research could focus on exploring novel visualization techniques that can effectively represent complex and large-scale data in a more intuitive and interpretable manner (e.g., for electromagnetic situational awareness). Additionally, advancements in data grouping and data pre-processing, such as clusterization and pre-drawing (visual elements grouping) may significantly optimize real-time data visualization algorithms that may improve performance.

ЛІТЕРАТУРА

1. Keim D. Solving problems with visual analytics: Challenges and applications. / D. Keim, Z. Leishi. // Proceedings of the 11th International Conference on Knowledge Management and Knowledge Technologies. – 2011. – С. 1–4.
2. Mastering the information age solving problems with visual analytics / D.Keim, J. Kohlhammer, G. Ellis, F. Mansmann., 2010.
3. Cohen L. Time-frequency distributions-a review / Leon Cohen. // Proceedings of the IEEE 77. – 1989. – №7. – С. 941–981.
4. An overview of the NTIA/NIST spectrum monitoring pilot program / [M. Cotton, J. Wepman, J. Kub та ін.]. // 2015 IEEE Wireless Communications and Networking Conference Workshops (WCNCW). – 2015. – С. 217–222.
5. Timelines revisited: A design space and considerations for expressive storytelling. / [M. Brehmer, B. Lee, B. Bach та ін.]. // IEEE transactions on visualization and computer graphics. – 2016. – №24. – С. 2151–2164.
6. Adnan M. Adnan, Muhammad, Mike Just, and Lynne Baillie. "Investigating time series visualisations to improve the user experience / M. Adnan, M. Just, L. Baillie. // Proceedings of the 2016 CHI Conference on Human Factors in Computing Systems. – 2016. – С. 5444–5455.
7. Visual methods for analyzing time-oriented data / [W. Aigner, S. Miksch, W. Müller та ін.]. // IEEE transactions on visualization and computer graphics. – 2007. – №14. – С. 47–60.
8. Visual analytics for electromagnetic situation awareness in radio monitoring and management. / [Y. Zhao, X. Luo, X. Lin та ін.]. // IEEE transactions on visualization and computer graphics. – 2019. – №26. – С. 590–600.
9. Ellis B. Real-time analytics: Techniques to analyze and visualize streaming data / Byron Ellis. –Indianapolis: John Wiley & Sons, 2014.

REFERENCES

1. Keim D. Solving problems with visual analytics: Challenges and applications. / D. Keim, Z. Leishi. // Proceedings of the 11th International Conference on Knowledge Management and Knowledge Technologies. – 2011. – pp. 1–4.
2. Mastering the information age solving problems with visual analytics / D.Keim, J. Kohlhammer, G. Ellis, F. Mansmann., 2010.
3. Cohen L. Time-frequency distributions-a review / Leon Cohen. // Proceedings of the IEEE 77. – 1989. – №7. – pp. 941–981.

4. An overview of the NTIA/NIST spectrum monitoring pilot program / [M. Cotton, J. Wepman, J. Kub et al.]. // 2015 IEEE Wireless Communications and Networking Conference Workshops (WCNCW). – 2015. – pp. 217–222.
5. Timelines revisited: A design space and considerations for expressive storytelling. / [M. Brehmer, B. Lee, B. Bach et al.]. // IEEE transactions on visualization and computer graphics. – 2016. – №24. – pp. 2151–2164.
6. Adnan M. Adnan, Muhammad, Mike Just, and Lynne Baillie. "Investigating time series visualisations to improve the user experience / M. Adnan, M. Just, L. Baillie. // Proceedings of the 2016 CHI Conference on Human Factors in Computing Systems. – 2016. – pp. 5444–5455.
7. Visual methods for analyzing time-oriented data / [W. Aigner, S. Miksch, W. Müller et al.]. // IEEE transactions on visualization and computer graphics. – 2007. – №14. – pp. 47–60.
8. Visual analytics for electromagnetic situation awareness in radio monitoring and management. / [Y. Zhao, X. Luo, X. Lin et al.]. // IEEE transactions on visualization and computer graphics. – 2019. – №26. – pp. 590–600.
9. Ellis B. Real-time analytics: Techniques to analyze and visualize streaming data / Byron Ellis. – Indianapolis: John Wiley & Sons, 2014.

Received 20.10.2023.

Accepted 24.10.2023.

Візуалізація даних у режимі реального часу для систем мереж IoT:

виклики та стратегії оптимізації продуктивності

Візуалізація даних в реальному часі стала невід'ємним інструментом для систем прийняття рішень у різних галузях, включаючи фінанси, охорону здоров'я, IoT та виробництво. Візуалізація даних в реальному часі дозволяє організаціям відслідковувати та аналізувати дані в момент їхнього створення. Проте візуалізація даних в реальному часі має перед собою кілька викликів, включаючи продуктивність, якість даних та складність візуалізації. У цій статті було розглянуто важливість візуалізації даних в реальному часі в системах IoT та пов'язані з цим виклики. Зокрема, у статті обговорено виклики візуалізації даних в реальному часі та ідеї щодо підвищення продуктивності. Стаття також надає комплексний аналіз впливу візуалізації даних в реальному часі на системи IoT та процеси прийняття рішень, висвітливши її переваги та потенційні недоліки. Стаття містить обговорення важливості візуалізації даних в реальному часі в системах IoT. Далі у статті проведено аналіз викликів, пов'язаних з візуалізацією даних в реальному часі, включаючи якість даних, складність візуалізації та продуктивність.

Стаття представляє детальний аналіз кожної проблеми, висвітлюючи потенційний вплив на системи візуалізації даних в реальному часі та процеси прийняття рішень. Крім того, у статті висвітлені ідеї щодо підвищення продуктивності візуалізації даних в реальному часі, включаючи впровадження високопродуктивної обчислювальної інфраструктури, оптимізацію обробки та аналізу даних, використання технік кешування, використання технік візуалізації, оптимізованих для продуктивності, впровадження стиснення даних. У висновку у статті підкреслено важливість візуалізації даних в реальному часі в системах Інтернету речей та необхідність вирішення пов'язаних із цим викликів. Стаття також надає рекомендації щодо впровадження систем візуалізації даних в реальному часі, висвітлюючи ключові питання та найкращі практики для забезпечення успішного впровадження та оптимальної продуктивності.

Ключові слова: візуалізація даних в реальному часі, IoT, прийняття рішень, оптимізація продуктивності.

Лук'янець Михайло Олександрович – аспірант кафедри програмного забезпечення комп'ютерних систем, Національний технічний університет України «Київський політехнічний інститут імені Ігоря Сікорського».

Сулема Євгенія Станіславівна – д.т.н., доцент, завідувач кафедри програмного забезпечення комп'ютерних систем, Національний технічний університет України «Київський політехнічний інститут імені Ігоря Сікорського».

Lukianets Mykhailo Oleksandrovykh - Postgraduate student of the Computer Systems Software Department, National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute".

Sulema Yevgeniya Stanislavivna – Doctor of Technical Sciences, Associate Professor, Head of the Computer Systems Software Department, National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute".

В.Ю. Песчанський, Є.С. Сулема

ПРОЄКТУВАННЯ АРХІТЕКТУРИ ПРОГРАМНОЇ СИСТЕМИ ДЛЯ СТВОРЕННЯ ЦИФРОВИХ ДВІЙНИКІВ МЕДИКО-БІОЛОГІЧНИХ ОБ'ЄКТІВ

Анотація. У статті запропоновано універсальну об'єктно орієнтовану архітектуру програмної системи, призначеної для створення двійників медико біологічних об'єктів на прикладі отоларингології. Архітектура ґрунтується на використанні поліморфізму для мінімізації роботи розробників програмного забезпечення у разі виникнення необхідності модифікації системи для використання в інших галузях медицини. Розглянуто групи компонентів системи та надано рекомендації, щодо їхнього розроблення.

Ключові слова: архітектура програмної системи, цифрові двійники, медичне програмне забезпечення.

Постановка проблеми. Останнім часом технології створення цифрових двійників стають дедалі більш поширеними, зокрема, у галузі медицини та біології. Це пов'язано з тим, що використання цифрових моделей медико-біологічних об'єктів дозволяє більш ефективно вивчати їх фізіологію, патологію та взаємодію з лікарськими засобами. Наприклад, створення цифрових двійників людського тіла дозволяє моделювати патології та експериментувати з різними методами лікування, що дозволяє більш точно визначити ефективність та безпеку лікарських засобів. Таким чином, постає науково-практична задача створення універсальної модульної архітектури програмного забезпечення, яке, за рахунок використання цифрових двійників полегшить стеження за внутрішнім органом людини чи їх групою, а також значно спростить передбачення реакцій на медичне втручання. При цьому важливою задачею є модульність архітектури, яка дозволить швидко масштабувати систему для моделювання різних внутрішніх органів або систем.

Аналіз останніх досліджень. Аналіз рівня застосування медичних двійників у галузі медицини показав, що, незважаючи на те, що тема не є новою [1], її практичне застосування ще залишається на досить початковому рівні. Станом на сьогодні, доволі багато стартапів намагаються розробити універсальну систему для догляду пацієнтів з використанням технології циф-

рових двійників, проте наразі таке універсальне рішення відсутнє. У зв'язку з цим постає актуальна науково-практична задача зі створення універсальної архітектури, яка надасть можливість моделювати органи людського організму у вигляді цифрових двійників медико-біологічних об'єктів.

Мета досліджень. Метою дослідження, представленого у цій статті, є створення універсальної архітектури програмного забезпечення точних та достовірних моделей медико-біологічних об'єктів, які можуть бути використані для проведення експериментів та досліджень в медицині та біології.

Викладення основного матеріалу досліджень. Перед початком проєктування архітектури, слід зазначити одну дуже важливу проблему, яку необхідно враховувати при роботі з медико-біологічними об'єктами, а саме їх величезну диференціацію [2]. Для кожного окремого внутрішнього органу пацієнта необхідно використовувати зовсім різні системи моделювання, які у більшості випадків будуть мало сумісні між собою, тому необхідно забезпечити модульність розробленої архітектури для забезпечення можливості масштабувати системи без необхідності перероблення вже існуючих компонентів.

Слід зазначити, що будь-який цифровий двійник медико-біологічного об'єкту є часово зв'язним та являє собою результат обробки потоку темпоральних даних. Тому до використовуваного у розроблюваній архітектурі сховища мають бути висунені вимоги щодо максимальної ефективності зберігання такого типу даних.

Також до будь-якої системи, що розроблюється для сфери охорони здоров'я, має бути висунуто вимоги щодо логування та збереження персональних даних, оскільки згідно з проведеними дослідженнями та існуючим законодавством [3] гарантування цілісності та безпеки персональних даних є обов'язковим для впровадження будь-якої подібної системи. Отже, стійка система захисту даних повинна бути запроваджена на етапі розроблення архітектури.

Таким чином, можна сформулювати такі основні вимоги до розроблюваної системи.

1. Модульність та розширюваність.

- Система повинна бути побудована з можливістю легкого додавання нових функціональних модулів або оновлення існуючих.
- Модулі повинні бути добре ізольовані, щоб забезпечити можливість розвитку окремих компонентів незалежно.

- Архітектура повинна бути готовою до масштабування, якщо з'являться нові медичні дані або інші вимоги.

2. *Універсальність обробки даних.*

- Система повинна забезпечувати можливість збору та обробки різних типів даних, включаючи зображення, звуки, показники приладів тощо.
- Повинна бути підтримка обробки даних в реальному часі для забезпечення актуальності і точності цифрових двійників.

3. *Адаптивний аналіз даних.*

- Система повинна мати можливість аналізувати накопичені дані для виявлення паттернів, аномалій, а також для здійснення діагностики медичних станів пацієнтів.
- Розробка алгоритмів аналізу та визначення повинна бути гнучкою для адаптації до різних сценаріїв.

4. *Безпека та конфіденційність.*

- Забезпечення конфіденційності медичних даних пацієнтів є обов'язковим і повинно включати механізми шифрування, аутентифікації та авторизації.
- Всі оброблювані дані повинні відповідати вимогам стандартів безпеки в медичній галузі.

5. *Інтуїтивний інтерфейс.*

- Система повинна мати інтуїтивний інтерфейс для взаємодії зі спеціалістами.
- Повинна бути можливість відображення цифрових двійників у різних форматах та підтримка інструментів для їх аналізу.

6. *Відповідність медичним стандартам.*

- Розроблена система повинна відповідати всім вимогам та стандартам, які регулюють використання медичних технологій та обробку медичних даних.

На основі вищезазначеного, визначимо архітектуру програмної системи для створення цифрових двійників медико-біологічних об'єктів.

Всього пропонується п'ять основних груп компонентів: GUI, Storage, Data processing, Data analysis & Digital Twin constructor та Logger. При чому перші три є модульними для забезпечення відповідності сформованим вимогам, тобто вони надають певну базову функціональність, до якої можливо додавати модулі відповідно до потреб щодо подальшого розвитку системи.

Розглянемо зазначені групи компонентів детально на прикладі створення цифрового двійника гортані людини для використання у дослідженнях з отоларингології [4].

Система у цілому ґрунтується на мікросервісній архітектурі, яка характеризується гнучкістю та придатністю до масштабування [5].

Перша група компонентів реалізує узагальнений графічний інтерфейс (UI) та відповідає за взаємодію між користувачем (лікарем, медичним спеціалістом) та системою. Пропонується модульна організація графічного інтерфейсу, причому кожен з модулів є розширенням базисного UI для відповідної групи компонентів з обробки даних медико-біологічних об'єктів. Передбачається, що графічний інтерфейс користувача розроблений засобами JavaScript та HTML та буде комунікувати з основними обробниками за допомогою REST API та сокетів. Узагальнену схему даної групи представлено на рис. 1.

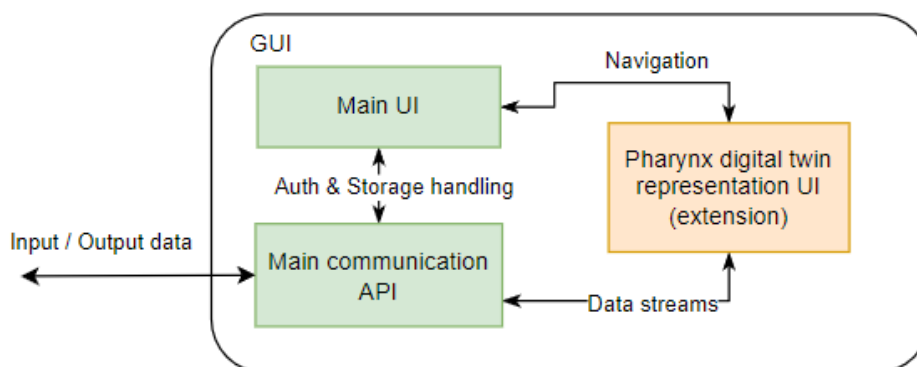


Рисунок 1 - Архітектура модулю графічного інтерфейсу

Друга група компонентів відповідає за обробку узагальнених даних для побудови цифрового двійника, а в окремих випадках і усіх мультимедійних даних, отриманих у реальному часі. Для їхнього зберігання рекомендується використовувати базу даних PostgreSQL, як одну з найбільш придатних для зберігання темпоральних даних [6].

Третя група компонентів відповідає за обробку та аналіз медичних даних для створення уніфікованого потоку даних з метою їх подальшого використання у створенні цифрових двійників. Результатом роботи даної частини архітектури є потік уніфікованих даних, які будуть придатні для аналізу та отримання з них відповідних метрик. Для підвищення відсотку перевикористання коду рекомендується створювати нові модулі обробки за допомогою шаблону “Прототип”. Узагальнену схему цієї групи представлено на рис. 2.

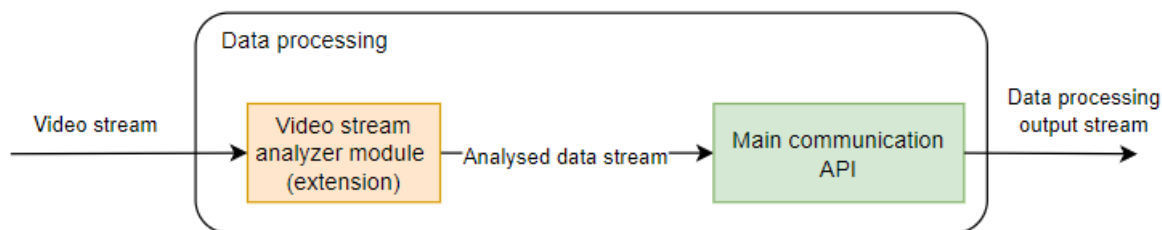


Рисунок 2 - Архітектура модулю обробки даних

Четверта група компонентів відповідає за аналіз та обробку медичних даних для побудови цифрового двійника пацієнта, а саме, виявлення аномалій, патологій та інших важливих характеристик на основі отриманих даних, а також за генерацію цифрового двійника медико-біологічного об'єкта пацієнта. Основна ідея функціонування цього модулю полягає у тому, що він надає основу функціональності та комунікації з іншими модулями, які потім розширюються за рахунок додаткових функціональних можливостей. Гнучкість реалізації досягається використанням шаблонів проєктування “Адаптер” та “Фасад”, які забезпечать зручний та зрозумілий базовий інтерфейс для взаємодії з будь-якою новою базою для двійника медико-біологічного об'єкта та потоком даних. Дотримання цих рекомендацій та застосування відповідних шаблонів проєктування допоможе створити модуль аналізу та побудови цифрового двійника, який буде гнучким, ефективним та забезпечить високу якість аналізу медичних даних у проєкті. Узагальнену схему цієї групи компонентів представлено на рис. 3.

Слід зазначити, що четверта та третя група тісно пов'язані, однак для полегшення розширення системи та внесення модифікацій до системи рекомендується їх розділити через те, що навіть різні галузі медицини використовують однакові принципи отримання потоку мультимодальних даних; прикладом цього може слугувати схожість пристроїв для отримання відеоданих у гастроентерології та оториноларингології [7]. Таким чином, забезпечується додаткова гнучкість пропонованої архітектури.

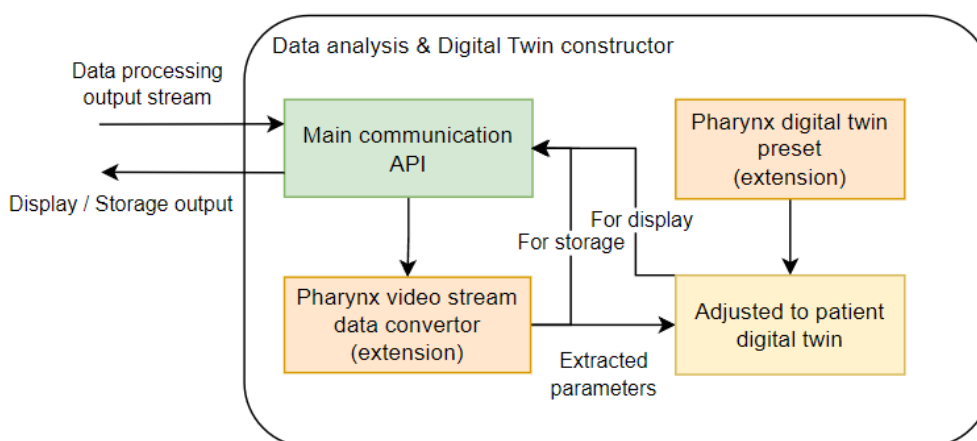


Рисунок 3 - Архітектура створення цифрових двійників

П'ята група компонентів відповідає за логування усіх дій системи та запис сесій користувачів. Також через розширення цієї частини системи можлива подальша інтеграція з іншими медичними системами.

Усі вищезазначені групи пропонується розроблювати мовою Python зі вставками коду C/C++ у разі необхідності роботи зі специфічними пристроями вводу [8]. Такий підхід дозволяє поєднати легкість розроблення та ефективність мови Python з можливостями оптимізованого виконання мови C для виконання задач з підвищеними вимогами щодо рівня оптимізації. Це особливо корисно у ситуаціях, коли потрібна висока швидкість обробки даних. Все вищезазначене робить обрані інструменти розроблення доцільними для запропонованої архітектури. Узагальнену архітектуру системи представлено на рис. 4.

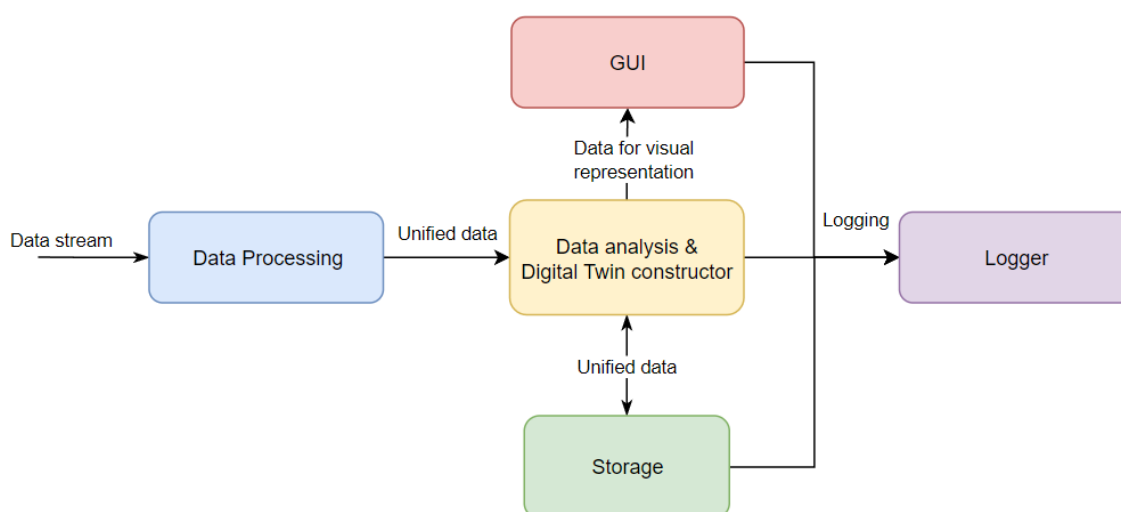


Рисунок 4 - Узагальнена архітектура програмної системи для створення
двійників медико-біологічних об'єктів

Висновки. Запропоновано нову універсальну архітектуру програмної системи програмної системи, призначеної для створення цифрових двійників медико-біологічних об'єктів, що дозволяє легко масштабувати та підтримувати систему, оскільки всі реалізації модулів реалізовані окремо та інкапсульовані. Практична цінність цього підходу полягає у забезпеченні можливості використання методів та алгоритмів штучного інтелекту для аналізу медичних даних, а також у можливості побудови цифрових двійників медико-біологічних об'єктів із забезпеченням масштабування програмної системи за потреби. Це створює нові можливості для більш точної діагностики та індивідуального підходу до лікування. Крім того, система забезпечує високий рівень безпеки та контролю завдяки модулю логування, що дозволяє відстежувати події та забезпечувати аудит доступу до даних. Це важливий аспект у медичній галузі, де конфіденційність та захист інформації мають вирішальне значення.

ЛІТЕРАТУРА

- Björnsson, B., Borrebaeck, C., Elander, N., Gasslander, T., Gawel, D.R., Gustafsson, M., Jörnsten, R., Lee, E.J., Li, X., Lilja, S. and Martínez-Enguita, D., 2020. Digital twins to personalize medicine. *Genome medicine*, 12, pp.1-4.
- Benson, M., 2023. Digital Twins for Predictive, Preventive Personalized, and Participatory Treatment of Immune-Mediated Diseases. *Arteriosclerosis, Thrombosis, and Vascular Biology*, 43(3), pp.410-416.
- Eva, G., Liese, G., Stephanie, B., Petr, H., Leslie, M., Roel, V., Martine, V., Sergi, B., Mette, H., Sarah, J. and Laura, R.M., 2022. Position paper on management of personal data in environment and health research in Europe. *Environment international*, 165, p.107334.
- Ihler, F. and Canis, M., 2019. The Role of the Internet for Healthcare Information in Otorhinolaryngology. *Laryngo-Rhino-Otologie*, 98(S 01), pp.S290-S333.
- Li, S., Zhang, H., Jia, Z., Zhong, C., Zhang, C., Shan, Z., Shen, J. and Babar, M.A., 2021. Understanding and addressing quality attributes of microservices architecture: A Systematic literature review. *Information and software technology*, 131, p.106449.
- Makris, A., Tserpes, K., Spiliopoulos, G. and Anagnostopoulos, D., 2019, March. Performance Evaluation of MongoDB and PostgreSQL for Spatio-temporal Data. In *EDBT/ICDT Workshops*.
- Papa, A., Mital, M., Pisano, P. and Del Giudice, M., 2020. E-health and wellbeing monitoring using smart healthcare devices: An empirical investigation. *Technological Forecasting and Social Change*, 153, p.119226.
- Luther, A.C., 1990. Digital video in the PC environment. McGraw-Hill, Inc.

REFERENCES

- Björnsson, B., Borrebaeck, C., Elander, N., Gasslander, T., Gawel, D.R., Gustafsson, M., Jörnsten, R., Lee, E.J., Li, X., Lilja, S. and Martínez-Enguita, D., 2020. Digital twins to personalize medicine. *Genome medicine*, 12, pp.1-4.
- Benson, M., 2023. Digital Twins for Predictive, Preventive Personalized, and Participatory Treatment of Immune-Mediated Diseases. *Arteriosclerosis, Thrombosis, and Vascular Biology*, 43(3), pp.410-416.
- Eva, G., Liese, G., Stephanie, B., Petr, H., Leslie, M., Roel, V., Martine, V., Sergi, B., Mette, H., Sarah, J. and Laura, R.M., 2022. Position paper on management of personal data in environment and health research in Europe. *Environment international*, 165, p.107334.
- Ihler, F. and Canis, M., 2019. The Role of the Internet for Healthcare Information in Otorhinolaryngology. *Laryngo-Rhino-Otologie*, 98(S 01), pp.S290-S333.
- Li, S., Zhang, H., Jia, Z., Zhong, C., Zhang, C., Shan, Z., Shen, J. and Babar, M.A., 2021. Understanding and addressing quality attributes of microservices architecture: A Systematic literature review. *Information and software technology*, 131, p.106449.
- Makris, A., Tserpes, K., Spiliopoulos, G. and Anagnostopoulos, D., 2019, March. Performance Evaluation of MongoDB and PostgreSQL for Spatio-temporal Data. In *EDBT/ICDT Workshops*.
- Papa, A., Mital, M., Pisano, P. and Del Giudice, M., 2020. E-health and wellbeing monitoring using smart healthcare devices: An empirical investigation. *Technological Forecasting and Social Change*, 153, p.119226.
- Luther, A.C., 1990. Digital video in the PC environment. McGraw-Hill, Inc.

Received 25.10.2023.

Accepted 28.10.2023.

Design of software system architecture for medical-biological objects digital twins creating

The paper introduces an innovative and universal object-oriented architecture specifically tailored for the creation of digital twins pertaining to medico-biological entities, with otolaryngology serving as the primary exemplar. The very essence of this architecture is its profound utilization of polymorphism. By capitalizing on this technique, the architecture offers a substantial reduction in the workload and complexity software developers might face when there's a necessity to adapt or modify the system for alternative medical specializations.

The genesis of this architecture stems from a recognized need to establish a more modular and adaptable framework that can cater to the rapidly evolving landscape of medical technology and digital twinning. As the medical realm continues to burgeon,

having such a flexible and scalable software blueprint becomes paramount, especially when aiming to maintain relevance and efficacy in diverse medical sectors.

The paper offers a comprehensive breakdown of the various component groups that constitute the system. Each group is dissected to provide clarity on its function, importance, and integration with the broader system. This hierarchical and systematic representation ensures that software developers, medical researchers, and other stakeholders have a clear roadmap when engaging with the architecture.

In summation, this research stands as a pivotal contribution to the intersection of medical technology and software development. By providing a robust and adaptable architectural blueprint for digital twinning in the medico-biological domain, it paves the way for future innovations and ensures that the medical community can harness the power of digital replication with increased accuracy, efficiency, and versatility.

Песчанський Владислав Юрійович – аспірант кафедри програмного забезпечення комп'ютерних систем Національний технічний університет України «Київський політехнічний інститут імені Ігоря Сікорського», Київ.

Сулема Євгенія Станіславівна – д-р техн. наук, доцент, завідувач кафедри програмного забезпечення комп'ютерних систем. Національний технічний університет України «Київський політехнічний інститут імені Ігоря Сікорського», Київ.

Vladyslav Peschanskii – Post-Graduate Student of Computer Systems Software Department National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute", Kyiv.

Yevgeniya Sulema – DSc, Associate Professor, Head of Computer Systems Software Department. National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute", Kyiv.

AN ALGORITHM FOR SOLVING A TWO-STAGE CONTINUOUS-DISCRETE LOCATION PROBLEM FOR MEDICAL LOGISTICS OPTIMIZATION

Abstract. The research paper focuses on logistics optimization, a critical component in supply chain management across various sectors, including healthcare. Efficient coordination of medical logistics is essential for maintaining public health and welfare, particularly during global emergencies where quick and effective distribution of medicine is crucial. This study aims to create and analyze a model and algorithm for a two stage continuous discrete location problem within medical logistics applications. We present a mathematical model tailored for a two stage continuous discrete location problem in medical logistics, considering the unique aspects of this field. The solution algorithm combines genetic methods with the optimal partition of sets theory. Additionally, we demonstrate the algorithm's effectiveness through a software application, using it to solve a representative model problem.

Keywords: continuous discrete two stage location problem, optimal set partitioning, genetic algorithm, medical logistics

Introduction. The optimization of logistics processes constitutes a critical component of supply chain management across various industrial sectors, notably in the healthcare domain. Effective medical logistics management is imperative in safeguarding public health and well-being, a priority that becomes paramount during global crises necessitating rapid and efficient distribution of medicinal products. Furthermore, proficient logistics organization is vital for delivering humanitarian aid, where the prompt provision of medical supplies and resources is often critical for survival. Contemporary technologies and algorithmic advancements are pivotal in enhancing medical logistics processes. A notable strategy that has garnered considerable acceptance is the utilization of genetic algorithms to address two-stage location challenges, optimizing the distribution and location of facilities. This methodology enables medical institutions and suppliers to make informed decisions, thereby augmenting the efficacy, expediency, and economic viability of medical logistics operations.

Literature review. Metaheuristic algorithms find extensive application in research about problems similar to those in medical logistics. The authors of [1] utilized a genetic algorithm to examine a two-stage transportation problem that involved a fixed route fee and the transport of goods. Research [2] introduced an iterative algorithm for a similar two-stage transportation challenge, incorporating a customer shuffling mechanism and an efficient method for duplicate elimination, proving effective across diverse data sets. Publication [3] is dedicated to a multi-stage reverse logistics network employing a genetic algorithm with priority encoding. The objective of enhancing spatial planning for public health services via the development of location-allocation and accessibility models is explored in [4]. This study aims to identify optimal locations for hospitals and other healthcare facilities, considering factors such as population demand, accessibility, and proximity to other medical establishments. The research, exemplified by the healthcare system in Lisbon, Portugal, demonstrated that the application of these methods significantly improved healthcare quality and cost efficiency. Article [5] discusses a two-stage transportation problem to minimize logistics expenses, considering the costs associated with establishing distribution centers and transportation charges between businesses, distribution centers, and customers. Research [6] addresses the complex issue of ICU bed allocation. The described problem is approached by decomposing it into two different stages to handle multiple sources of uncertainty effectively. The study [7] focuses on hospital capacity allocation planning for operating rooms. It presents an integrative solution approach combining a two-stage stochastic recourse programming model and a stochastic optimization model with a mean absolute deviation risk measure. Publication [8] assesses a location-allocation model for healthcare services. It employs a two-stage optimization strategy using a multi-period capacitated maximal-covering model, considering interservice referral and equity access. Research [9] aims at optimal demand coverage in healthcare facility location problems. It is proposed to primarily use optimization techniques with a focus on stochastic characteristics important for facility location decisions. The paper [10] explores capacity allocation and scheduling in two-stage service systems. The authors propose a simple and easy-to-implement capacity allocation and scheduling policy, establishing its asymptotic optimality for the stochastic system. Article [11] presents a new idea for the ambulance location problem in environments under uncertainty: it addresses the problem with a focus on robust planning for path and average speed uncertainty. The work [12] investigates resource allocation optimization in natural disasters with multiple secondary hazards. A two-stage stochastic optimization model is proposed

to simulate scenarios of random occurrence and severity of disasters. The papers [13, 14] develop a framework for supplier selection and order allocation. It introduces a hybrid multi-step long short-term memory network for demand forecasting and a new fuzzy SWOT model for supplier evaluation, followed by a multi-objective model using results from earlier stages. More information about multi-stage transportation problems for different areas, such as passenger traffic or goods transportation, can be found in [15, 16]. A more detailed analytic review of multi-stage location problems is listed in [17]. Substantial research in the field of discrete optimization is associated with genetic and stochastic approaches. However, this approach can be challenging due to the issue of scalability: in discrete scenarios, only small to medium-sized problems can be resolved efficiently. In situations requiring the location of many centers and considering population demand in facility location, continuum approaches to location problems become pertinent. Study [18] focuses on multi-stage location tasks, evaluating various facility groups and potential locations. This paper underscores the significance of considering continuum scenarios in such tasks and presents a model for two-stage location.

Problem statement. In times of crisis within a region, it becomes necessary to quickly distribute essential medical resources (medicines and medical equipment) to the people. To aid in this process, each region has a certain number of subregional centers (SRCs) that act as primary distribution points. However, due to logistical and resource constraints, only a certain number of these SRCs can be activated by the government. The chosen SRCs then redistribute the resources to multiple distribution centers (DCs) throughout the region, directly supplying the required medicines to the local population within their respective service areas. The main objective is to find the best combination of SRCs, positions of DCs and an optimal transportation strategy. The goal is to convey the medicines and medical supplies from the SRCs to the DCs and ultimately to the consumers while minimizing transportation costs. To achieve efficient distribution, careful consideration is required for factors such as transportation expenses, distance, and capacity limitations at each stage.

Figure 1 illustrates a complex transportation network, depicting a system with 4 SRCs and 8 DCs. This diagram is designed to provide a clear visual representation of the logistics and supply chain infrastructure. The active SRCs are highlighted with vibrant colors, symbolizing their operational status, whereas the inactive SRCs are differentiated by a more subdued, muted color palette, indicating their non-operational status.

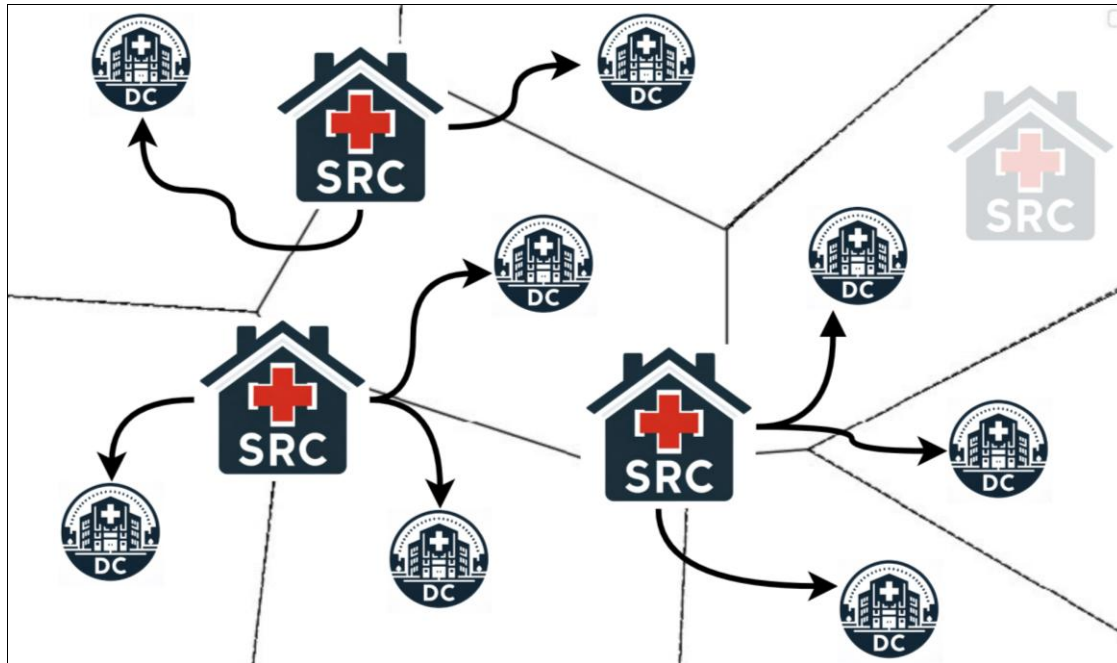


Figure 1 – Transportation schema for problem statement

Mathematical model. Moving forward, let us proceed with developing a mathematical model. We introduce the following definitions.

- Ω – the area where medical products are distributed to consumers.
- Ω_i – service area of the i -th DC, $i = \overline{1, N}$.
- N – the required number of DCs.
- b_j – capacity of j -th SRC, $j = \overline{1, M}$.
- J – set of SRCs that are available for activation: $J = \{\tau_1^{\text{II}}, \tau_2^{\text{II}}, \dots, \tau_M^{\text{II}}\}$.
- $c_i^{\text{I}} = c(x, \tau_i^{\text{I}})$ – transportation costs for volume weight unit of medicines and medical equipment from DC with coordinates τ_i^{I} to consumer $x \in \Omega$, $i = \overline{1, N}$.
- $c_{ij} = c(\tau_i^{\text{I}}, \tau_j^{\text{II}})$ – transportation costs for volume weight unit of medicines and medical equipment from SRC with coordinates τ_j^{II} to DC with coordinates τ_i^{I} , $i = \overline{1, N}$. $j = \overline{1, M}$.
- $\rho(x)$ – demand for medicines and medical equipment at the point x in area Ω .

- A_j – operating expenses associated with the activation of a subregional center j , $j = \overline{1, M}$.
- $\tau_i^r = (\tau_{i1}^r, \tau_{i2}^r)$ – DC's coordinates ($r=I$) or SRC's coordinates ($r=II$).
- v_{ij}^I – the volume weight units number of medicines and medical equipment transported from SRC j to DC i , $i = \overline{1, N}$, $j = \overline{1, M}$.
- $\lambda_j = \begin{cases} 1, & \text{if } DC \ j \text{ is activated,} \\ 0, & \text{otherwise} \end{cases}, j = \overline{1, M}$.

The total expenses for activating sub-regional centers and transportation of medicines and medical equipment can be expressed as follows:

$$\sum_{j=1}^M A_j \lambda_j + \sum_{i=1}^N \int_{\Omega_i} c_i^I(x, \tau_i^I) \rho(x) dx + \sum_{i=1}^N \sum_{j=1}^M c_{ij} v_{ij}^I \lambda_j.$$

Under the following constraints:

The number of volume weight units of medicines and medical equipment delivered from i -th DC is equal to its demands:

$$\sum_{j=1}^M v_{ij}^I \lambda_j = \int_{\Omega_i} \rho(x) dx, i = \overline{1, N}.$$

The number of volume weight units of medicines and medical equipment transported from j -th SRC does not exceed its capacity:

$$\sum_{i=1}^N v_{ij}^I \leq \lambda_j b_j, j = \overline{1, M}.$$

The maximum number of activated SRCs does not exceed L , ($L > 0$):

$$\sum_{j=1}^M \lambda_j \leq L.$$

The service areas of distribution centers include the entire region Ω :

$$\bigcup_{i=1}^N \Omega_i = \Omega.$$

Each customer is served by only one distribution center:

$$\Omega_i \cap \Omega_j = \emptyset, i \neq j, i, j = \overline{1, N}.$$

Considering this, the mathematical model can be expressed as follows.

Problem I. Minimize

$$\sum_{j=1}^M A_j \lambda_j + \sum_{i=1}^N \int_{\Omega_i} c_i^I(x, \tau_i^I) \rho(x) dx + \sum_{i=1}^N \sum_{j=1}^M c_{ij} v_{ij}^I \lambda_j, \quad (1)$$

under the following constraints

$$\sum_{j=1}^M v_{ij}^I \lambda_j = \int_{\Omega_i} \rho(x) dx, i = \overline{1, N}, \quad (2)$$

$$\sum_{i=1}^N v_{ij}^I \leq \lambda_j b_j, j = \overline{1, M}, \quad (3)$$

$$\sum_{j=1}^M \lambda_j \leq L, \quad (4)$$

$$\bigcup_{i=1}^N \Omega_i = \Omega, \quad (5)$$

$$\Omega_i \cap \Omega_j = \emptyset, i \neq j, i, j = \overline{1, N}, \quad (6)$$

$$v_{ij}^I \geq 0, \lambda_j \in \{0; 1\}, i = \overline{1, N}, j = \overline{1, M}, \quad (7)$$

$$\tau^I = (\tau_1^I, \tau_2^I \dots \tau_N^I), \tau^I \in \Omega^N. \quad (8)$$

Solution approach. We propose to use a combination of a genetic algorithm and an approach from the optimal set partitioning theory. Given this, we can divide the solution of problems (1) - (8) into two stages.

At the first stage, distribution centers are located with the determination of their service areas by solving optimal partition of set problem in formulation (9) - (11), where N – required number of DCs. The solution to the optimal set partitioning problem begins with the algorithm initializing the optimization domain. This domain is bounded by a parallelepiped oriented along the coordinate axes. Initial conditions such as the grid size and the algorithm's step size are established. The parallelepiped is covered with a grid that allows for calculations at discrete points. A preliminary approximation for the variables is then set.

The core of the optimization step involves the use of Shor's r-algorithm, which employs the concept of subgradient descent with a significant variation in the calculation of the step multiplier. In classical subgradient methods, the step size often decreases following a fixed rule. In contrast, the r-algorithm adaptively updates the step size based on previous iterations, allowing the algorithm to better account for the specifics of the problem at hand.

The main innovation of the algorithm lies in its use of regularization to enhance the characteristics of the optimization problem. The addition of regularization

helps to address the non-smoothness of the function and improves the convergence rate to the minimum. The r-algorithm is particularly effective for solving high-dimensional problems, where traditional optimization methods may need help with computational complexity and slow convergence.

The process of solving the optimal set partitioning problem is iterative, with each step updating the variable values based on prior calculations and adjusting the iterative process. The algorithm concludes its operation once a convergence criterion is met, either through the variable distance or the difference in functional values. For a more detailed version of the algorithm refer to [19].

Minimize

$$\sum_{i=1}^N \int_{\Omega_i} c_i^I(x, \tau_i^I) \rho(x) dx, \quad (9)$$

under the following constraints:

$$\bigcup_{i=1}^N \Omega_i = \Omega, \quad (10)$$

$$\Omega_i \cap \Omega_j = \emptyset, i \neq j, i, j = \overline{1, N}, \quad (11)$$

At the second stage, we solve a discrete location problem (12) – (18):

$$\sum_{j=1}^M A_j \lambda_j + \sum_{i=1}^N \sum_{j=1}^M c(\tau_i^I, \tau_j^{II}) \lambda_j v_{ij}^I \rightarrow \min, \quad (12)$$

under the following constraints:

$$\sum_{j=1}^M v_{ij}^I \lambda_j = b_i^*, i = \overline{1, N}, \quad (13)$$

$$\sum_{i=1}^N v_{ij}^I \leq \lambda_j b_j, j = \overline{1, M}, \quad (14)$$

$$\sum_{j=1}^M \lambda_j \leq L, \quad (15)$$

$$v_{ij}^I \geq 0, \lambda_j \in \{0; 1\}, i = \overline{1, N}, j = \overline{1, M}, \quad (16)$$

where $\tau^I = (\tau_1^I, \tau_2^I \dots \tau_N^I)$, $\tau^I \in \Omega^N$ – locations of the DCs obtained as a result of the first stage; b_i^* – determined capacity of DCs:

$$b_i^* = \int_{\Omega_i} \rho(x) dx, i = \overline{1, N}, \quad (17)$$

Additionally, there is a solvability condition:

$$\int_{\Omega} \rho(x) dx \leq \sum_{j=1}^M b_j \lambda_j, \quad (18)$$

An algorithm utilizing a genetic approach with priority encoding is developed to address the issue (12) - (18). The algorithm's general description is as follows.

1. Initially, the population $P(t)$ is set up through a priority-based encoding method, creating the first generation of potential solutions.
2. The fitness level of each individual chromosome within this population is then assessed.
3. For the purpose of reproduction, chromosomes are chosen based on the roulette wheel selection technique.
4. The reproductive process involves the crossover of selected chromosomes to yield new offspring.
5. The mutation process is applied to certain chromosomes to introduce new variations.
6. The next generation $P(t+1)$ is created by combining these offspring with some members of the current population, effectively replacing the previous generation.
7. The algorithm checks for the termination condition. If it's not met, the process loops back to step 2, now working with the updated population $t+1$.
8. Once the termination condition is met, the genetic algorithm finishes. The most efficient chromosome is decoded to reveal the objective function's value and the proposed transportation plan.

End of the algorithm.

A more detailed description of each step can be found in [20]. The following modifications to the genetic algorithm procedures should be highlighted.

- Each transportation plan between SRCs and DCs is represented as a chromosome using a priority-based coding scheme.
- The chromosome is checked to satisfy all model constraints when calculating the fitness function.
- A roulette selection procedure is used: a roulette wheel is created with segments proportional to these fitness scores. Selection is made by spinning the wheel and choosing individuals based on where it lands, favoring those with higher fitness but still allowing chances for less fit individuals.

- The weighted crossover method is used to cross chromosomes: first, two parent chromosomes are selected based on their fitness values, then each component of these decisions is assigned a weight that reflects its importance or effectiveness, which is determined by past performance, knowledge of the subject area and other heuristic methods. During crossbreeding, these weights determine the combination of parental decisions to create offspring. Components with higher weights are more likely to be passed on to offspring. It results in offspring that ideally inherit the best traits from both parents.

- The use of mixed mutation is proposed: with a probability of 0.5, swap or insertion mutation is used. For the first type of mutation, two unique elements are selected from each part of the chromosomes and swapped with each other. For the insertion mutation, a random index is chosen for each part of the chromosome, from which the element is rearranged to another random location.

Software implementation. We implemented a software application for numerical experiments designed to tackle the two-stage continuous-discrete location problem. This application was developed in Python 3 using the Qt5 framework for the user interface and map visualization. Qt5 provides the tools for creating interactive maps, while Python3 handles mathematical calculations and the core logic. C++ language and Eigen library are used for optimal partition numerical implementation. The development is done in PyCharm IDE. The calculations are performed using the following hardware: CPU: AMD Ryzen 7 5800X 3.8-4 GHz and 32 gigabytes of DD4-3200 RAM.

Model task. We use the proposed mathematical model, solving approaches and algorithm to optimize the medical logistics of the Dnipropetrovsk region, Ukraine. Here are the input parameters of the problem.

- Number of SRCs: $M = 7$.
- Number of potential DCs: $N = 25$.
- SRCs maximum number limit: $L = 6$.
- Vector that limits capacities of SRCs: [40, 40, 39, 40, 41, 39, 40].
- Total demand in DCs: 230.
- Vector of SRCs activation costs: [46, 50, 41, 48, 49, 42, 50].

The coordinates of the subregional centers are shown in Figure 2 as blue markers and correspond to the locations: Verkhnodniprovsk, Bozhedarivka, Tomakivka, Novomoskovsk, Synelnykove, Prosiana, Ternivka.

In the first stage, we solve the continuous-discrete location problem: a rectangle that delimits the geographical location is defined with the following coordinates.

- Left lower corner: (47.51611, 33.37109).
- Right upper corner: (48.95391, 36.93066).
- The map used in visualization is centered around Dnipro, Ukraine: (48.450001, 34.983334).

We use the Euclidean metric as a distance function for continuous cases for the described model task. The defined rectangle is shown as gray borders in Figure 2. As a result of the solution, we have the DCs locations as green markers in Figure 2. Along with locating each Distribution Center (DC), we also obtained its service area (optimal partition) marked as polygons with gray borders.

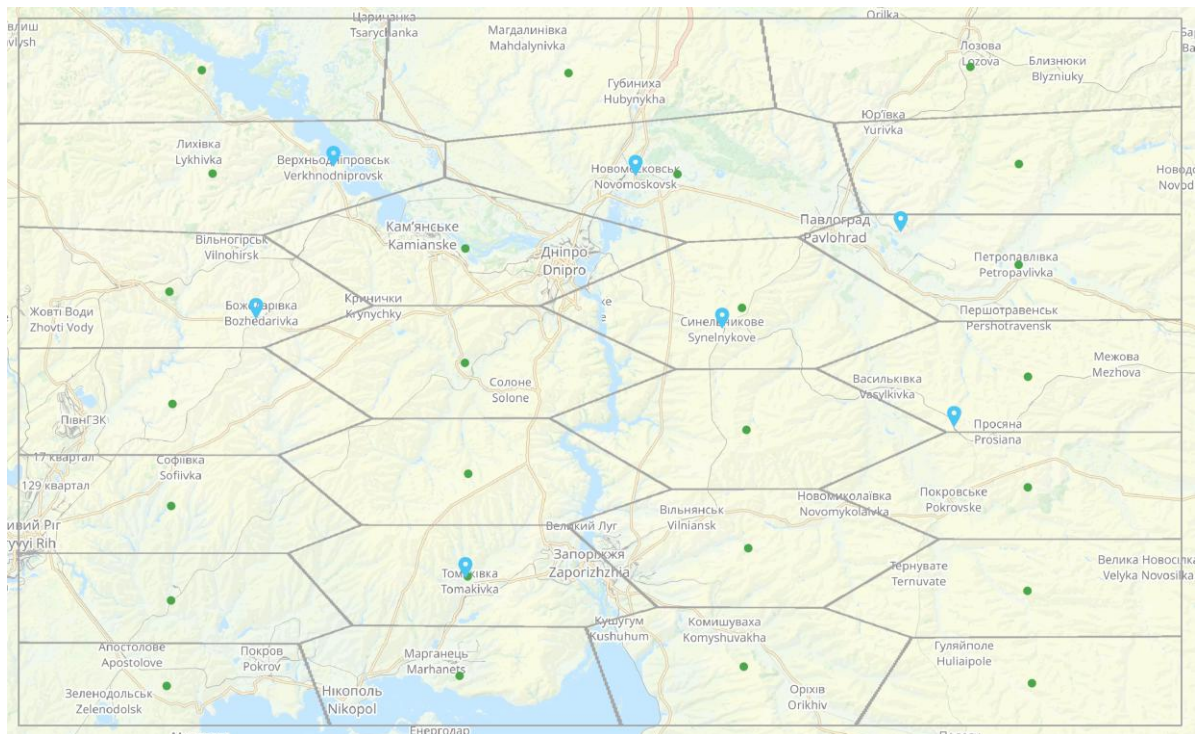


Figure 2 – Original positions of subregional centers and located distribution centers

Let's move on to the second stage. At this stage we are applying a genetic algorithm to solve discrete activation problem (12) - (18). In distance terms we use the geographic distance between the points as now we deal with geographical coordinates. The genetic algorithm will be executed with the following parameters:

- Population size: 50.
- Size of the chromosome: 32.
- Crossover probability: 0.9.
- Mutation probability: 0.15.
- Initial population generation: pseudo-randomly generated.

-Termination criteria: terminates after exceeding the number of 100 generations.

The overall result of the solution is shown in Figure 3 with the following results:

-The objective function value for the most effective solution to the problem is 1298.65.

-6 subregional centers were activated, except for the center in Novomoskovsk.

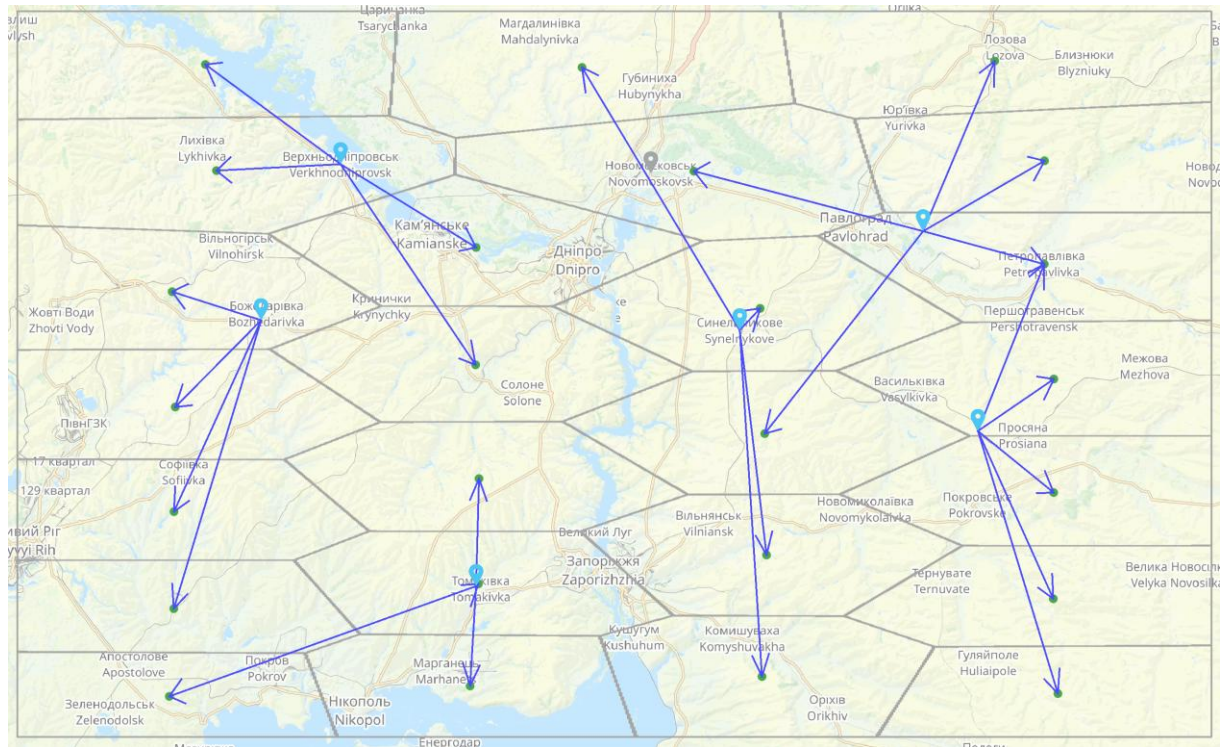


Figure 3 – Solution for the model task of two-stage continuous location problem

The total execution time of the algorithm (with visualization) took 13 seconds.

Each activated SRC is visually represented by a blue pin, indicating its active status and readiness to receive and process service requests. In contrast, any deactivated SRC is denoted by a gray pin, signifying its current inactivity and unavailability for handling service requests. Arrows emanate from each DC, illustrating the specific DC the particular SRC will service.

Conclusions. The research paper delves into a two-stage continuous-discrete location problem, imposing constraints on the maximal number of centers, primarily within the medical logistics sector. This area deals with distributing medicines and medical equipment, necessitating compliance with stringent requirements. The paper introduces a mathematical model that integrates optimal set partitioning with

metaheuristic techniques, intending to streamline the solution process into two distinct stages. Initially, the optimal set partitioning method is employed to determine the locations of distribution centers. Subsequently, the acquired coordinates and their corresponding service areas are utilized to address a discrete activation problem. The chosen metaheuristic method is a genetic algorithm, enhanced with priority-based encoding and weight mapping crossover, alongside specific adaptations. A specially designed fitness function evaluates the adherence to the model's constraints. Moreover, a mixed mutation approach is implemented in each iteration, based on a predetermined probability. The algorithm's output encompasses the established distribution centers, activated subregional centers, and a comprehensive transportation plan. A software tool was developed in Python 3 to demonstrate the algorithm's functionality, incorporating the Qt5 framework. The paper features a case study focused on the medical logistics in the Dnipropetrovsk region of Ukraine. It exemplifies the practical implementation of the proposed model and algorithm, highlighting their potential to enhance logistics and distribution efficiency in a specific geographic context.

Future research. Future research should focus on refining the location problem by linking subregional center positions with distribution centers, optimizing algorithm complexity, and applying the model to real-world logistics, using authentic data for regional specificity. Further, the model could be expanded to include distance and time constraints, enhancing its applicability in scenarios like rapid medical supply distribution, and promoting cost-effective, sustainable logistics.

Acknowledgements. This publication was prepared as a part of the scientific theme 0123U100011 "Problems of analysis, modeling, and optimization of technological processes in complex systems of different nature", which is implemented on the System Analysis and Control Department at Dnipro University of Technology.

REFERENCES

- 1) 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. DOI: 10.1016/j.cor.2011.09.020.
- 2) 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. DOI: 10.1016/j.procs.2019.12.066 (date of access: 15.11.2023).

- 3) 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. DOI: 10.1016/j.cie.2008.09.021 (date of access: 15.11.2023).
- 4) 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. e0119190. DOI: 10.1371/journal.pone.0119190 (date of access: 15.11.2023).
- 5) Gen M., Altiparmak 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. DOI: 10.1007/s00291-005-0029-9 (date of access: 15.11.2023).
- 6) Two-stage multi-objective optimization for ICU bed allocation under multiple sources of uncertainty / F. Wan et al. *Scientific reports*. 2023. Vol. 13, no. 1. DOI: 10.1038/s41598-023-45777-x (date of access: 22.11.2023).
- 7) Lalmazloumian M., Baki M. F., Ahmadi M. A two-stage stochastic optimization framework to allocate operating room capacity in publicly-funded hospitals under uncertainty. *Health care management science*. 2023. DOI: 10.1007/s10729-023-09644-5 (date of access: 22.11.2023).
- 8) Salami A., Afshar-Nadjafi B., Amiri M. A two-stage optimization approach for healthcare facility location-allocation problems with service delivering based on genetic algorithm. *International journal of public health*. 2023. Vol. 68. DOI: 10.3389/ijph.2023.1605015 (date of access: 22.11.2023).
- 9) An empirical evaluation of a two-stage healthcare facility location approach using simulation optimization and mathematical optimization / S. Mohd et al. *SSRN electronic journal*. 2022. DOI: 10.2139/ssrn.4246268 (date of access: 22.11.2023).
- 10) Capacity allocation and scheduling in two-stage service systems with multi-class customers / Z. Zhong et al. *SSRN electronic journal*. 2023. DOI: 10.2139/ssrn.4431264 (date of access: 22.11.2023).
- 11) Akincilar A., Akincilar E. A new idea for ambulance location problem in an environment under uncertainty in both path and average speed: absolutely robust planning. *Computers & industrial engineering*. 2019. Vol. 137. P. 106053. DOI: 10.1016/j.cie.2019.106053 (date of access: 22.11.2023).
- 12) A two-stage stochastic optimization for disaster rescue resource distribution considering multiple disasters / B. Wang et al. *Engineering optimization*. 2022. P. 1–17. DOI: 10.1080/0305215x.2022.2144277 (date of access: 22.11.2023).
- 13) Islam S., Hassanzadeh Amin S., Wardley L. J. Supplier selection and order allocation planning using predictive analytics and multi-objective programming. *Comput-*

ers & industrial engineering. 2022. P. 108825. DOI: 10.1016/j.cie.2022.108825 (date of access: 22.11.2023).

14) Islam M. T., Miah M. M. Optimization of interval cost multi-objective transportation problem in uncertain parameters: a modified least-cost method. *Journal of bangladesh academy of sciences*. 2022. Vol. 46, no. 2. P. 155–164. DOI: 10.3329/jbas.v46i2.62168 (date of access: 23.11.2023).

15) The forecasting of dangerous goods transport by rail in Poland in terms of environmental security / M. Stajniak et al. *European research studies journal*. 2022. Vol. XXV, Issue 2B. P. 359–368. DOI: 10.35808/ersj/2968 (date of access: 23.11.2023).

16) Wang X., Song J., Qunqi W. A novel bi-level programming model for structure sustainability optimization of passenger transport corridor. *PLoS ONE*. 2022. Vol. 17, no. 7. P. 1–21.

17) Serhieiev O., Us S. Analysis of modern approaches to solving discrete and continuous multi-stage allocation problems. *Information technology: computer science, software engineering and cyber security*. 2023. No. 2. P. 50–58. DOI: 10.32782/it/2023-2-7.

18) Us S. A., Koriashkina L. S., Stanina O. D. An optimal two-stage allocation of material flows in a transport-logistic system with continuously distributed resource. *Radio electronics, computer science, control*. 2019. No. 1. P. 256–271. DOI: 10.15588/1607-3274-2019-1-24 (date of access: 17.11.2023).

19) Kiselova O., Koriashkina L. *Nepervni zadachi optymalnoho rozbytta mnozhyn ta r-alhorytmy : monohrafiia*. Dnipro : Naukova dumka, 2015. 400 p.

20) Serhieiev O. S., Us S. A. Modified genetic algorithm approach for solving the two-stage location problem. *Radio electronics, computer science, control*. 2023. No. 3. P. 159–170. DOI: 10.15588/1607-3274-2023-3-16 (date of access: 16.11.2023).

Received 01.11.2023.

Accepted 15.11.2023.

Алгоритм розв'язання двоетапної неперервно-дискретної задачі розміщення на прикладі оптимізації медичної логістики

Оптимізація логістичних процесів є одним із важливих завдань управління ланцюгами поставок у різних сферах діяльності, включаючи медицину. Ефективна координація у сфері медичної логістики має важливе значення для забезпечення громадського здоров'я та процвітання. Це стає особливо актуальним в умовах глобальних надзвичайних ситуацій, коли швидке та ефективне розповсюдження медикаментів має вирішальне значення. Метою роботи є побудова моделі та роз-

робка алгоритму для розв'язання двоетапної неперервно-дискретної задачі розміщення у контексті проблеми медичної логістики з подальшим аналізом їх застосування на модельних прикладах. В роботі розглянута проблема транспортування медикаментів та виробів медичного призначення в Україні та визначена постановка задачі в предметній області. На першому етапі необхідно активувати субрегіональні центри з набору доступних, що є дискретною задачею оптимізації. Другий етап складається з розміщення центрів дистрибуції та визначення їх зон обслуговування. Авторами запропоновано математичну модель для двоетапної неперервно-дискретної задачі розміщення у контексті медичної логістики з врахуванням особливостей галузі. Розроблено алгоритм розв'язання, що базується на поєднанні генетичних підходів та теорії оптимального розбиття множин. Використано пріоритетне кодування з міксованою процедурою мутації. На першому кроці розв'язується задача оптимального розбиття множин із розміщенням центрів підмножин у кількості, що відповідає кількості необхідних центрів дистрибуції. На другому кроці, використовуючи координати розміщених центрів та визначені зони обслуговування, застосовується генетичний алгоритм для визначення ефективного розв'язку. Авторами розроблено програмну імплементацію запропонованого алгоритму з використанням мов програмування Python3 та C++ та бібліотек Qt5 для інтерфейсу користувача та Eigen для математичних обчислень. Запропонований підхід використано для розв'язання модельної задачі медичної логістики у Дніпропетровській області, Україна.

Ус Світлана Альбертівна – к.ф.-м.н., доцент, професор кафедри системного аналізу і управління НТУ «Дніпровська політехніка»

Сергєєв Олексій Сергійович - аспірант кафедри системного аналізу і управління НТУ «Дніпровська політехніка»

Us Svitlana Albertivna – candidate of physical and mathematical sciences, associate professor, professor of the Department of System Analysis and Control, NTU Dnipro Polytechnic, Dnipro, Ukraine.

Serhieiev Oleksii Serhiiovych – post-graduate student of the Department of System Analysis and Control, NTU Dnipro Polytechnic, Dnipro, Ukraine.

О.П. Іванов, В.В. Скалозуб, В.М. Горячкін, В.І. Шинкаренко

**ВИЗНАЧЕННЯ ЗДАТНОСТІ ШТУЧНОГО ІНТЕЛЕКТУ ДО ВСТАНОВЛЕННЯ
АВТОРСТВА ХУДОЖНЬОГО УКРАЇНОМОВНОГО ТЕКСТУ
ЗА ЗНАЧНИМИ ФРАГМЕНТАМИ**

Анотація. Штучний інтелект стає невід'ємною частиною побутового життя та професійної діяльності людини. Представлена робота має на меті дослідити ефективність визначення авторства художніх текстів за допомогою сучасних інструментів штучного інтелекту, а саме інтелектуальна пошукова система Bing. Для експерименту було вибрано десять українських авторів з багатим доробком художніх творів, які відображають різні аспекти української культури та історії. Випадкові фрагменти довжиною до 500 слів були вибрані з різних творів цих авторів. Було проведено експеримент з визначення авторства 360 фрагментів. Використовуючи мову програмування Python та пакет skru, було створено програмне забезпечення, яке надсилає запити та отримує відповіді від Bing бота, вбудованого в Microsoft Skype. В текстах відповідей перевірялася наявність імені автора фрази та відповідної назви твору. Встановлено, що довгі фрагменти дозволяють з високою точністю, але не завжди визначити автора художнього україномовного тексту.

Ключові слова: визначення авторства, природомовний текст, штучний інтелект, генеративні мовні моделі, ChatGPT, бот Bing, Skype.

Огляд предметної області. Визначення авторства тексту є однією з важливих і складних проблем, з якими стикаються дослідники в галузі обробки природної мови. Ця задача полягає в тому, щоб встановити, хто є реальним автором даного тексту, або ж, якщо текст має декілька авторів, як розподілити внесок кожного з них. Визначення авторства може мати різні застосування, такі як перевірка на плагіат, аналіз стилю письма, ідентифікація особистості, атрибуція літературних творів тощо. Однак більшість існуючих методів визначення авторства працюють ефективно лише для коротких текстів, таких як електронні листи, твіти або блог-пости. Для великих фрагментів тексту, таких як книги, наукові статті або романи, ця задача стає значно складнішою, оскільки потребує більш глибокого аналізу семантики, синтаксису і лексики тексту. У

статті розглядається можливість вирішення задачі визначення авторства великих фрагментів тексту з використанням засобів штучного інтелекту. Така задача має важливе значення для літературознавства, культурології та права інтелектуальної власності.

В останній час відбувається значний розвиток великих мовних моделей, які є потужним інструментом штучного інтелекту, вони здатні генерувати текст, розуміти мову та відповідати на запитання. Основна принципова відмінність таких моделей від традиційного розпізнавання природної мови полягає в тому, що вони базуються на нейронних мережах та глибокому машинному навчанні [3, 4].

Одна з найвідоміших великих мовних моделей – GPT (Generative Pre-trained Transformer), розроблена компанією OpenAI. Ця модель має вражаючі можливості в генерації текстів, розумінні мови та відповіді на запитання. Існують також інші великі мовні моделі, такі як BERT (Bidirectional Encoder Representations from Transformers), RoBERTa (Robustly Optimized BERT), XLNet (eXtreme-Long Transformer), T5 (Text-to-Text Transfer Transformer) та багато інших. Кожна з цих моделей має свої особливості та застосування [5, 6].

Сучасні великі мовні моделі вже застосовуються в багатьох сферах, включаючи машинний переклад, генерацію текстів, відповіді на запитання, семантичний аналіз текстів, автоматичну обробку мовлення та багато іншого [1]. Вони також використовуються в інтернет-пошуку, особистих асистентах, для розпізнавання мови та інших додатках, які потребують розуміння та генерації людської мови.

Великі мовні моделі, такі як GPT, навчаються на великих корпусах текстів, що дозволяє їм засвоювати граматику, лексику та контекстуальні залежності мови. Це дозволяє моделям генерувати тексти, які максимально наближені до людського мовлення. Великі мовні моделі мають свої обмеження. Наприклад, вони можуть виробляти неправдиву або неетичну інформацію, якщо їм будуть надані такі дані. Також вони можуть бути вразливі до атак, які можуть використовувати їх для поширення дезінформації або зловживання. Усупереч обмеженням, великі мовні моделі є потужними інструментами для автоматизації розуміння та генерації текстів, а їхні можливості продовжують зростати з кожним новим випуском моделей та покращенням алгоритмів навчання.

Компанія Microsoft створила чат Bing та Copilot, які є інноваційними продуктами, що використовують штучний інтелект на базі ChatGPT для надання корисної інформації та генерації творчого контенту [7]. Bing – це пошукова си-

стема, яка дозволяє знаходити відповіді на будь-які запитання, а чат Copilot – це чат-бот, який може спілкуватися з користувачами на різних мовах, створювати вірші, код, історії, пісні та багато іншого. Обидва продукти базуються на технології OpenAI, яка є одним з найпередовіших проектів у галузі штучного інтелекту. OpenAI – це некомерційна організація, яка має на меті створити дружній і загальнодоступний штучний інтелект, який буде працювати на благо людства. Microsoft також інтегрує Bing у свій браузер Edge, месенджер Skype і інвестує мільярди, щоб привести його в свої продуктивні додатки Office [8].

Мета дослідження. Метою роботи є підтвердження або спростування гіпотези що штучний інтелект здатний з точністю більше 95% встановити авторство відомих україномовних авторів за значними фрагментами їх творів. Bing – це інтелектуальна пошукова система, яка може допомогти встановити авторство українського літературного тексту. Bing дозволяє шукати дані про уривок та письменника, але пошукові результати можуть бути неповними або неточними.

Методика досліджень. Було обрано десять українських письменників, які написали значну кількість творів про різні сторони української культури та історії. База авторів та творів співпадає з представленою у роботах [9, 10]. Фрагменти обирались випадково із відібраних творів, з урахуванням пунктуації та структурних елементів тексту. Відбиралися фрагменти довжиною від 50 до 500 слів. Загальна кількість уривків представлена в Таблиці 1, з розподілом довжин представленим на рис. 1. Відповідь пошукової системи може змінюватися в залежності від довжини пошукового запиту, але це не завжди пряма залежність.

Використано мову програмування Python та бібліотеку skpy [2], щоб створити програму, яка взаємодіє зі Skype. За допомогою цієї програми передаються запитання і отримуються відповіді від бота Bing, вбудованого в програму Skype. Skpy є бібліотекою для Python, яка надає зручний спосіб взаємодії з API Skype. Вона дозволяє автоматизувати надсилання повідомлень і отримання відповідей з чатів або розмов.

У експерименті боту Bing пропонувалося відповісти на запитання у форматі "Дай коротку відповідь, хто автор фрагменту:%?".

Загальна кількість фрагментів відібраних до експерименту по авторам

Автор	9
Іван Багрянний	40
Остап Вишня	40
Марко Вовчок	28
Олександр Довженко	37
Михайло Коцюбинський	40
Панас Мирний	40
Всеволод Нестайко	39
Валер'ян Підмогильний	26
Іван Франко	38
Микола Хвильовий	40
Загалом	368

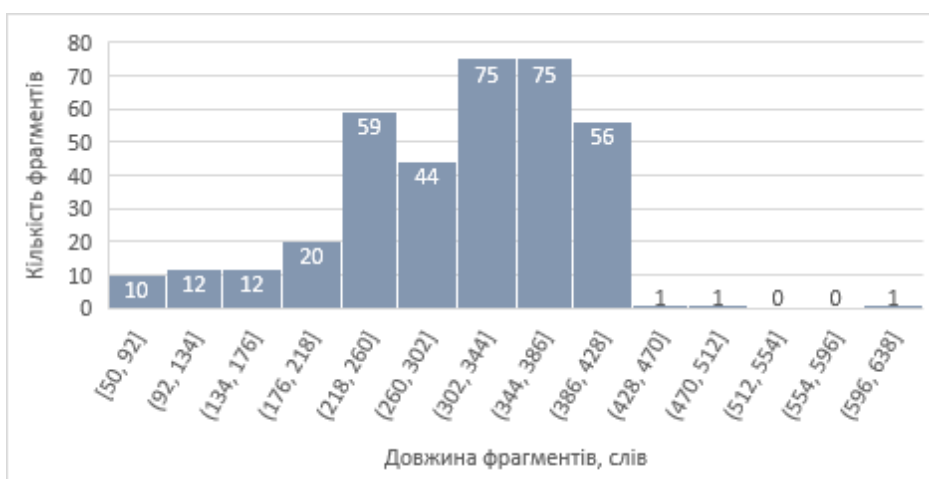


Рисунок 1 – Гістограма довжин фрагментів в експерименті

Результати досліджень. Усі отримані відповіді було зібрано в єдину загальну таблицю та проаналізовано. У текстах відповідей перевірено наявність імені автора фрагменту і відповідної назви твору.

Наступна таблиця представляє кількість відповідей, де було вірно визначено ім'я автора та назва твору, або вірно було визначено тільки ім'я автора. За отриманими результатами, Іван Франко має самий високий відсоток точних відповідей, не розпізнано було всього 13% фрагментів, визначено вірно автора та назву твору в 71% і в 16% вірно було розпізнано тільки ім'я автора. Найнижчий результат отримав Іван Багрянний 53% невизначених фрагментів.

Таблиця 2

Результат визначення авторства фрагментів творів в кількості фрагментів

Автор фрагменту	Точний результат	Тільки автор	Не визначено
Іван Багряний	13	6	21
Марко Вовчок	13	2	13
Микола Хвильовий	23	2	15
Михайло Коцюбинський	20	6	14
Остап Вишня	24	3	13
Валер'ян Підмогильний	14	4	8
Олександр Довженко	22	4	11
Панас Мирний	23	5	12
Всеволод Нестайко	21	7	11
Іван Франко	27	6	5
Загалом	200	45	123

Визначення назви твору за фрагментом може бути ускладненим, якщо існують серії творів одного автора, це може бути спричинено такими факторами:

- схожість тематики – якщо автор написав серію творів на одну тему або з одними й тими ж персонажами, може бути важко визначити, до якого саме твору належить конкретний фрагмент;
- схожість стилю – якщо автор має виразний стиль, який він використовує в усіх своїх творах, це також може ускладнити визначення конкретного твору;
- сюжетні зв'язки – якщо твори в серії мають сюжетні зв'язки, фрагменти з одного твору можуть здаватися подібними до фрагментів з іншого;
- повторення мотивів – якщо автор повторює певні мотиви або образи в різних творах;
- відсутність унікальних відмінностей – якщо фрагмент не містить унікальних відмінностей, які відрізняють один твір від іншого в серії.

Ці фактори можуть зробити визначення твору за фрагментом складним процесом, який вимагає глибокого знання творчості автора та контексту його творів.

Всеволод Нестайко був відомим українським письменником, особливо відомим своїми дитячими творами. Він написав багато книг, які стали популярними серед дітей та дорослих. “Тореадори з Васюківки” – це відома трилогія Всеволода Нестайка, яка розповідає про пригоди двох друзів.

За результатами експерименту, саме у Нестайка високий відсоток (18%), невірно визначених назв творів, фрагменти його творів були невірно визначені саме по назві твору (див. рис 2.).

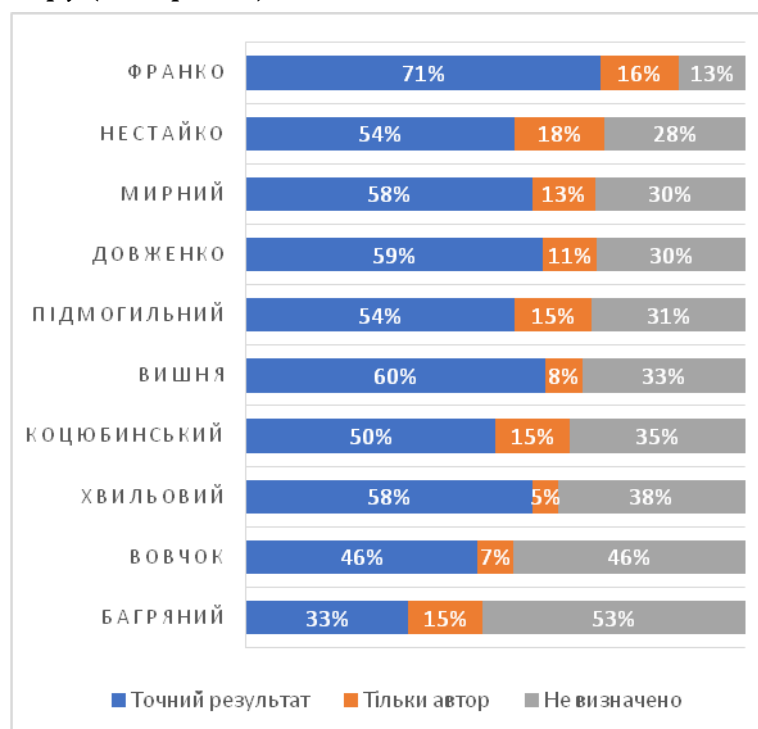


Рисунок 2 – Точність встановлення авторства фрагментів текстів

Псевдоніми письменників також можуть ускладнити визначення авторства фрагментів творів, особливо якщо автор використовував багато різних псевдонімів або якщо псевдонім не був широко відомим. Наприклад, Іван Багряний, справжнє ім'я Іван Миколайович Лозов'яга, був відомим українським письменником і політичним діячем. Іван Багряний використовував декілька псевдонімів у своїй творчості. Окрім відомого “Іван Багряний”, він також писав під псевдонімами “Полярний” та “Дон Кочерга”. Результати визначення авторства фрагментів творів, написаних Іваном Багряним, показали найнижчий результат: 53% невірного визначення, отже наявність псевдонімів може ускладнити процес.

В роботі [9] представлені дослідження відповідей бота Bing для визначення авторства надкоротких фрагментів довжиною в 3-7 слів. Порівняння точності встановлення авторства приведено в таблиці 2. Визначення авторства Івана

Франко має саму високу точність, що за короткими, що за довгими фрагментами (див. Таблиця 3).

Точність визначення авторства зазвичай підвищується з розміром фрагментів. Більші фрагменти надають більше контексту та інформації, які можуть допомогти визначити авторський стиль, тематику, мову та інші характеристики, що є відмінними для конкретного автора.

Таблиця 3

Підвищення точності визначення авторства при порівняння коротких та довгих фрагментів

Автор	Надкороткі фрагменти	Довгі фрагменти	Зміна точності
Марко Вовчок	47%	54%	15%
Панас Мирний	56%	70%	25%
Іван Багряний	36%	48%	30%
Іван Франко	65%	87%	33%
Остап Вишня	51%	68%	34%
Всеволод Нестайко	40%	72%	78%
Михайло Коцюбинський	27%	65%	139%
Михайло Хвильовий	24%	63%	156%
Валер'ян Підмогильний	25%	69%	182%
Олександр Довженко	23%	70%	202%

Однак, для авторів з вираженим стилем, точність може підвищуватися менше з розміром фрагментів. Це тому, що їх унікальний стиль може бути впізнаний навіть у невеликих фрагментах. Наприклад, автори, які використовують унікальні обороти мови, метафори або інші стилістичні елементи, можуть бути впізнані за допомогою цих особливостей. Важливо зауважити, що точність визначення авторства також може залежати від інших факторів, таких як якість тексту, наявність помилок, контекст та інших факторів.

Популярність творів автора також впливає на точність визначення авторства. Чим більше творів автора відомо і доступно для аналізу, тим більше даних є для визначення його унікального стилю та манери написання. Ці фактори впливають на точність визначення авторства твору або на коректне визначення назви твору.

Найнижчий приріст точності визначення характеризує однаково впізнаваний стиль на будь якій довжині фрагментів. Наприклад, Іван Франко творив майже в усіх літературних жанрах: поезії; поеми (соціально-побутові, сатирично-політичні, гумористичні, історичні, біографічні, філософські) казки. Його творчий доробок, писаний українською (більшість текстів), польською, німецькою, російською, болгарською, чеською мовами, за приблизними оцінками налічує кілька тисяч творів загальним обсягом понад 100 томів.

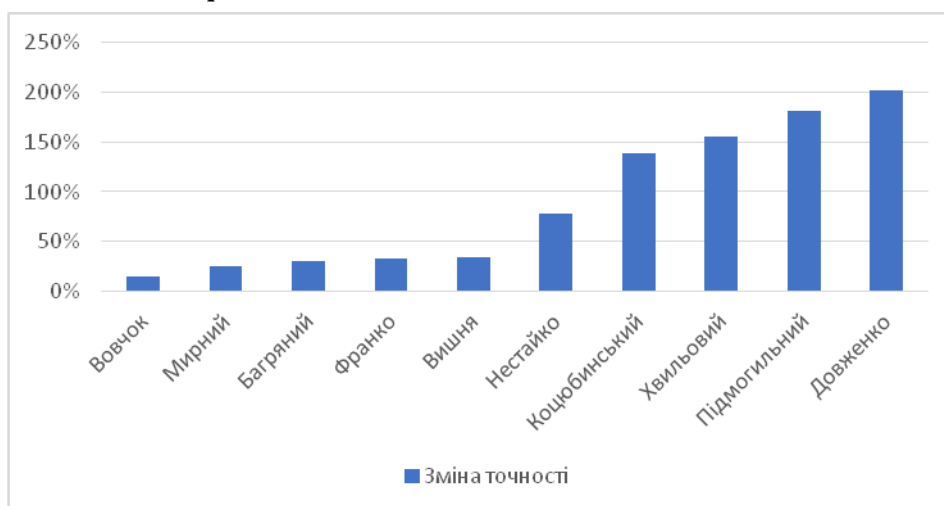


Рисунок 3 – Приріст точності визначення авторства при використанні довгих фрагментів

В відповідях до фрагментів з невірно визначеним авторством було відібрано ім'я помилково визначеного автора. Кількість відповідей з помилково призначеними авторами приведено на рисунку 4.

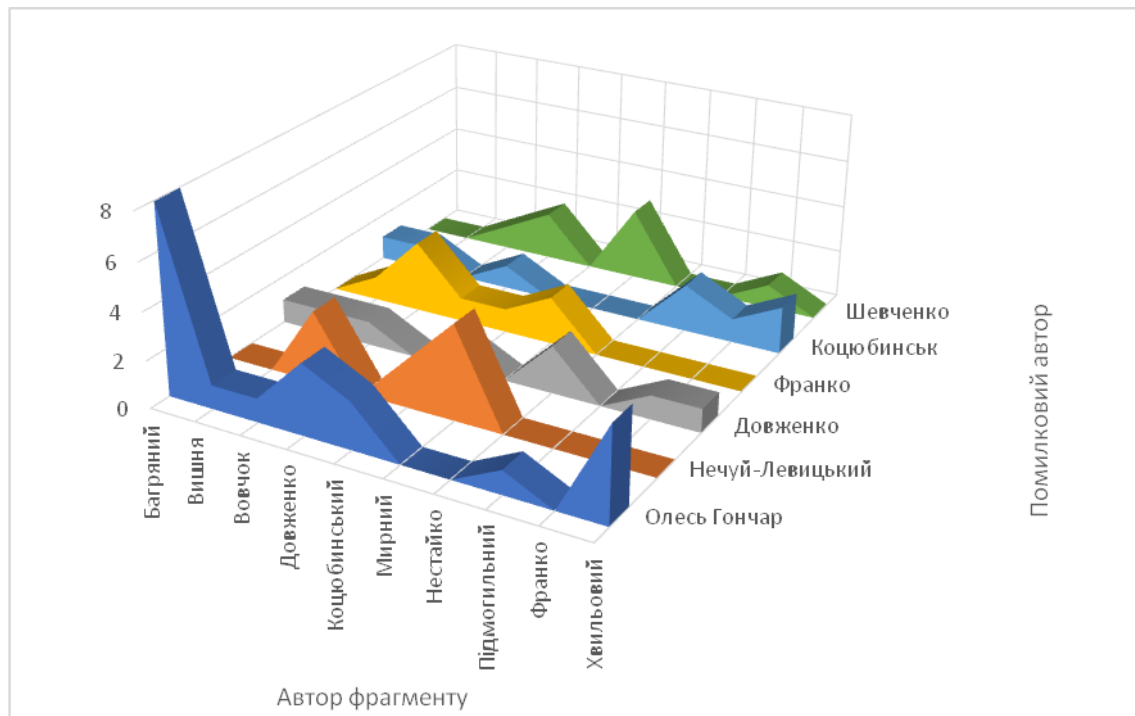


Рисунок 4 – Кількість помилкового визначення авторства фрагменту

Автори Олександр Довженко, Іван Багряний, Микола Хвильовий жили в один час з Олесем Гончаром, працювали в один історичний період. Вони всі були видатними представниками української літератури, кожен зі своїм стилем і внеском. Саме до авторства Олесь Гончара були віднесені помилково фрагменти з творів декількох авторів. Олесь Гончар був одним з тих письменників, які відомі, завдяки сильному глобальному мовному мисленню, мав потужний неолексикон. Його твори відображають вагу його лексики в розвитку сучасної української мови. Олександр Довженко, Іван Багряний та Микола Хвильовий також були видатними представниками української літератури, їхні твори теж відображають різні аспекти української культури, історії та суспільства. Всі вони використовували сучасну українську мову як засіб для створення потужних образів і висловлювання глибоких думок.

Висновки. Встановлено, що довгі фрагменти дозволяють з високою точністю визначити автора художнього україномовного тексту.

Використання систем рейтингування та генеративних мовних моделей може допомогти у визначенні авторства тексту, проте цей метод не є абсолютно точним. Системи рейтингування дозволяють здійснювати пошук інформації про текст та його автора, але результати такого пошуку можуть бути неточні-

ми або неповними. З іншого боку, генеративні мовні моделі можуть створювати текст, який імітує стиль певного автора, але не є його авторським твором.

Висунута гіпотеза щодо ефективності штучного інтелекту зі встановлення авторства творів не виправдалась. Штучний інтелект має дещо нижчу ефективність ніж очікувалось, що опосередковано викриває засоби його роботи. А саме те що при встановленні авторства не виконується послідовне дослідне співставлення запропонованого фрагменту з банком творів що поширені в інтернет середовищі. Він загалом базується на знаннях отриманих під час глибинного навчання. Слід очікувати ще нижчу ефективність встановлення авторства текстів менш відомих та плодovitих авторів.

ЛІТЕРАТУРА

1. Chowdhery A. PaLM: Scaling Language Modeling with Pathways, Aakanksha Chowdhery, Sharan Narang, Jacob Devlin, Maarten Bosma, Gaurav Mishra, Adam Roberts.// <https://doi.org/10.48550/arXiv.2204.02311> – 2022.
2. An unofficial Python library for interacting with the Skype HTTP API. [Електронний ресурс] – Режим доступу: <https://pypi.org/project/SkPy/>.
3. Bengio Y. «Neural net language models». Scholarpedia. Dr. Yoshua Bengio 3. doi:10.4249/scholarpedia.3881. – 2008 – P. 3881.
4. Perez E. True few-shot learning with language models. Advances in Neural Information Processing Systems, Ethan Perez, Douwe Kiela, and Kyunghyun Cho. – 2021 – P. 34.
5. Google Bard: How to Use Google's AI Chatbot - How-To Geek. [Електронний ресурс] – Режим доступу: <https://www.howtogeek.com/880668/google-bard-how-to-use-googles-ai-chatbot/>.
6. How Bing AI Search Uses Website Content - Search Engine Journal. [Електронний ресурс] – Режим доступу: <https://www.searchenginejournal.com/how-bing-ai-search-uses-web-content/480643/>.
7. Microsoft's next-gen Bing uses a 'much more powerful' language model than ChatGPT [Електронний ресурс] – Режим доступу:
8. <https://www.engadget.com/microsofts-next-gen-bing-more-powerful-language-model-than-chatgpt-182647588.html>.
9. Reinventing search with a new AI-powered Microsoft Bing and Edge, your copilot for the web. [Електронний ресурс] – Режим доступу:
10. <https://blogs.microsoft.com/blog/2023/02/07/reinventing-search-with-a-new-ai-powered-microsoft-bing-and-edge-your-copilot-for-the-web/>.

11. Іванов О.П. Визначення авторства художнього українського тексту засобами штучного інтелекту за надкороткими уривками О.П. Іванов, В.І. Шикаренко, В.В. Скалозуб, А.А. Косолапов // Наука та прогрес транспорту № 2(102) – 2023 – С. 45-53
12. Shynkarenko V. I. Methods and software for significant indicators determination of the natural language texts author profile V.I. Shynkarenko, I.M. Demydovych. // Prombles in programming/ 3 – 2023 – С. 22-29.
13. Shynkarenko V. I. A Dual Approach to Establishing the Authority of Technical Natural Language Texts and Their Components V. I. Shynkarenko, I. M. Demidovich, O. S. Kuropiatnyk // Science and Transport Progress № 2(102) – 2023 – С. 71-85.

REFERENCES

- Chowdhery A. PaLM: Scaling Language Modeling with Pathways, Aakanksha Chowdhery, Sharan Narang, Jacob Devlin, Maarten Bosma, Gaurav Mishra, Adam Roberts.// <https://doi.org/10.48550/arXiv.2204.02311> – 2022.
- An unofficial Python library for interacting with the Skype HTTP API. [Electronic resource] – Access mode: <https://pypi.org/project/SkPy/>.
- Bengio Y. «Neural net language models». Scholarpedia. Dr. Yoshua Bengio 3. doi:10.4249/scholarpedia.3881. – 2008 – P. 3881.
- Perez E. True few-shot learning with language models. Advances in Neural Information Processing Systems, Ethan Perez, Douwe Kiela, and Kyunghyun Cho. – 2021 – P. 34.
- Google Bard: How to Use Google's AI Chatbot - How-To Geek. [Electronic resource] – Access mode: <https://www.howtogeek.com/880668/google-bard-how-to-use-googles-ai-chatbot/>.
- How Bing AI Search Uses Website Content - Search Engine Journal. [Electronic resource] – Access mode: <https://www.searchenginejournal.com/how-bing-ai-search-uses-web-content/480643/>.
- Microsoft's next-gen Bing uses a 'much more powerful' language model than ChatGPT [Electronic resource] – Access mode: <https://www.engadget.com/microsofts-next-gen-bing-more-powerful-language-model-than-chatgpt-182647588.html>.
- Reinventing search with a new AI-powered Microsoft Bing and Edge, your copilot for the web. [Electronic resource] – Access mode: <https://blogs.microsoft.com/blog/2023/02/07/reinventing-search-with-a-new-ai-powered-microsoft-bing-and-edge-your-copilot-for-the-web/>.

- Ivanov O. P. Determining the Authorship of a Ukrainian-Language Literary Text by Means of Artificial Intelligence from Ultra-Short Excerpts O. P. Ivanov, V.I. Shynkarenko, V.V. Skalozub, A.A. Kosolapov. Science and Transport Progress № 2(102) – 2023 – P. 45-53.
- Shynkarenko V.I. Methods and software for significant indicators determination of the natural language texts author profile V.I. Shynkarenko, I.M. Demydovych. // Prombles in programming/ 3 – 2023 – C. 22-29.
- Shynkarenko V.I. A Dual Approach to Establishing the Authority of Technical Natural Language Texts and Their Components V. I. Shynkarenko, I. M. Demidovich, O. S. Kuropiatnyk // Science and Transport Progress № 2(102) – 2023 – C. 71-85.

Received 08.11.2023.

Accepted 15.11.2023.

***Determining the ability of artificial intelligence to establish authorship
of artistic ukrainian texts using significant fragments***

Artificial intelligence is becoming an integral part of everyday life and professional activity of a person. Bing, as an intelligent search system, can serve as a tool for determining the authorship of artistic text in Ukrainian. Bing helps to uncover information about a text fragment and its author, although the search results may be inaccurate or incomplete.

This work aims to explore the effectiveness of determining the authorship of artistic texts using modern artificial intelligence tools based on significant fragments of works. Ten Ukrainian authors with a rich body of artistic works, reflecting various aspects of Ukrainian culture and history, were selected for the experiment. Random fragments up to 500 words in length were selected from various works of these authors. An experiment was conducted to determine the authorship of 360 fragments.

Using the Python programming language and the skpy package, software was created that sends queries and receives responses from the Bing bot embedded in Microsoft Skype. The presence of the author's name and the corresponding title of the work were checked in the response texts.

This work introduces, for the first time, a method of verifying the authorship of Ukrainian-language text fragments using the Bing bot equipped with artificial intelligence. A comparative analysis was conducted and experiments were carried out to identify the authorship of significant long fragments.

It was found that long fragments allow the author of the artistic Ukrainian text to be determined with high accuracy. Ivan Franko has the highest percentage of responses where the author's name and the title of the work were mentioned - 87%.

The proposed hypothesis regarding the effectiveness of artificial intelligence in establishing authorship of works has not been confirmed. Artificial intelligence has slightly lower efficiency than expected, which indirectly exposes its means of operation. Namely, when establishing authorship, a sequential research comparison of the proposed fragment with a bank of works that are widespread in the Internet environment is not performed.

Іванов Олександр Петрович – доцент, к.т.н., доцент кафедри комп’ютерних інформаційних технологій, Український державний університет науки і технологій.

Скалозуб Владислав Васильович – професор, д.т.н., професор кафедри комп’ютерних інформаційних технологій, Український державний університет науки і технологій.

Горячкін Вадим Миколайович – доцент, к.т.н., завідувач кафедри комп’ютерних інформаційних технологій, Український державний університет науки і технологій.

Шинкаренко Віктор Іванович – професор, д.т.н., професор кафедри комп’ютерних інформаційних технологій, Український державний університет науки і технологій.

Ivanov Oleksandr – candidate of Technical Sciences, Associate Professor of the Computer Information Technologies Department, Ukrainian State University of Science and Technology.

Skalozub Vladislav – Doctor of Technical Sciences, professor, Professor of the Computer Information Technologies Department, Ukrainian State University of Science and Technology.

Horiachkin Vadym – candidate of Technical Sciences, Head of Department Computer Information Technologies, Ukrainian State University of Science and Technology.

Shynkarenko Viktor – Doctor of Technical Sciences, professor, Professor of the Computer Information Technologies Department, Ukrainian State University of Science and Technology.

А.О. Щербаков, Т.А. Ліхоузова

ПРОГРАМНЕ ЗАБЕЗПЕЧЕННЯ РОЗУМНОГО ГОДИННИКА ДЛЯ САМОСТІЙНИХ ЗАНЯТЬ СПОРТОМ ТА ФІТНЕСОМ

Анотація. За наявною статистикою, найбільш популярним розумним годинником серед користувачів є Apple Watch. На сьогодні існує більше 100 мільйонів унікальних користувачів цього девайсу, 75% з яких використовують його для занять спортом. Маючи широкий набір різноманітних датчиків для відстеження фізичних параметрів користувача, ані компанія Apple, ані сторонні розробники поки не розробили програмного забезпечення для систематизації усіх зібраних даних задля покращення фізичних параметрів спортсмена та досягнення особистих спортивних цілей.

Метою дослідження є пошук можливості покращення фізичних параметрів починаючого спортсмена за рахунок комплексного аналізу зібраних розумним годинником даних його активності та створенні на основі цих даних більш персоналізованих рекомендацій.

Запропоновано програмне забезпечення – застосунок для розумних годинників, який аналізує зібрану статистику тренувань та на її основі надає рекомендації по проведенню поточного тренування.

Ключові слова: wearable devices, Xcode, watches, health and fitness, HealthKit.

Смарт-годинники та фітнес-браслети стали не просто модним аксесуаром, а невід'ємною частиною повсякденного життя багатьох людей. Більшість з них можуть здійснювати моніторинг тренувань, збираючи інформацію з різних датчиків, таких як GPS, пульсометр, акселерометр [1]. Однак не всі такі пристрої можуть надавати користувачеві докладну аналітику та рекомендації з покращення тренування, що є важливим для досягнення успіху в спорті.

Метою дослідження є пошук можливості покращення фізичних параметрів починаючого спортсмена за рахунок комплексного аналізу зібраних розумним годинником даних його активності та створенні на основі цих даних більш персоналізованих рекомендацій під час тренування.

Аналіз останніх досліджень і публікацій. Для Apple Watch доступно багато програм для фітнесу, кожна з яких має свої унікальні функції та особливості. Наведемо деякі популярні програми для фітнесу та їх особливості.

Apple Fitness [2] – власна програма для фітнесу від Apple, розроблена спеціально для Apple Watch. Застосунок пропонує широкий спектр тренувань, включаючи біг, йогу, їзду на велосипеді та танці, а також показує показники фізичної активності.

MyFitnessPal [3] – популярна програма для відстеження калорій і фізичних вправ. MyFitnessPal інтегрується з Apple Watch, щоб відстежувати щоденні кроки та тренування. Застосунок також пропонує харчовий щоденник, де ви можете записувати свої прийоми їжі та відстежувати споживання калорій.

Strava [4] – застосунок для соціального фітнесу, який дозволяє вам відстежувати свої тренування та ділитися ними з друзями. Він також пропонує виклики та змагання, які допоможуть мотивувати вас досягти ваших фітнес-цілей.

Nike Training Club [5] – фітнес-застосунок від Nike, який пропонує різноманітні тренування, зокрема силові вправи, циклічні та йогу. Він надає персоналізовані плани тренувань і дозволяє відстежувати ваш прогрес.

FitOn [6] – застосунок пропонує набір вправ та програму тренувань для зменшення ваги, занять йогою та виконання повсякденних занять спортом.

Загалом відмінності між цими фітнес-застосунками зводяться до їхніх особливостей і типів тренувань, які вони пропонують. Деякі програми зосереджені на соціальних функціях і змаганнях, а інші пропонують персоналізовані плани тренувань або зосереджені на певних типах тренувань, як-от їзда на велосипеді чи йога.

Однією з ключових проблем фітнес-застосунків для смарт годинників є недостатня обробка та аналіз зібраних даних безпосередньо під час тренування. У більшості випадків, користувачі самостійно повинні оцінювати показники та здійснювати аналіз даних, що може бути складним для новачків у спорті.

Існують виробники розумних годинників, які більше орієнтовані на девайси для занять спортом, такі як Garmin [7] чи Suunto. Garmin у своїх девайсах пропонує функцію «Daily Suggested Workouts», яка запропонує спортсмену темп та дистанцію для його наступного тренування. Ця функція є унікальною для GarminOS, що унеможливорює її аналіз та використання на інших платформах.

На жаль, жоден проаналізований застосунок не надає достатньої інформації щодо правильного виконання тренувань та збору показників, що є недоліком для користувачів, які хочуть займатися спортом без ризиків для свого здоров'я.

Результати та основний матеріал дослідження. Однією з головних проблем розробки програмного забезпечення для розумних годинників є відсутність широких можливостей для монетизації. Це призвело до відсутності інтересу великих компаній та команд до такого виду програмного забезпечення. Також проблемами є невеликий розмір екрану та обмежена обчислювальна потужність, порівняно з настільними або мобільними пристроями. Це означає, що розробники повинні оптимізувати як користувацький інтерфейс, так і алгоритми свого програмного продукту для підвищення швидкодії застосунку. Іншою проблемою є різноманітність ринку розумних годинників, де різні пристрої працюють на різних операційних системах і мають різні апаратні характеристики. Доводиться враховувати особливості та обмеження кожного пристрою та платформи та переконатися, що програми оптимізовані для кожного з них, або зосередитись на створенні своїх програмних продуктів тільки для обмеженого ряду пристроїв.

Через відсутність інструментів для розробки застосунків одразу для декількох платформ, було обрано шлях розробки застосунку тільки для однієї платформи – для WatchOS. Також однією з цілей є створення повністю автономного від смартфона застосунку. Для цього розроблений простий та мінімалістичний інтерфейс та простий алгоритм для аналізу даних тренування.

Єдиним доступним варіантом IDE для розробки застосунків для Apple Watch є XCode. Для розробки користувацького інтерфейсу є два популярних фреймворки – SwiftUI та UIKit. У проєкті використано SwiftUI та мову програмування Swift, бо вони пропонують більш сучасний і спрощений підхід до створення компонентів інтерфейсу користувача, ніж UIKit [8].

Для роботи з даними, які відносяться до фізичних характеристик та здоров'я користувача, компанія Apple пропонує використовувати фреймворк HealthKit [9], який надає розробникам набір інструментів для безпечного зберігання, керування та отримання даних про здоров'я та фізичну форму користувача. Це включає дані про фізичну активність, харчування, сон, пульс та артеріальний тиск. HealthKit також дозволяє користувачам ділитися своїми даними зі сторонніми додатками для аналізу фізичної активності [10].

Головною функцією запропонованого програмного забезпечення є формування порад спортсмену під час тренування, більше функцій можна побачити на рисунку 1.

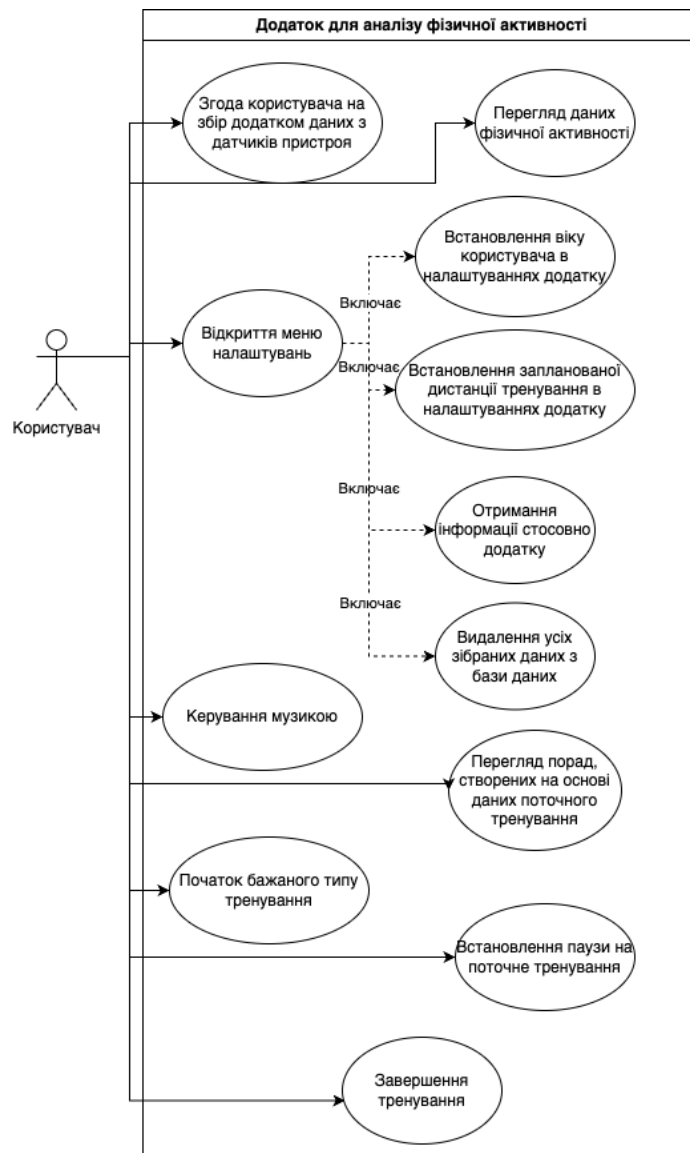


Рисунок 1 – Діаграма варіантів використання

Застосунок орієнтований на заняття бігом. Декілька перших тренувань буде збиратися статистика. Після старту активності застосунок почне виводити усю стандартну інформацію, як темп, пульс, час активності, пройденої відстань. Ці дані надходять з сенсорів та модулів розумного годинника. Коли базова статистика по тренуванням буде зібрана, в додатку почнуть відображатися поради щодо продовження тренування відповідно до спеціального алгоритму (рисунок 2).

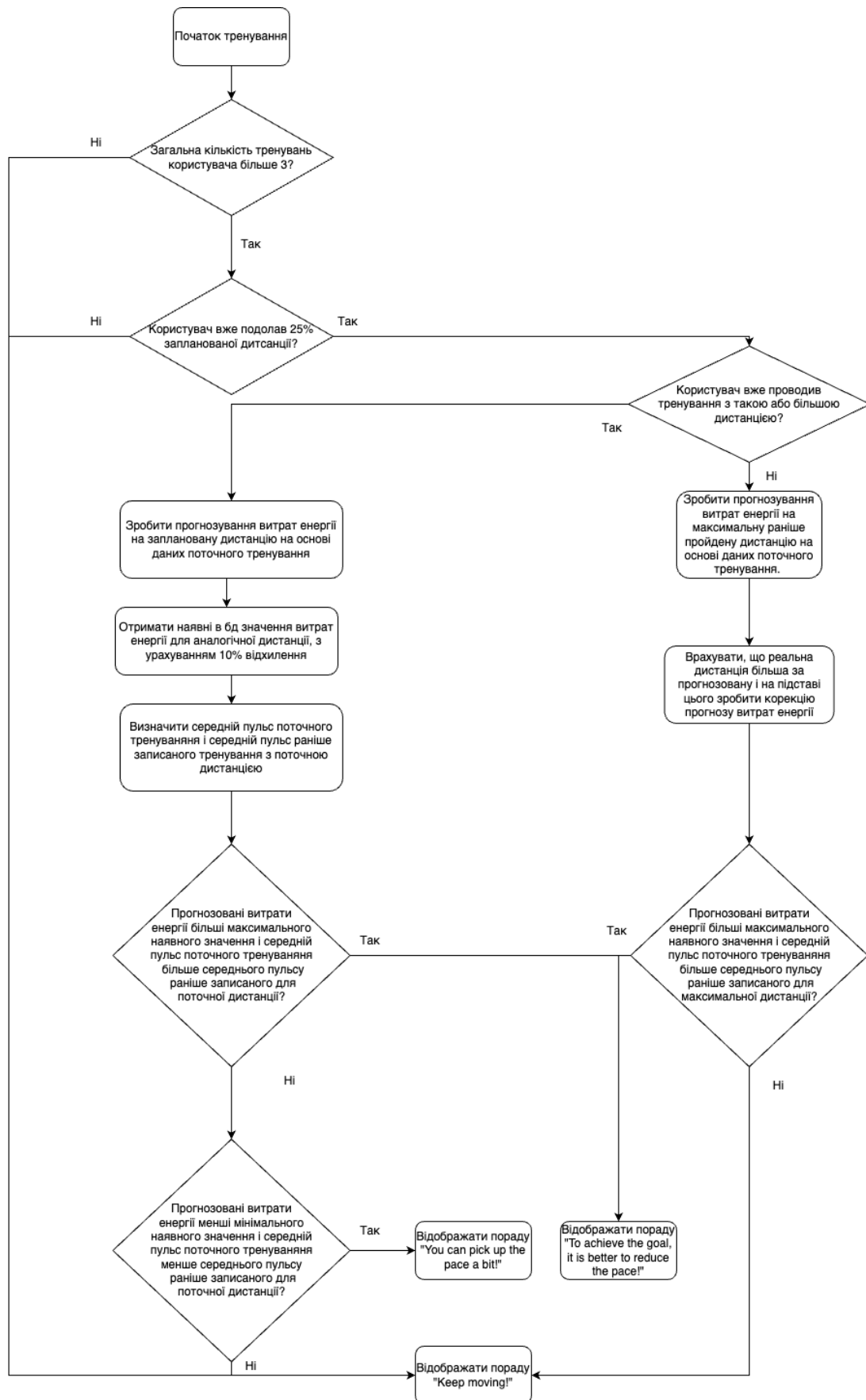


Рисунок 2 – Блок-схема алгоритму формування порад користувачу

Перед початком тренування користувач має встановити цільову дистанцію у метрах. Для аналізу витрат енергії на досягнення цілі алгоритм буде прогнозувати витрачену енергію у кілокалорія на основі даних поточного тренування. Через те, що данні і витрат енергії і дистанції поступово збільшуються, а спортсмени завжди намагаються тримати рівний темп або плавно його змінювати, для прогнозування витрат енергії добре підійде алгоритм лінійної регресії. Після отримання прогнозованого значення, буде виконано порівняння цього значення з вже наявними значеннями по раніше проведеним тренуванням. Наприклад, якщо раніше на 5 км дистанції спортсмен витрачав 200 кілокалорій, а зараз за прогнозом витратить 300, застосунок йому порадить зменшити темп для досягнення цілі.

Також можлива ситуація, коли встановлена користувачем ціль буде більше, ніж максимальна дистанція раніше проведених тренувань. В такому разі застосунок прорахує витрату калорій для вже відомої дистанції і буде показувати поради, орієнтуючись на це число, але з корегуванням, що реальна запланована дистанція буде більше.

Застосунок також буде попереджати про досягнення максимального рекомендованого пульсу. Максимальний рекомендований пульс можливо поррахувати з різниці числа 220 та віку користувача. Хоча кожен організм має індивідуальні показники, дослідження показують лінійну залежність між віком та максимальним пульсом людини. Слід зазначити, що дослідники не радять спиратися на цю формулу, рахуючи максимальний пульс дітей до 10 років [9].

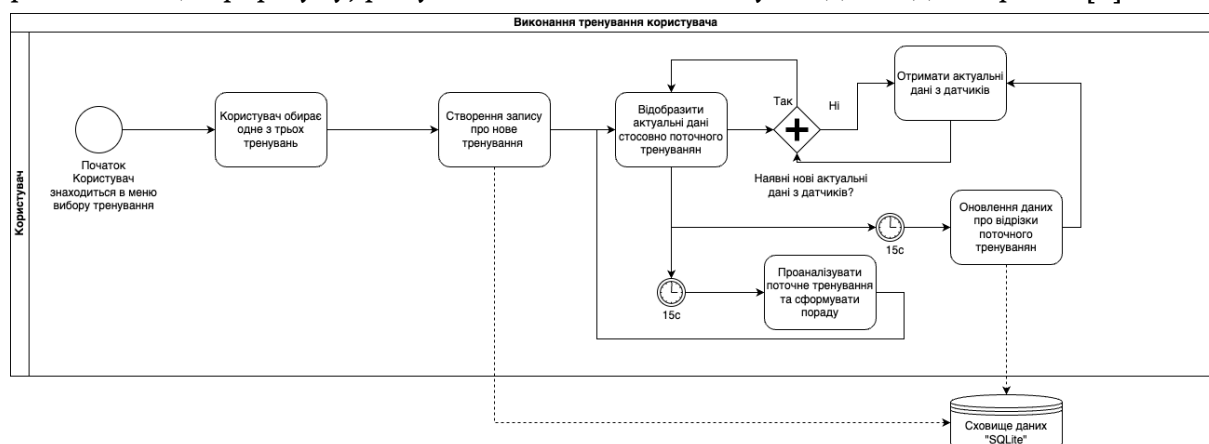


Рисунок 3 – Схема бізнес-процесу виконання тренування користувачем

Архітектура всього застосунку розподілена на 3 блоки: Model, ViewModel та View, як і в самому дизайн патерні MVVM. У блоці View односторонні стрілки означають, що перехід між представленнями можливий тільки в одну сто-

рону, через натискання кнопки. Двостороння стрілка означає, що можливий перехід між представленнями в обидва боки під час користування застосунком.

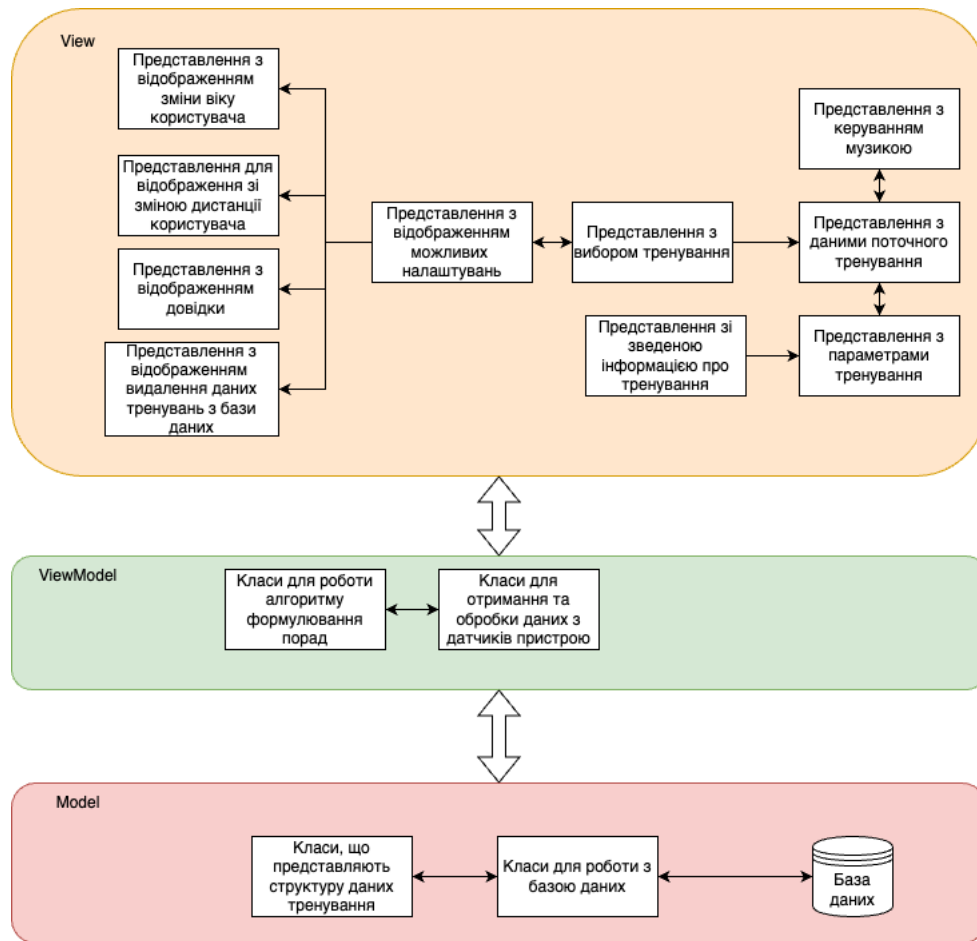


Рисунок 4 – Архітектура застосунку

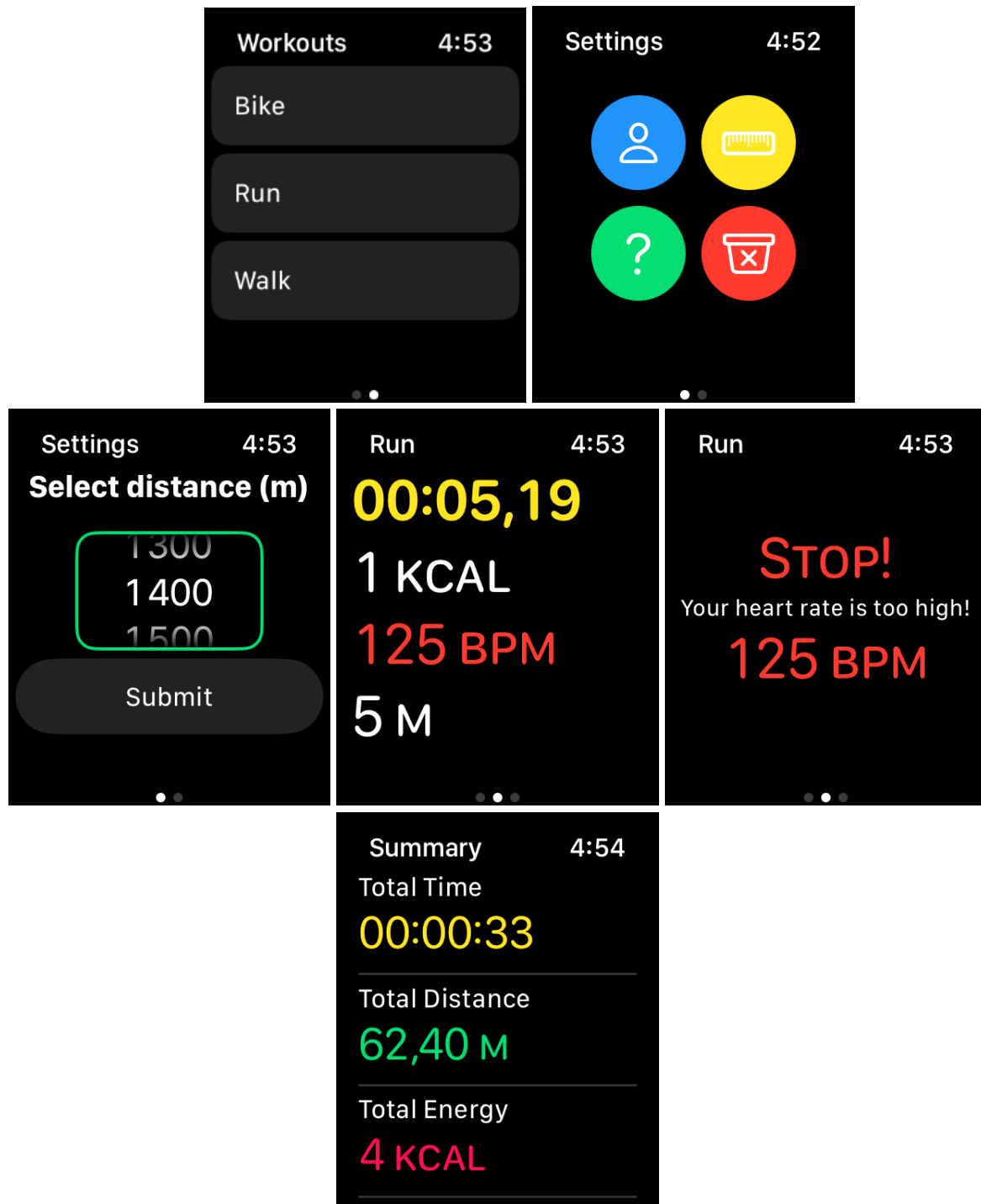


Рисунок 5 – Приклади екранних форм за стосунку

Висновки. Незважаючи на широкий асортимент застосунків для відстеження та запису тренувань користувача, жоден не аналізує дані безпосередньо у процесі тренування та не формує відповідні поради та рекомендації. Усі наведені застосунки нерозривно пов'язані з застосунком на смартфоні. Ключовою відмінністю запропонованого застосунка є можливість аналізувати дані та

формувати поради та застереження для користувача безпосередньо у процесі тренування, коли вони найбільше потрібні.

ЛІТЕРАТУРА / REFERENCES

- Gregersen E. Smartwatch [Electronic resource] / Erik Gregersen – Access mode: <https://www.britannica.com/technology/smartwatch>.
- Apple Fitness [Electronic resource]. Access mode: <https://www.apple.com/apple-fitness-plus/>
- MyFitnessPal [Electronic resource]. Access mode: <https://www.myfitnesspal.com>
- Strava [Electronic resource]. Access mode: <https://www.strava.com>
- Nike Training Club [Electronic resource]. Access mode: <https://www.nike.com/ntc-app>
- FitOn [Electronic resource]. Access mode: <https://fitonapp.com/>
- Garmin Daily Suggested Workouts [Electronic resource].
- SwiftUI [Electronic resource]. Access mode: <https://developer.apple.com/XCode/swiftui/>.
- HealthKit [Electronic resource]. Access mode: <https://developer.apple.com/documentation/healthkit>.
- Robergs R. THE SURPRISING HISTORY OF THE “HR_{max}=220-age” / R. Robergs, R. Landwehr. – 2002. – С. 2–4.

Received 12.11.2023.
Accepted 15.11.2023.

Smart watch software for independent sports and fitness

According to available statistics, the most popular smart watch among users is the Apple Watch. Today, there are more than 100 million unique users of this device, 75% of whom use it for sports. With a wide array of different sensors to track a user's physical parameters, neither Apple nor third-party developers have yet developed software to systematize all the collected data to improve an athlete's physical parameters and achieve personal athletic goals.

The purpose of the research is to find the possibility of improving the physical parameters of a novice athlete by means of a comprehensive analysis of his activity data collected by a smart watch and creating more personalized recommendations during training based on this data.

There are many fitness apps available for the Apple Watch, each with its own unique features and features. Unfortunately, none of the analyzed applications provide sufficient information regarding the correct execution of training and the collection of in-

dicators, which is a disadvantage for users who want to do sports without risks to their health.

One of the main challenges of software development for smartwatches is the lack of extensive monetization opportunities. This led to the lack of interest of large companies and teams in this type of software. Small screen size and limited processing power compared to desktop or mobile devices are also issues. This means that developers must optimize both the user interface and the algorithms of their software product to increase the speed of the application. Another challenge is the diversity of the smartwatch market, where different devices run on different operating systems and have different hardware specifications. You have to consider the specifics and limitations of each device and platform and make sure that your apps are optimized for each of them, or focus on building your software products for only a limited number of devices.

Due to the lack of tools for developing applications for several platforms at once, the path of developing an application for only one platform - for WatchOS - was chosen. Also, one of the goals is to create an application that is completely autonomous from a smartphone. For this, a simple and minimalistic interface and a simple algorithm for analyzing training data have been developed.

The proposed application is focused on running. Statistics will be collected during the first few training sessions. After starting the activity, the application will start displaying all the standard information, such as pace, heart rate, activity time, distance traveled. This data comes from the sensors and modules of the smart watch. When basic training statistics are collected, the app will begin to display tips on how to continue training.

Keywords: wearable devices, Xcode, watches, health and fitness, HealthKit.

Ліхоузова Тетяна Анатоліївна – к.т.н., доцент кафедри інформатики та програмної інженерії, Національний технічний університет України «Київський політехнічний інститут імені Ігоря Сікорського».

Щербаков Антон Олександрович – студент кафедри інформатики та програмної інженерії, Національний технічний університет України «Київський політехнічний інститут імені Ігоря Сікорського».

Likhouzova Tetiana – PhD, associate professor, Department of Informatics and Software Engineering National Technical University of Ukraine «Igor Sikorsky Kyiv Polytechnic Institute».

Shcherbakov Anton – student, Department of Informatics and Software Engineering National Technical University of Ukraine «Igor Sikorsky Kyiv Polytechnic Institute».

Л.О. Кіріченко, Д.С. Хацько, П.П. Зінченко

КЛАСТЕРИЗАЦІЯ ТРАЄКТОРІЙ БРОУНІВСЬКОГО РУХУ ЗА ДОПОМОГОЮ МАШИННОГО НАВЧАННЯ

Анотація. Стаття присвячена виявленню пасток, які захоплюють броунівську частинку, за допомогою машинного навчання. Проведено моделювання траєкторії броунівської частинки за допомогою моделі броунівського руху із дрейфом, що охоплює як вільну дифузію, так і рух у пастці. Для кластеризації часових даних використовуються метод кластеризації на основі щільності DBSCAN. Універсальність цього метода дозволяє ідентифікувати кластери без попереднього знання їхньої кількості або форми, що робить його придатним для виявлення пасток. Проведене дослідження показує, що застосування методу DBSCAN дозволяє досягнути середньої точності 95.0%.

Ключові слова: Броунівський рух з дрейфом, рух частинок у пастках, кластеризація за допомогою машинного навчання, DBSCAN

Вступ і мета

Броунівський рух виступає математичною моделлю для багатьох фізичних, хімічних і біологічних явищ [1-3]. Рух молекул у клітинах, процеси дифузії та взаємодії між біологічними компонентами є фундаментальними процесами, які визначають біологічні функції та реакції організмів. Вивчення молекулярної динаміки в клітинах за допомогою броунівського руху відкриває нові шляхи для розуміння клітинних процесів і механізмів, які впливають на функціональність живих організмів.

У контексті траєкторії броунівського руху пастками називають просторові області, в яких броунівська частинка утримується протягом певного часу. Ці пастки можуть виникати природним чином або можуть бути штучно створені в лабораторіях для детального вивчення властивостей частинок [4-6]. Дослідження динаміки броунівського руху в системах з пастками сприяє глибшому розумінню цього руху та виявленню складних закономірностей, що лежать в його основі. Виявлення пасток є окремим і складним завданням, яке було предметом різних досліджень із застосуванням як ансамблевих підходів [6,7], так і аналізу траєкторій окремих частинок [8-11]. Багато з цих методів використову-

ють ковзаючі вікна для виявлення пасток [10-12], хоча деякі підходи передбачають цілісне дослідження траєкторій частинок [8].

Для вирішення проблеми виявлення пасток було застосовані методи машинного навчання, як показано в [12], проте такі дослідження все ще перебувають на стадії розвитку. У нашій роботі ми пропонуємо застосування методу кластеризації часових рядів на основі методів машинного навчання [13-17]. Тому основною метою цього дослідження є кластеризація часової траєкторії броунівської частинки для виявлення пасток – областей, де рух частинки знає значного сповільнення.

Для виконання завдання кластеризації було обрано метод DBSCAN [17-20]. Цей метод спирається на розподіл щільності об'єктів даних у просторі, що дозволяє ідентифікувати області з високою щільністю, які можуть відповідати кластерам. Крім того, метод не залежить від апріорних знань про кількість і форму кластерів, що дозволяє ідентифікувати пастки без попередньої інформації про їхню кількість і часові інтервали.

Моделювання траєкторій частинок

Для моделювання траєкторії руху частинки було обрано модель броунівського руху з дрейфом. Модель броунівського руху з дрейфом відрізняється від стандартного броунівського руху тим, що вводить поняття дрейфу або середнього спрямованого зміщення частинки. У стандартному броунівському русі частинка випадковим чином рухається в різних напрямках без будь-якого переважного напрямку. Однак у моделі дрейфу на частинку також впливають зовнішні сили або градієнти, які надають певного напрямку її руху. Модель броунівського руху з дрейфом широко використовується для моделювання реальних процесів у різних галузях. Ці процеси охоплюють різні галузі, такі як фізика, хімія, біологія, фінанси та інженерія. Універсальність моделі броунівського руху з дрейфом робить її безцінним інструментом для розуміння складності динамічних систем та явищ [21, 22]. Наведемо означення як для стандартного броунівського руху, так і для броунівського руху з дрейфом.

Стохастичний процес $\{W(t), t \geq 0\}$ називається стандартним броунівським рухом, якщо він задовольняє наступним умовам: $W(0) = 0$; $W(t) - W(s)$ відповідає нормальному розподілу із середнім значенням 0 та дисперсією $t - s$; $W(t) - W(s)$ не залежить від $W(v) - W(u)$, якщо ці часові інтервали не перетинаються.

Броунівський рух з дрейфом $\{B(t), t \geq 0\}$ розширює поняття стандартного броунівського руху, додаючи до нього дрейфовий член. Цей дрейф μ вносить систематичний зсув або спрямовану складову в рух частинки на додаток до випадкових флуктуацій. Процес $\{B(t), t \geq 0\}$ задовольняє наступним умовам: $B(0) = 0$; має незалежні та стаціонарні прирости; процес нормально розподілений $B(t) \sim N(\mu t, t)$.

Процес броунівського руху з дрейфом, що включає обидва дрейфи μ і коефіцієнт дифузії σ , в загальному вигляді можна виразити так:

$$B_{(\mu, \sigma)}(t) = \mu t + \sigma W(t). \quad (1)$$

Для моделювання траєкторії броунівської частинки в середовищі з пастками інтервали руху частинки в пастках вибираються випадковим чином, з урахуванням конкретних обмежень, що диктуються конкретною задачею, яка розглядається. Рух частинки в пастці також моделювався за допомогою броунівського руху з дрейфом, але з істотним зменшенням швидкості частинки, що досягалося за рахунок зменшення коефіцієнта дифузії в M раз:

$$B_{(\mu, \sigma)}(t) = \mu t + \frac{\sigma}{M} W(t). \quad (2)$$

Таким чином, для моделювання траєкторії частинки використовували (1) для опису руху частинки під час вільної дифузії, а (2) – для опису її руху в межах пастки. Такий підхід дозволив отримати повне уявлення про поведінку частинки, враховуючи як необмежений броунівський рух, так і вплив пастки.

Застосування методу кластеризації DBSCAN для виявлення пасток

Використання методу кластеризації на основі щільності показало багатообіцяючий потенціал у сфері кластеризації часових рядів. У нашому дослідженні ми використали можливості методу DBSCAN (Density-Based Spatial Clustering of Applications with Noise) для ефективної кластеризації часових даних.

У нашому дослідженні застосування DBSCAN дозволило виявити складні часові патерни і структури всередині траєкторій, що поглибило наше розуміння основної динаміки системи.

Метод DBSCAN визначає кластери як області з високою щільністю, відокремлені від областей з низькою щільністю. DBSCAN здатний виявляти кластери різної форми та розміру, що робить його особливо ефективним, коли кластери мають різну щільність або складну форму.

DBSCAN використовує два основні параметри: радіус і мінімальну кількість сусідів. Починаючи з випадкової точки з набору даних, метод розглядає

всі точки в заданому радіусі від цієї точки. Якщо кількість точок у цій області перевищує мінімальну кількість сусідів, область вважається щільною і розглядається як один кластер. Потім цей процес поширюється на сусідні точки, які також задовольняють умові щільності, і так далі. Точки, які не належать до жодного кластера і не відповідають критерію щільності, вважаються шумом [13, 23].

У контексті представленого дослідження всі точки траєкторії були класифіковані як такі, що належать до одного з двох кластерів. Перший кластер охоплював точки, динаміка яких відповідала броунівському руху, тоді як другий кластер складався з точок, що демонстрували рух у пастці.

Пов'язуючи певні точки траєкторії з кластером пастки, можна було розшифрувати динаміку і тривалість утримання частинок у різних ділянках експериментальної установки. Ця часова інформація дозволила отримати повне уявлення про явище пастки та його зміни в часі.

Експеримент і результати

Дослідження фокусується на двовимірному броунівському русі, проте кластеризація проводилася незалежно для кожної координати, щоб підвищити точність виявлення пасток. Часова траєкторія моделювалася згідно з рівняннями (1) і (2). Кількість пасток, їхня тривалість та часове розташування визначалися випадковим чином, з використанням рівномірного розподілу, з дотриманням попередньо визначених обмежень. Наприклад, були накладені обмеження, щоб запобігти перекриттю позицій пасток у часовому просторі.

На рис. 1 вгорі показано ілюстрацію змодельованого двовимірного броунівського руху у фазовому просторі X та Y . Області, що відповідають пасткам, для наочності виділено червоним кольором. У нижній частині рис. 1 подано проєкції цього руху на осі X та Y , що дає повне зображення траєкторії частинки в обох координатах.



Рисунок 1 - Двовимірний броунівський рух у фазовому просторі (вгорі)
та проекції на осі X і Y (внизу)

Вхідними даними для кластеризації були модельні траєкторії з пастками, а вихідними – анотовані відрізки цієї траєкторії, пов'язані з пастками. Подальший кластерний аналіз мав на меті розбити цю траєкторію на окремі кластери, які відповідали різній динаміці, що її демонструвала частинка.

У процесі моделювання певні значення параметрів моделі використовувалися як базова лінія: $\mu = 0.1$, $\sigma = 20$ та $M = 10$. Крім того, було проведено серію експериментів, в яких значення цих параметрів систематично змінювалися.

Оскільки ми використовували реалізацію на основі моделей, ми змогли оцінити точність процесу кластеризації. Для оцінки точності роботи методів були використані матриці помилок для бінарної класифікації. На рис. 2 представлено одну з типових матриць помилок для проведеного експерименту. В даному випадку точність склала 96%. Порівнюючи передбачені кластерні розподіли з істинною природою точок траєкторії (класифіковані як такі, що належать до пастки чи ні), ці матриці дозволили нам кількісно оцінити успіх алгоритму кластеризації в правильній ідентифікації точок траєкторії, що належать і не належать до пастки.

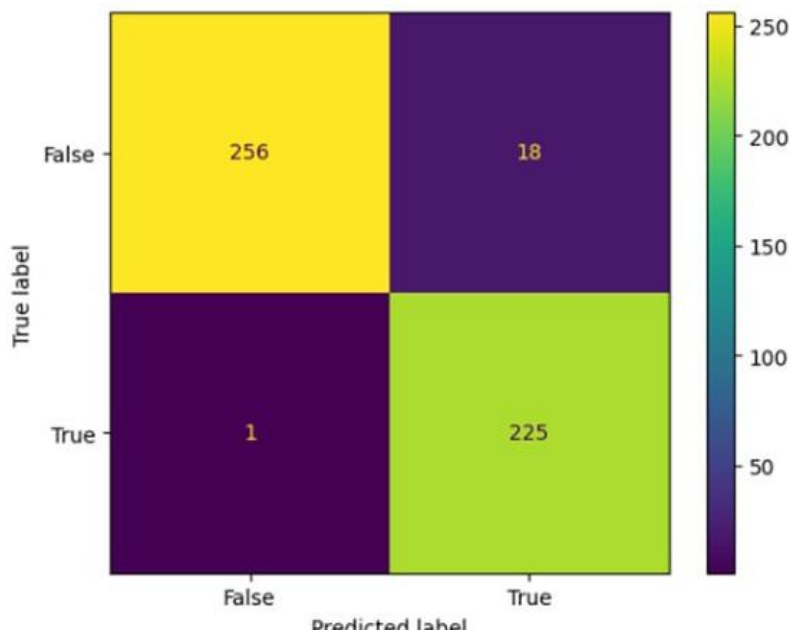


Рисунок 2 - Матриця помилок одного з результатів кластеризації

Для побудови моделі броунівського руху та її кластеризації була використана мова програмування Python, яка відома своєю великою колекцією бібліотек, пристосованих для математичних задач, машинного навчання та візуалізації результатів. Універсальність Python у наукових обчисленнях зробила його ідеальним вибором для нашого дослідження, дозволивши нам ефективно реалізувати та проаналізувати складну динаміку броунівських траєкторій руху.

На рис. 3 представлено ілюстрацію одного з остаточних результатів класифікації, отриманих за допомогою методу DBSCAN. Виділені області зображують окремі кластери, що представляють ідентифіковані пастки в даних траєкторії.

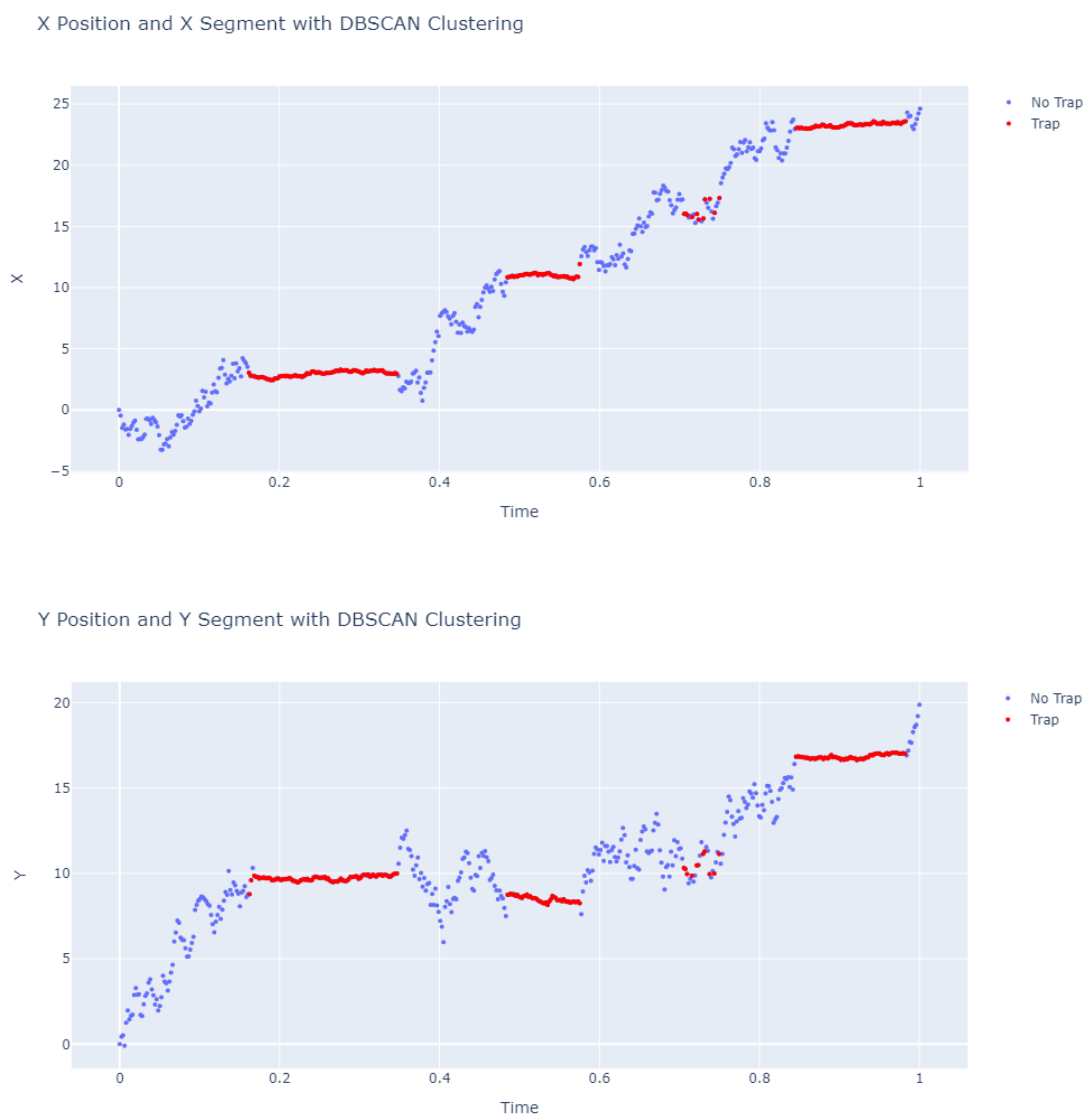


Рисунок 3 - Результати кластеризації методом DBSCAN
для проекції на ось X (вгорі) та Y (внизу)

Варто зазначити, що метод DBSCAN продемонстрував неабияку точність у визначенні точок, які потрапляли в пастки на траєкторії. Тим не менш, він був схильний до випадкової помилкової класифікації, особливо в сценаріях, де різниця між вільною дифузією і рухом, пов'язаним з пастками, була менш чіткою. Зокрема, представлений метод є схильними до помилкової класифікації сегментів вільного руху як руху, пов'язаного з пасткою. Ця відмінність є дуже

важливою, оскільки вона відображає чутливість методу до складних просторових патернів і шуму в даних.

Для оцінки ефективності методу кластеризації була проведена комплексна оцінка за допомогою серії з 100 випробувань. Аналіз виявив помітні відмінності в показниках точності. Метод DBSCAN показав середню точність 95.0%.

Моделювання та кластеризація проводилися в діапазоні значень дрейфу μ , коефіцієнта дифузії σ та коефіцієнта m . Оцінка не виявила помітної залежності ефективності методів від дрейфу μ . Як показали експерименти, точність методів знижується лише при малих значеннях σ . При менших значеннях коефіцієнта m точність кластеризації зменшується. Це явище пояснюється тим, що для малих ділянок рух частинки в пастці наближається до її безперешкодного руху. Отже, відмінні риси утримання стають малопомітними, що створює проблему для методів кластеризації ефективно виявляти такі аномалії в русі частинок.

Висновки

У цій статті досліджено моделювання траєкторій частинок на основі броунівського руху з дрейфом і вивчено ефективність методу кластеризації DBSCAN на основі машинного навчання, з метою виявлення пасток. Модель броунівського руху охоплювала як вільну дифузію, так і рух у пастках. Використання методу DBSCAN продемонструвало його потенціал для точної ідентифікації областей пасток у даних траєкторії.

Дослідження показало, що DBSCAN продемонстрував загалом високу середню точність, зокрема, досягнув середньої точності 95.0%. Комплексні експерименти, проведені в рамках цього дослідження, дозволили зрозуміти вплив параметрів моделі на результати кластеризації.

Майбутні дослідження в цьому напрямку можуть заглибитися в застосування більш досконалих методів машинного навчання для виявлення пасток і дослідити інтеграцію методів багатовимірної кластеризації для фіксації складної поведінки частинок.

ЛІТЕРАТУРА

1. E. Frey and K. Kroy, "Brownian Motion: a Paradigm of Soft Matter and Biological Physics," Ann. d. Phys. (Leipzig), vol. 14, pp. 20-50, 2005. [Online]. Available: arXiv:cond-mat/0502602.

2. M.A. Pinsky and S. Karlin, Brownian Motion and Related Processes[Online]. Available:https://faculty.ksu.edu.sa/sites/default/files/an_introd_to_stoch_modeling_4th_ed.pdf.
3. L. Kirichenko and R. Lavrynenko, "Probabilistic Machine Learning Methods for Fractional Brownian Motion Time Series Forecasting," *Fractal and Fractional*, vol. 7, no. 7, pp. 517, 2023. DOI:10.3390/fractalfract7070517.
4. C. Bustamante, J. Liphardt, and F. Ritort, "The nonequilibrium thermodynamics of small systems," *Physics Today*, 2005. [Online]. Available: <https://arxiv.org/ftp/cond-mat/papers/0511/0511629.pdf>.
5. L. P. Faucheux et al., "Optical thermal ratchets," *Physical Review Letters*, vol. 74, no. 8, pp. 1504-1507, 1995. DOI: 10.1103/PhysRevLett.74.1504.
6. V. Briane et al., "A computational approach for detecting micro-domains and confinement domains in cells: A simulation study," *Phys. Biol.*, vol. 17, p. 025002, 2020. DOI 10.1088/1478-3975/ab5e1d.
7. V. Briane, V. Kervrann, and M. Vimond, "Statistical analysis of particle trajectories in living cells," *Phys. Rev. E*, vol. 97, p. 062121, 2018. DOI:10.1103/PhysRevE.97.062121.
8. Y. Lanoiselée et al., "Detecting Transient Trapping from a Single Trajectory: A Structural Approach," *Entropy*, vol. 23, p. 1044, 2021. DOI:0.3390/e23081044.
9. P. Koo and J. K. Mochrie, "Systems-level approach to uncovering diffusive states and their transitions from single-particle trajectories," *Phys. Rev. E*, vol. 94, p. 052412, 2016. DOI:10.1103/PhysRevE.94.052412.
10. N. Meilhac et al., "Detection of confinement and jumps in single-molecule membrane trajectories," *Phys. Rev. E*, vol. 73, p. 011915, 2006. DOI:10.1103/PhysRevE.73.011915.
11. A. R. Vega et al., "Multistep Track Segmentation and Motion Classification for Transient Mobility Analysis," *Biophys. J.*, vol. 114, pp. 1018–1025, 2018. DOI: 10.1016/j.bpj.2018.01.012.
12. P. Dosset et al., "Automatic detection of diffusion modes within biological membranes using back-propagation neural network," *BMC Bioinform.*, vol. 17, p. 197, 2016. DOI:10.1186/s12859-016-1064-z.
13. S. Aghabozorgi, A. Shirkhorshidi, and T. J. Wah, "Time-series clustering. A Decade Review," *Information systems*, vol. 53, pp. 16-38, 2015.
14. L. Kirichenko, O. Pichugina, and H. Zinchenko, "Clustering time series of complex dynamics by features," *CEUR Workshop Proceedings*, vol. 3132, pp. 83-93, 2022.

15. S. Khlamov et al., "Statistical modeling for the near-zero motion detection of objects in series of images from the data stream," in 12th International Conference on Advanced Computer Information Technologies (ACIT 2022), 2022, pp. 1-4. DOI: 10.1109/ACIT54803.2022.9913151.
16. L. Kirichenko, T. Radivilova, and A. Tkachenko, "Comparative analysis of noisy time series clustering," CEUR Workshop Proceedings, vol. 2362, 2019.
17. A. Alqahtani et al., "Deep Time-Series Clustering: A Review," Electronics, vol. 10, no. 23, p. 3001, 2021. DOI:10.3390/electronics10233001.
18. R. J. G. B. Campello et al., "Density-Based Clustering Based on Hierarchical Density Estimates," in Advances in Knowledge Discovery and Data Mining, PAKDD 2013, 2013, pp. 305-316. DOI:10.1007/978-3-642-37456-2_14.
19. L. Kirichenko et al., "Forecasting weakly correlated time series in tasks of electronic commerce," in Proceedings of the 12th International Scientific and Technical Conference on Computer Sciences and Information Technologies, CSIT 2017, 2017, vol. 1, pp. 309-312, doi: 10.1109/STC-CSIT.2017.8098793.
20. L. McInnes and J. Healy, "Accelerated Hierarchical Density Based Clustering," in 2017 IEEE International Conference on Data Mining Workshops (ICDMW), New Orleans, LA, USA, 2017, pp. 33-42, doi: 10.1109/ICDMW.2017.12.
21. S. Kadloor, R. S. Adve, and A. W. Eckford, "Molecular communication using Brownian motion with drift," IEEE Transactions on NanoBioscience, vol. 11, no. 2, pp. 89-99, 2012. DOI:10.1109/tnb.2012.2190546.
22. S. Zhu and C. Ravishankar, "A scalable approach to approximating aggregate queries over intermittent streams," in Proceedings. 16th International Conference on Scientific and Statistical Database Management, 2004., Santorini, Greece, 2004, pp. 85-94, doi: 10.1109/SSDM.2004.1311196.
23. S. Louhichi, M. Gzara and H. Ben Abdallah, "A density based algorithm for discovering clusters with varied density," 2014 World Congress on Computer Applications and Information Systems (WCCAIS), Hammamet, Tunisia, 2014, pp. 1-6, doi: 10.1109/WCCAIS.

Received 17.11.2023.

Accepted 21.11.2023.

Clustering Brownian motion trajectories using machine learning

The article is dedicated to detecting traps encountered by a Brownian particle based on machine learning methods. The trajectory of the Brownian particle was modeled using a drift-extended Brownian motion model, encompassing both free diffusion and particle movement within a trap. The density-based spatial clustering of applica-

tions with noise (DBSCAN) method was employed for clustering the motion trajectory. The versatility of this method allows the identification of clusters without prior knowledge of their quantity or shape, making it suitable for trap detection. The conducted research demonstrates that the application of the DBSCAN method achieves an average accuracy of 95.0%

Кіріченко Людмила Олегівна – д.т.н., професор кафедри прикладної математики Харківського національного університету радіоелектроніки.

Хацько Дарина Сергіївна – магістрантка кафедри прикладної математики Харківського національного університету радіоелектроніки.

Зінченко Петро Петрович – аспірант кафедри прикладної математики Харківського національного університету радіоелектроніки.

Kirichenko Lyudmila Olehivna - Ph.D., professor of the Department of Applied Mathematics of the Kharkiv National University of Radio Electronics.

Khatsko Daryna Serhiyivna - is a master's student at the Department of Applied Mathematics of the Kharkiv National University of Radio Electronics.

Zinchenko Petro Petrovych - is a graduate student of the Department of Applied Mathematics of the Kharkiv National University of Radio Electronics.

І.А. Шедловський, В.В. Гнатушенко, Я.І. Шедловська, В.М. Горєв

ІМІТАЦІЙНА МОДЕЛЬ ПЛАСКОГО ПОВІТРЯНОГО СОНЯЧНОГО КОЛЕКТОРА

Анатоція. Системи опалення, що використовують сонячні теплові колектори регулюють температуру підігрітого в сонячному колекторі повітря тільки зміною швидкості повітря, яке проходить через нього. Використання інформаційно управляючої системи керування (зокрема в складі IoT) для поточного керування та прогнозних обчислень потребує використовувати узагальнюючу імітаційну модель. Проведені дослідження по казали що модель, яка описує залежність температури повітря на виході з повітряного колектора від швидкості його потоку – нелінійна. Динаміка нагріву повітря визначається динамічною ланкою першого порядку. При розробці комп'ютерної системи керування необхідно враховувати використання датчиків температури повітря на вході і на виході сонячного колектора. Також потрібно використовувати дані з датчика поточної потужності сонячного випромінювання.

Ключові слова: повітряний сонячний колектор, імітаційна модель, комп'ютерна система, керування, нелінійний, динаміка.

Постановка завдання. Повітряний сонячний колектор використовується в системах опалення приміщень (житлових та виробничих). Враховуючи періодичність потоку сонячної енергії, сонячні колектори розглядаються як допоміжні джерела теплової енергії, які використовуються виключно в комплексі з основною системою опалення приміщень [1]. Робота двох систем опалення в обов'язковому порядку обладнується комп'ютерною інформаційно-управляючою системою автоматичного керування, яка дозволяє максимально використовувати сонячну енергію завдяки чому економія енергоносіїв може досягати досить значних величин. До основних функцій використовуваних систем керування входять алгоритми підтримки необхідних температурних параметрів у опалюваному приміщенні, контроль витрат енергоносіїв, також регулювання теплової потужності основної системи опалення в залежності від теплової потужності яку може забезпечити сонячний колектор. Регулювання роботи сонячного колектора повинно в обов'язковому порядку враховувати показники потужності сонячного випромінювання, температурних умов зовніш-

нього середовища, особливості опалюваного приміщення, інерційність об'єктів що використовуються у системі опалення [2]. Особливості систем, що використовують сонячні теплові колектори, є в тому, що регулювання температури підігрітого в сонячному колекторі повітря можливе тільки зміною швидкості повітря, яке проходить через колектор. Для цього використовуються вентилятори. При управлінні процесом переносу тепла від сонячного повітряного колектора до опалюваного приміщення необхідно мати математичну модель сонячного колектора як об'єкта управління.

Актуальним є також те, що при використанні інформаційно-управляючої системи керування (наприклад в складі IoT) для поточного керування та прогнозних обчислень потрібна узагальнююча імітаційна модель, до складу якої обов'язково повинна входити модель повітряного сонячного колектора.

Аналіз останніх досліджень і публікацій. Теплові процеси у повітряних сонячних колекторах на теперішній час є досить глибоко дослідженою темою. Зазвичай використовуються рівняння теплового балансу системи [3]. Відповідно до конструкції сонячних колекторів розроблені узагальнюючі залежності, визначені значення основних параметрів колектора, що надає можливість досить просто визначити можливу ефективність системи сонячного опалення ще на етапі попереднього етапу проектування [4].

Системи, де використовуються теплові сонячні колектори зазвичай обладнуються досить простими системами керування. Ці системи характеризуються тим, що в них використані статичні математичні моделі і управління процесами роботи виконується періодичними включеннями та відключеннями виконуючих пристроїв. З точки зору використання більш ефективних, безперервних, локальних систем управління сонячний колектор є досить складним об'єктом. По перше, керування можливе тільки за рахунок регулювання продуктивності вентилятора, що забезпечує циркуляцію теплоносія. По друге, забезпечення температури теплоносія на вході в опалюване приміщення також забезпечується регулюванням продуктивності вентилятора. В реальних умовах надходження сонячної енергії до колектора – процес, що визначається багатьма випадковими факторами. Тому робота системи керування повинна забезпечувати відповідне реагування на такі зміни [5]. Методи, які дозволяють розробити відповідну математичну модель об'єкта та системи керування, визначити її параметри, досить детально розроблені теорією автоматичного керування. Відповідно до цього перед розробкою математичної моделі та визначенням її

характеристик та параметрів обов'язково потрібно визначитися з принципами організації автоматичного керування роботою сонячного колектора [6].

Мета дослідження. Розробити методику побудови математичної (імітаційної) моделі динаміки роботи повітряного сонячного колектора. Визначити параметри та особливості моделі з урахуванням експлуатаційних характеристик сонячного колектора.

Виклад основного матеріалу. Функція сонячного колектора – нагрів повітря. Як відомо перетворення енергії сонячного випромінювання в теплову енергію супроводжується тепловими втратами через корпус сонячного колектора, які залежать від різниці температур абсорбера і навколишнього середовища. Втратами потужності випромінювання сонця при проходженні через прозоре покриття. Втратами пов'язаними зі значенням коефіцієнта поглинання абсорбера. Втратами тепла через ізоляцію повітроводів. Це основні види теплових втрат.

Максимально можлива температура на виході сонячного колектора ($T_{\text{макс.}}$) залежить від потужності сонячного випромінювання ($P_{\text{сон.}}$), сумарної потужності теплових втрат ($P_{\text{випр.}}$), початкової температури повітря, що є теплоносієм ($T_{\text{мін.}}$), часу нагріву (t).

$$T_{\text{макс}} = f(P_{\text{сон}}, P_{\text{випр}}, T_{\text{мін}}, t) . \quad (1)$$

Якщо припустити, що $P_{\text{сон.}} = \text{Const}$, $T_{\text{мін}}$ – визначається властивостями приміщення яке обігрівается, час t – можемо регулювати, $P_{\text{випр.}}$ – для кожного елемента сонячного колектора визначається різницею температур абсорбера і навколишнього повітря.

В цілому, за допомогою сучасних датчиків досить складно визначити $P_{\text{випр.}}$, $P_{\text{сон.}}$ – можна визначити за показаннями спеціалізованих датчиків, але вони в теперішній час досить коштовні.

Приміщення що опалюється також характеризується великою кількістю особливостей і чинників, які впливають на значення $T_{\text{мін.}}$

Енергію сонячного випромінювання, яка перетворюється в теплову визначається наступним чином:

$$P_{\text{еф}} = P_{\text{сон}} - P_{\text{випр}} . \quad (2)$$

Далі в розрахунках будемо використовувати значення теплової потужності $P_{\text{еф}}$. Робота системи опалення в режимі максимальної ефективності визначається наступним чином:

$$P_{ef} = P_{prim} , \quad (3)$$

де P_{prim} - теплова потужність, яка споживається приміщенням, яке обігрівається.

Теплова потужність, яка витрачається на обігрів приміщення визначається формулою:

$$P_{prim} = \frac{c \cdot m \cdot (T_{max} - T_{min})}{t} , \quad (4)$$

де c – теплоємність повітря, Дж/кг·К; m - маса повітря, кг; t – час, с; T_{max} - T_{min} – різниця температур (ΔT , К).

Нагрівання повітря в сонячному колекторі описується аналогічним рівнянням [3]:

$$P_{ef} = \frac{c \cdot m \cdot \Delta T}{t} . \quad (5)$$

Необхідно зазначити, що ефективна теплова потужність сонячного колектора P_{ef} визначається різницею між тепловою потужністю сонячного випромінювання і тепловою потужністю втрат, а P_{prim} теплова потужність що враховує усі теплові втрати приміщення. З позиції використання сонячного повітряного колектора для того щоб забезпечити в приміщенні потрібну температуру це дозволяє значно спростити математичні вирази.

У рівнянні (5) виходячи з значення потужності P_{ef} , ми повинні створити умови роботи, при яких вся енергія яка надходить в сонячний колектор передавалася в опалюване приміщення. Ці умови можуть скластися за рахунок зміни швидкості проходження повітря через канали абсорбера колектора, при цьому температура на виході колектора:

$$T_{max} = T_{min} + \Delta T . \quad (6)$$

Кількість тепла (Q), яке надходить в приміщення від сонячного колектора дорівнює:

$$Q = c \cdot m \cdot \Delta T . \quad (7)$$

В процесі прогріву приміщення температура T_{min} буде підвищуватися і в якийсь період передача тепла буде здійснюватися за рахунок потоку повітря, що має досить високу температуру і невелику швидкість, у разі коли нема системи, яка регулює швидкість потоку повітря і як наслідок його температуру.

Враховуючи, що робота сонячного колектора можлива тільки вдень і при ясному сонці, то виникає необхідність регулювання для забезпечення найбільшої швидкості прогріву приміщення, стабілізації температури в приміщенні,
ISSN 1562-9945 (Print) 123
ISSN 2707-7977 (Online)

режиму провітрювання (примішування до нагрітого повітря з приміщення частини зовнішнього повітря) та інші можливі варіанти роботи [7].

Маса повітря, яка проходить через сонячний колектор визначиться:

$$\frac{m}{t} = S_{абс} \cdot v_{абс} \cdot \rho, \quad (8)$$

де $S_{абс}$ - площа каналів колектора, m^2 , $v_{абс}$ - швидкість потоку повітря в колекторі м/с; ρ - об'ємна щільність повітря (кг).

$$S_{абс} \cdot v_{абс} \cdot \rho = S_{нов} \cdot v_{нов} \cdot \rho. \quad (9)$$

Швидкість руху повітря в каналах абсорбера висловимо через значення швидкості потоку повітря в повітропроводі ($v_{нов}$) і площі перетину повітропровода ($S_{нов}$):

$$v_{абс} = \frac{S_{нов}}{S_{абс}} \cdot v_{нов} \quad (10)$$

Час, за який елементарний об'єм повітря переміщується від входу до виходу абсорбера:

$$t = \frac{v_{абс}}{l_{абс}} = \frac{S_{нов} \cdot v_{абс}}{S_{абс} \cdot l_{абс}}. \quad (11)$$

Час перебування елементарного об'єму повітря в абсорбері сонячного колектора (t) це параметр, значенням якого ми можемо керувати, задаючи за допомогою вентилятора швидкість руху повітря $v_{абс}$. При цьому температура, до якої нагріється повітря на виході колектора буде залежати від t . Для створення ефективної системи управління необхідно визначити критерій, за яким слід організувати роботу системи.

Найбільша потреба в додатковому джерелі опалення виникає в холодну пору року. При цьому, у зв'язку зі зростаючими тепловими втратами температура, до якої максимально може бути розігріте повітря не може бути дуже високою, в порівнянні з температурою, яку можуть забезпечити електрокалорифери. Тому система автоматичного управління сонячним опаленням повинна забезпечувати задану температуру що подається в приміщення. Для визначення значення температури на виході колектора, запропонована схема, представлена на рис.1. На сонячному колекторі розташовуються два датчика температури. T_{min} - температура вхідного повітря; T_{max} - температура підігрітого повітря.

Сонячні повітряні колектори, особливо в зимовий час, коли сонце знаходиться низько над горизонтом встановлюються або вертикально, або з невеликим відхиленням від вертикального положення. Внаслідок нагрівання елементарного об'єму повітря, в процесі його руху від входу до виходу сонячного колектора, його температура підвищується.

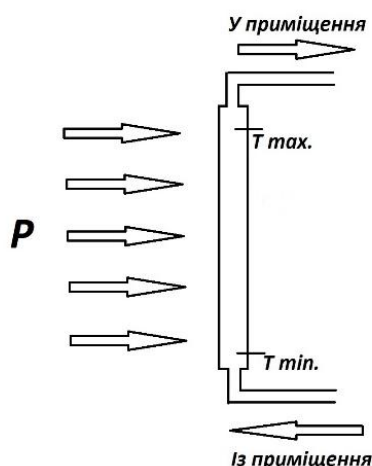


Рисунок 1 – Схема розташування датчиків температури на сонячному колекторі

Управління здійснюється за допомогою зміни швидкості потоку повітря в сонячному колекторі. Але при цьому необхідно враховувати інерційність колектора, вентилятора і аеродинамічного опору системи [8,9]. Для визначення виду математичної моделі повітряного сонячного колектора з управлінням швидкістю потоку повітря приймемо припущення, що модель повинна бути адаптована до використання узагальненої структурної схеми системи управління з урахуванням головного від'ємного зворотного зв'язку [6].

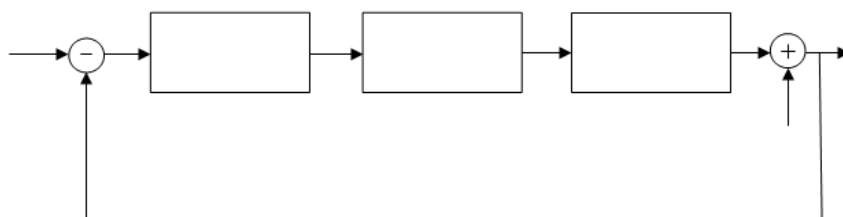


Рисунок 2 – Структура системи управління

Як об'єкт управління розглянемо наступні елементи:

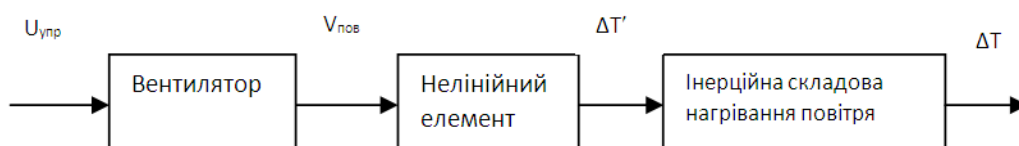


Рисунок 3 – Узагальнена структура об'єкта управління

Вентилятор з регулятором потужності можемо розглядати як без інерційну підсилювальну ланку [6]. Це припущення базується на попередньому порівнянні інерційності колекторів у цілому і інерційності вентиляторів порівняно невеликої потужності. Передавальний коефіцієнт вентилятора $K_{вент}$ визначиться зі співвідношення:

$$K_{вент} = \frac{v_{пов}}{U_{упр}}, \frac{M / c}{B}. \quad (12)$$

Передавальний коефіцієнт сонячного колектора є величиною не постійною і нелінійною. Тому в структурній схемі об'єкта керування не лінійність виділена окремим блоком. Покажемо співвідношення, що визначають нелінійну залежність температури $\Delta T'$ від швидкості $v_{пов}$:

$$\Delta T' = \frac{P_{ef}}{c \cdot \frac{m}{t}}. \quad (13)$$

Як зазначалося раніше P_{ef} – це потужність випромінювання, яка перетворюється безпосередньо в теплову енергію. Але з підвищенням температури абсорбера зростають втрати. При цьому $\Delta T'$ стабілізується.

$$\frac{m}{t} = v_{abc} \cdot S_{abc} \cdot \rho, \quad (14)$$

де ρ - об'ємна вага повітря ($\text{кг}/\text{м}^3$), m/t - масова швидкість повітря в абсорбері.

Тоді:

$$\Delta T' = \frac{P_{ef}}{c \cdot v_{abc} \cdot S_{abc} \cdot \rho} \quad (15)$$

Так як швидкість потоку повітря буде вимірюватися в повітроводі, необхідно враховувати (9):

$$\Delta T' = \frac{P_{ef}}{c \cdot v_{пов} \cdot S_{пов} \cdot \rho} \quad (16)$$

Вираз (16) дозволяє визначити область зміни коефіцієнта передачі сонячного колектора в залежності від значення потужності P_{ef} і швидкості потоку повітря в системі (рис.4).

Коефіцієнт пропорційності:

$$K_{np} = \Delta T / v \quad (17)$$

Коефіцієнт пропорційності може змінюватися в дуже широких межах від 1 до 140.

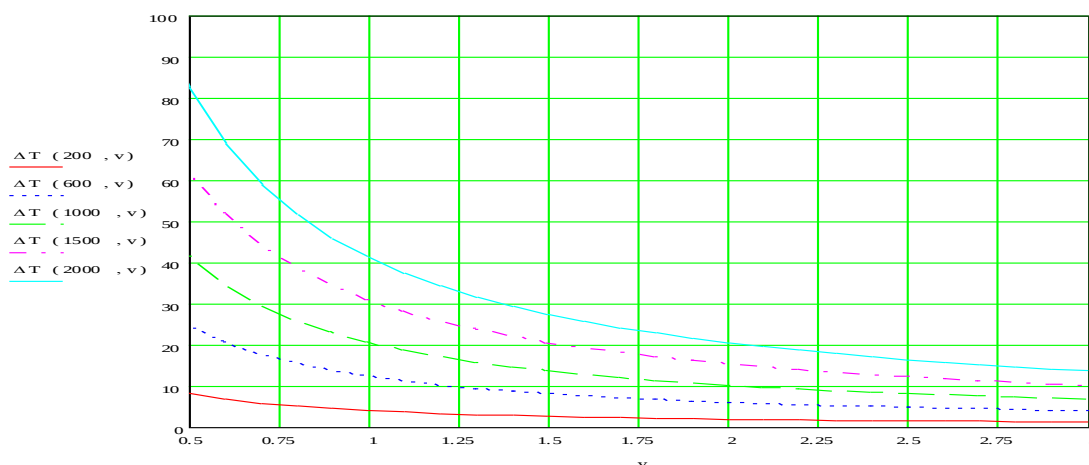


Рисунок 4 – Залежність ΔT від швидкості повітря в повітропроводі v , при значенні ефективної теплової потужності від 200 до 2000 ват

Сонячний колектор є інерційним об'єктом. Нагрівання поверхні абсорбера, потім передавання теплової енергії теплоносію – повітрю проходить досить повільно. Розглядаючи елементарні залежності (3.14) швидкості нагріву від маси тіла можемо вважати, що найбільша постійна часу сонячного колектора буде визначатися процесом нагрівання абсорбера та інших конструкційних елементів колектора. Передача тепла від абсорбера до повітря в його каналах буде в стільки разів швидшою, в скільки разів маса повітря менша за масу абсорбера. Зважаючи на метал і повітря, конструктивні особливості повітряних колекторів вага повітря що нагрівається приблизно на 2 порядки менша ніж вага абсорбера. Тому в першому наближенні можемо припустити, що модель інерційності колектора – це передаточна функція першого порядку [9].

Але в цілому перехідна характеристика може бути апроксимована як динамічна ланка першого порядку.

$$W(p) = \frac{K}{T \cdot p + 1} \quad (18)$$

Подальше використання математичної моделі об'єкта управління для синтезу локальної системи автоматичного управління та вибору типу регулятора та визначення його параметрів потребує характеризувати сам об'єкт.

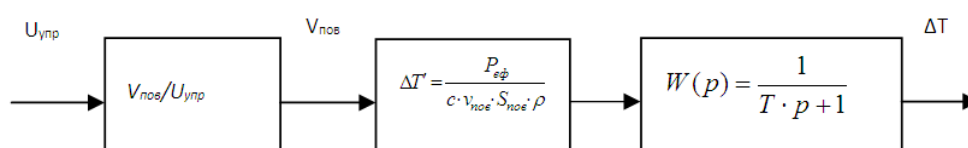


Рисунок 5 – Структура математичної моделі об'єкта управління

Головна особливість об'єкта управління в тому, що це нелінійний об'єкт. Загальний коефіцієнт пропорційності може приймати значення в межах від 0.2 до 30. Визначення числового значення постійної часу T інерційної ланки першого порядку можливо здійснити тільки за умови стабілізації швидкості потоку повітря через колектор і з урахуванням значення $\Delta T'$, яке можна розрахувати. Для цього також необхідно визначитися з значенням P_{ef} , для чого необхідний вимірювальний пристрій потужності сонячного випромінювання та дані виробника з значеннями відповідних коефіцієнтів втрат тепла. При зазначених умовах виконаний експеримент та подальша ідентифікація параметрів динамічної ланки дозволить використовувати отримані параметри для практичного використання.

Висновки і перспективи подальших досліджень. Повітряні сонячні колектори використовуються в комплексних системах обігріву, одночасна робота кількох різних систем ефективна тільки під керуванням комп'ютерної інформаційно-управляючої системи. Математична модель повітряного сонячного колектора повинна відображати не тільки пропорційні залежності а і динаміку нагрівання повітря. В роботі побудовано імітаційну модель сонячного колектора з урахуванням його використання у складі сучасних комп'ютерних інформаційних та керуючих систем. Проведені дослідження показали, що модель, яка описує залежність температури повітря на виході з повітряного колектора від швидкості його потоку – нелінійна. Динаміку нагріву повітря можна описувати динамічною ланкою першого порядку. Для роботи комп'ютерної системи керування необхідно використовувати датчики температури повітря на вході і на виході сонячного колектора. Також потрібно використовувати дані з датчика поточної потужності сонячного випромінювання.

ЛІТЕРАТУРА

- [1] Енергоефективність та відновлювані джерела енергії / Під заг. ред. А. К. Шидловського – К. : Українські енциклопедичні знання, 2007. – 560 с.
- [2] Кузнецов М. П. Особливості комбінованих енергосистем з відновлюваними джерелами енергії: монографія / М. П. Кузнецов. — Київ: ІВЕ, 2022. — 142 с.
- [3] John A. Duffie, William A. Beckman. Solar Engineering of Thermal Processes. 2013 John Wiley & Sons, Inc. 2013. DOI:10.1002/9781118671603.
- [4] Пасічник П.О. Проблеми застосування повітряних геліосистем / П.О. Пасічник, О.В. Приймак// Науково-технічний збірник «Енергоефективність в будівництві та архітектурі». Випуск 6. – К.: КНУБА, 2014 р. – С. 322 - 327.

- [5]Бабак В.П. Автоматизація і контроль регулювання теплоспоживання з використанням сонячних колекторів та акумуляції тепла / В.П. Бабак, А.О. Назаренко // Неруйнівний контроль і технічна діагностика: наук конф., 20 – 23 листопада 2012 р.: тези доп. – К., 2012. – С. 229-231.
- [6]Шаруда В. Г. Практикум з теорії автоматичного управління [Текст] : навч. посібник / В. Г. Шаруда. - Д. : НГУ, 2002. - 414 с.: - ISBN 966-7476-74-X
- [7]Дорошенко О.В. Сонячні плоскі металополімерні колектори / О.В. Дорошенко, С.С. Титар, Б.Є. Молчанський // Вісник Вінницького політехнічного інституту. – 2010. – №4. – С.32-35.
- [8]Назаренко А.О. Система управління теплоспоживанням з сонячними колекторами / А.О. Назаренко // Метрологія та прилади. – 2013. – №2 (40). – С. 151-158.
- [9]Назаренко А.О. Система керування теплоспоживанням будівель з комбінованим теплопостачанням і використанням сонячної енергії. – *Рукопис. Дисертація на здобуття наукового ступеня кандидата технічних наук за спеціальністю 05.14.01 – Енергетичні системи та комплекси. – Інститут технічної теплофізики Національної академії наук України, Київ, 2016.*

REFERENCES

1. Energy efficiency and renewable energy sources / Under general ed. A.K. Shidlovsky - K.: Ukrainian encyclopedic knowledge, 2007. - 560 p.
2. Kuznetsov M.P. Peculiarities of combined energy systems with renewable energy sources: monograph / M.P. Kuznetsov. — Kyiv: IVE, 2022. — 142 p.
- John A. Duffie, William A. Beckman. Solar Engineering of Thermal Processes. 2013 John Wiley & Sons, Inc. 2013. DOI:10.1002/9781118671603.
4. Beekeeper P.O. Problems of using air solar systems/ P.O. Pasichnik, O.V. Pryimak// Scientific and technical collection "Energy efficiency in construction and architecture". Issue 6. - K.: KNUBA, 2014 - pp. 322 - 327.
5. Babak V.P. Automation and control of regulation of heat consumption using solar collectors and heat accumulation / V.P. Babak, A.O. Nazarenko // Non-destructive testing and technical diagnostics: science conference, November 20-23, 2012: abstracts - K., 2012. - P. 229-231.
6. Sharuda V. G. Workshop on the theory of automatic control [Text]: training. manual / V. G. Sharuda. - D.: NSU, 2002. - 414 p.: fig. - ISBN 966-7476-74-X
7. Doroshenko O.V. Solar flat metal-polymer collectors / O.V. Doroshenko, S.S. Tytar, B.E. Molchansky // Bulletin of the Vinnytsia Polytechnic Institute. – 2010. – No. 4. - P.32-35.

8. Nazarenko A.O. Heat consumption control system with solar collectors / A.O. Nazarenko // Metrology and devices. – 2013. – No. 2 (40). – P. 151-158
9. Nazarenko A.O. A system for controlling the heat consumption of buildings with combined heat supply and the use of solar energy. – Manuscript. Dissertation for obtaining the scientific degree of candidate of technical sciences in the specialty 05.14.01 - Energy systems and complexes. – Institute of Technical Thermophysics of the National Academy of Sciences of Ukraine, Kyiv, 2016.

Received 17.10.2023.
Accepted 23.10.2023.

Simulation model of a flat plate air solar collector

Analysis of recent research and publications. According to the design of solar collectors, generalizing dependencies are known, the values of the main parameters of the collector are determined, which makes it possible to determine the possible efficiency of the solar heating system quite simply at the stage of the preliminary design stage. Systems where thermal solar collectors are used are usually equipped with fairly simple control systems. These systems are characterized by the fact that they use static mathematical models and the management of work processes is performed by periodic switching on and off of the executing devices. Thus using more effective, continuous, local control systems, the solar collector is a rather complex object. Firstly, the control is possible only by adjusting the efficiency of the fan, which provides the circulation of the coolant. Secondly, the temperature of the coolant at the entrance to the heated room is also adjusted by the regulation of the fan. Considering that in real conditions, the arriving of the solar energy to the collector is a process determined by many random factors, the operation of the control system must provide an appropriate response to such changes.

Problem formulation. The operation of two heating systems is necessarily equipped with a computer information-control system of automatic control, which allows the maximum usage of solar energy. As a result of that the energy savings can reach quite significant values. The main functions of the used control systems include: the algorithms for maintaining the necessary temperature parameters in the heated room, the energy consumption control, the regulation of the thermal power of the main heating system depending on the thermal power that the solar collector can provide. Regulation of the solar collector work must necessarily takes into account the indicators of solar radiation power, the temperature conditions of the external environment, the features of the heated room, the inertia of the objects used in the heating system.

Main material. To build a simulation model of a solar collector for heating a room where air is used as a heat carrier, known dependencies which describe thermal processes

were considered. It is shown that the effective thermal power of the solar collector is determined by the difference between the thermal power of solar radiation and the thermal power of losses. Taking into account that the operation of the solar collector is possible only during the day and when the sun is clear, there regulation is necessary to provide the highest rate of the room heating, stabilization of the room temperature, ventilation mode (mixing of the heated air from the room with the part of the outside air) and other possible work options.

As a control object, we will consider the following elements: a circulation fan, the dependence of the air temperature at the outlet of the collector (on the power of the solar radiation, the heat losses, the flow rate of the heat carrier), the inertial component of the heat transfer process to the heat carrier.

We can consider a fan with a power regulator as a non-inertial element. This assumption is based on a preliminary comparison of the collectors inertia in common and the inertia of fans of relatively low power. The transmission coefficient of the solar collector is a non-constant and non-linear value. Therefore, in the structure of the collector as a control object, non-linearity is highlighted by a separate block. The heat transfer from the absorber to the air in its channels will be as faster as the mass of the air is smaller than the mass of the absorber. Considering the mass of metal and air, design features of air collectors, the weight of heated air is approximately 10^2 times less than the weight of the absorber. Therefore, in the first approximation, we can assume that the inertia model of the collector is a transfer function of the first order. Further use of the mathematical model of the control object for the local automatic control system synthesis and selection of the regulator type and its parameters determination requires of the object characterizing. The main feature of the control object is that it is a non-linear object. The general proportionality coefficient can take values in the range from 0.2 to 30.

Conclusions and further research. The mathematical model of an air solar collector should display not only proportional relationships but also the dynamics of air heating. It is shown that the model that describes the dependence of the air temperature at the outlet of the air collector on the speed of its flow is non-linear. The dynamics of air heating can be described by a dynamic element of the first order. For the computer control system operation, it is necessary to use air temperature sensors at the inlet and outlet of the solar collector. It is also necessary to use data from the sensor of the current power of solar radiation.

Шедловський Ігор Анатолійович - к.т.н., доцент, доцент кафедри інформаційних технологій та комп'ютерної інженерії національного технічного університету «Дніпровська політехніка» (м. Дніпро).

Гнатушенко Володимир Володимирович - д.т.н., професор, завідувач кафедри інформаційних технологій та комп'ютерної інженерії національного технічного університету «Дніпровська політехніка» (м. Дніпро).

Шедловська Яна Ігорівна - к.т.н., доцент кафедри інформаційних технологій та комп'ютерної інженерії національного технічного університету «Дніпровська політехніка» (м. Дніпро).

Горєв Вячеслав Миколайович - завідувач кафедри фізики національного технічного університету «Дніпровська політехніка» (м. Дніпро).

Shedlovsky Igor - candidate of technical sciences, associate professor of department of information technologies and computer engineering, Dnipro University of Technology, Dnipro, Ukraine.

Hnatushenko Volodymyr - doctor of technical science, professor, head of department of information technologies and computer engineering, Dnipro University of Technology, Dnipro, Ukraine.

Shedlovska Yana - candidate of technical sciences, associate professor of department of information technologies and computer engineering, Dnipro University of Technology, Dnipro, Ukraine.

Gorev Vyacheslav - candidate of technical sciences, head of department of Physics, Dnipro University of Technology, Dnipro, Ukraine.

THE CONCEPT OF ASSOCIATIVE GRAPHICAL INTERFACE IN THE WORKFLOW AUTOMATION SYSTEM

Abstract. The article is devoted to the topical problem of developing an associative graphical interface for workflow automation systems. Based on the analysis of modern technologies and research methods, the authors set a goal to develop a new interface concept that provides optimal efficiency and ease of use. The result is the creation of Draw & GO, a new tool for automating workflows. As part of the study, a plug in architecture was used, which makes it easy to integrate new functions and optimize the operation of the automation system. Key findings highlight the potential of the associative GUI in improving productivity and streamlining workflows.

Keywords: Draw & GO, automation, macros, performance, optimization, plugin architecture.

Statement of the problem. Information technology and computerized devices have become an integral part of the daily life of mankind and the organization of business processes. To save time and minimize monotonous work, an urgent task is to automate often repeated, identical, or similar actions. Automation software can help reduce human error, improve the efficiency of operations, increase speed, and free up time for more important tasks.

Currently, almost all users of computerized devices conduct a dialogue with the operating system using a graphical interface. The graphical interface has made file and device management easier by increasing the number of users and the number of applications for personal computers.

However, in our fast-paced present, the GUI is becoming more and more cumbersome and difficult to understand – it has many nested menu items and hidden settings. Therefore, a relevant area of research is the search for new ways of human-machine interaction.

Since modern devices are capable of perceiving a stylus, touch, and mouse movements, it becomes relevant to derive useful value from more than just clicks. One such method is the use of gesture or shape recognition, namely: the user draws a continuous shape, and the computer determines what has been done by matching

it with a predefined library of shapes. After the computer defines the form, any action can be specified for it.

Analysis of the latest research and publications. The workflow engine software market is highly competitive. The largest and most developed players in this market are Microsoft Flow [1] and WorkflowCore [2]. These two products are the de facto industry standards for workflow automation. Microsoft Flow is a product that allows users to automate workflows across multiple apps and services. While the service offers a wide range of features, it can be expensive and difficult to use for many users. As a result, there is a growing number of competitors offering similar services at a lower cost. WorkflowCore is an open-source workflow engine written in .NET Core, providing an easy way to implement complex workflows in any .NET application. While WorkflowCore is gaining popularity in the industry, other options offer different benefits: Make [3], Zapier [4], MacroMaker [5], JetBit MacroRecorder [6], etc. These programs allow users to set macros to make it easier to call programs.

The disadvantage of existing solutions is the lack of gesture-graphic control. Graphical control is currently used mainly only in applications and security shells. Well-known tools are Mazelock [7] and Windows Picture Password [8]. In addition to their advantages, these programs have a significant drawback – the limited gestures that are tied to certain shapes, that is, the user can guess the password after a while.

Objective. The research of this work is aimed at creating cross-platform software for automating workflows Draw & GO, which implements gesture-graphic interaction to achieve a minimum interface and maximum intuitiveness in control.

Presentation of the main material of the research. Draw & GO is an innovative cross-platform workflow system that is built on the microkernel architecture of plugins written on .NET MAUI and Blazor technologies, implements interaction based on graphic gesture recognition, allows users to create powerful and secure workflows customized according to the needs and preferences of each client based on specific conditions, actions and triggers in a secure environment of industry standards, Assists users in automating processes and numerous tasks. The software consists of a multi-step workflow that consists of user-defined plugins and connectors for connecting systems. The app combines workflow automation, visual design, and web-based APIs to provide a powerful and flexible solution for all users.

Draw & GO has several features that set it apart from other workflow mechanisms:

- Graphic Gesture Recognition: Draw & GO is one of the first workflow engines to use gesture recognition as an input method, users can draw any shapes on their screens to create workflows without typing or making mouse clicks;
- Visual Designer: Users can drag and drop plugins and connectors to define conditions, actions, triggers, loops, forks, and more;
- Real-time synchronization: users can see their changes instantly on all devices;
- Customization: Users can create their plugins or use existing ones from the library, including connectors that connect to a range of applications, including Microsoft, Google, and Dropbox.

The idea of using graphic gestures (Fig. 1) is based on a person's tendency to associative thinking [9]. The main role of associations in memorization is that we tie new knowledge to information we already know. The research of the German scientist G. Ebbinghaus shows that a person remembers objects and images better than sequences of symbols. In this way, personalized graphic code provides convenient and quick access to macros and Helps to optimize the workflow.



Figure 1 – Example of a user-drawn gesture to run a macro

To create and recognize graphic code, the DrawGo.GestureLibrary was created, which consists of the following 3 classes:

- PointPatternAnalyzer.cs – Performs an actual comparison of the gesture with a list of trained gestures.
- PointPatternMatchResult.cs – Represents the result of comparing a single gesture, including the macro ID and match probabilities.
- PointPatternMath.cs – contains helper functions for performing point distribution, angle calculation, and distance calculation (for point distribution, not comparison).

Before dividing, the gesture contains an array of unevenly spaced points and may contain more or less of the desired precision.

Let's take a look at the graphic gesture recognition algorithm, which is implemented in the PointPatternAnalyzer.GetPointPatternMatchResult (Macro compareTo, Point[] points method:

Let's break the graphic gesture into segments and find the desired length of each of them:

$$L_w = \frac{\sum_{i=1}^{n-1} \sqrt{(x_{i+1} - x_i)^2 + (y_{i+1} - y_i)^2}}{N}, \quad (1)$$

where n is the actual number of points, and N is the desired number of points. N has been chosen to be 100, but the more points we distribute for the gesture, the more accurate the comparison should be. However, the more points we check for each gesture, the slower the overall comparison will be. The distribution of a gesture implies the ability to significantly increase or decrease the number of points in a gesture to a certain fixed value. So, for example, a simple gesture like a straight line can be only 2 points, but a more complex gesture like M will take a minimum of 5 points. Most gestures will have many more points because they are made with a mouse or finger.

Let's set the initial values: the current length of the segment, the auxiliary point equals the first input point, which is also the first point we are looking for. $L_c = 0$

Let's set the boundaries of the indices $1 \leq i < n$ and $1 \leq j \leq N$

Let's set the current point as P_i

Let's find the increment of the current length of the segment as the distance between the auxiliary point and the current one:

$$\text{increment} = \sqrt{(x_{i+1} - x_i)^2 + (y_{i+1} - y_i)^2} \quad (2)$$

Let's find the test length of the segment:

$$l = L_c + \text{increment} \quad (3)$$

If $l < L_w$, then $L_c = l$, replace the auxiliary point with the current one, and go to step 4.

Let's calculate the proportionality factor to find a new point:

$$\text{koef} = \frac{L_w - L_c}{\text{increment}} \quad (4)$$

Let's find the coordinates of the desired point (Fig. 2)

$$X_j = (1 - \text{coef}_i) * x_{i-1} + \text{coef}_i * x_i \quad (5)$$

$$Y_j = (1 - \text{coef}_i) * y_{i-1} + \text{coef}_i * y_i \quad (6)$$

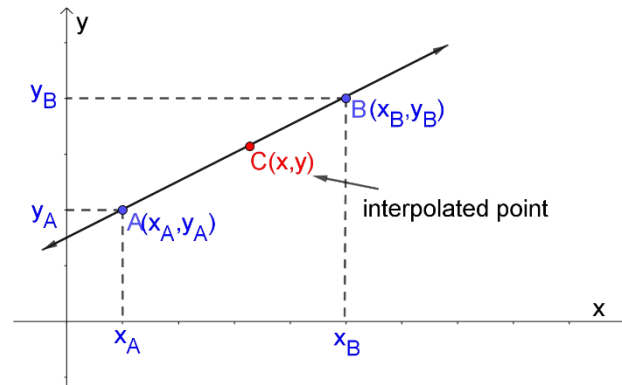


Figure 2 – Search for an interpolated point

If $j < N$, set the auxiliary point to the desired one, $L_c = 0$, reduce the index "i" by 1 and return to step 4.

The distribution is complete.

Once distributed, the movement shapes are preserved, and all points are evenly distributed and contain the desired number of points to compare.

When the distribution is complete, we convert the array of points into an array of angles. We do this because we are comparing the angular difference between the two gestures. We get an array of angles for each gesture.

$$\text{angle}_j = \tan^{-1}(Y_{j+1} - Y_j, X_{j+1} - X_j) \quad (7)$$

Now that we have angles for each gesture, we can iterate over each corner in the gesture and calculate the angular difference between the two corners. Let's place them in the created array.

$$\text{angleDifferences}_j = |\text{angle}_{j+1} - \text{angle}_j| \quad (8)$$

If $\text{angleDifferences}_j > \pi$

$$\text{angleDifferences}_j = \pi - (\text{angleDifferences}_j - \pi) \quad (9)$$

Use an array to calculate the mean angular difference between the two gestures, and convert it to a match accuracy (from 0% to 100%).

$$\text{probability} = 100 - \frac{\sum_{i=1}^n \text{angleDifferences}_i}{n} * \frac{100}{\pi} \quad (10)$$

That's all it takes to accurately recognize the difference between the two gestures. Once we have our probability, let's put everything in a PointPatternMatchResult, which contains the gesture and probability ID.

The above is a comparison of an array of points with a single pattern for a list of saved gestures. The next step is a calculation of the probability for each gesture in the sample set, which will help to evaluate them for the best match based on the highest probability.

Draw & GO is designed to be powerful and intuitive, providing a drag-and-drop-based graphical user interface for developing complex workflow systems (Figure 3).

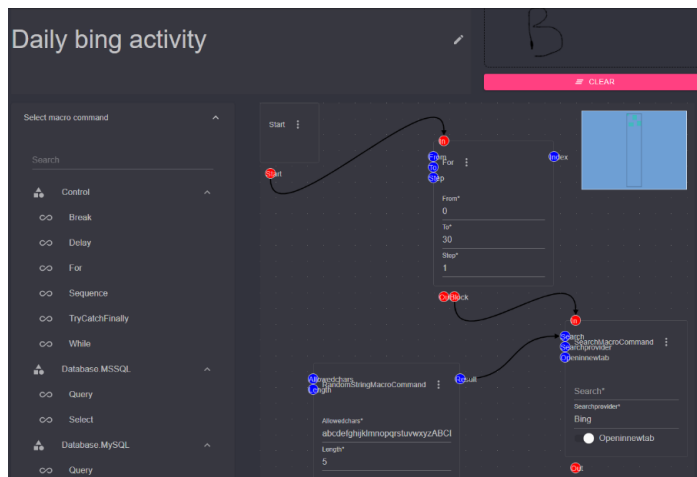


Figure 3 – Macro Configurator

Developers can create and customize flowcharts using drag-and-drop components, and the platform allows them to create customized applications that meet their specific business needs.

The software product "Draw & GO" consists of two modules, one of which is used to create and organize macros, and associate them with graphic code, and the second is used to open any exported macro and conduct online consultations. The created program allows users to increase the speed of working with a computer, reproducing several identical actions to facilitate the user's work and to protect private data. The "Give-a-Hand" subroutine allows users to organize online consultations, control another computer, and share files remotely.

Conclusions. Draw & GO is an innovative cross-platform workflow engine with microkernel plugin architecture and graphic gesture recognition. The mi-

crokernel architecture provides scalability, and plugin support allows users to customize the platform to suit their needs. With an easy-to-use visual workflow designer, users can quickly create and override complex workflows for many purposes. Regardless of what needs to be automated, whether user processes are simple or company workflows are complex, Draw & GO can help everyone achieve their goals. Due to the flexibility of its plugin architecture, it can be integrated with any existing third-party services such as Slack, Google Calendar, Dropbox, and many more. Draw & GO is designed to run on all major operating systems, including Windows, MacOS, iOS, Android, and Tizen. It works effectively on desktop and mobile devices, as well as web browsers, providing an intuitive user experience.

Users can get acquainted with the functionality of the Draw & GO software product and find out about available licenses on the official page <https://drawgo.azurewebsites.net> [10].

REFERENCES

1. Power Automate [Electronic resource]. – Access Mode: <https://powerautomate.microsoft.com/en-us/>.
2. Workflow Core [Electronic resource]. – Access Mode: <https://workflow-core.readthedocs.io/en/latest/>.
3. Make [Electronic resource]. – Access Mode: <https://www.make.com/en>.
4. Zapier [Electronic resource]. – Access Mode: <https://zapier.com/>.
5. MacroMaker [Electronic resource]. – Access Mode: <https://macromaker.en.softonic.com/>.
6. JitBit Macro Recorder [Electronic resource]. – Access Mode: <https://www.jitbit.com/macro-recorder/>.
7. Mazelock [Electronic resource]. – Access Mode: <https://eusing.com/mazelock/pclock.htm>.
8. Gao, H., Jia, W., Liu, N., Li, K. (2013). The Hot-Spots Problem in Windows 8 Graphical Password Scheme. In: Wang, G., Ray, I., Feng, D., Rajarajan, M. (eds) Cyberspace Safety and Security. CSS 2013. Lecture Notes in Computer Science, vol 8300. Springer, Cham. https://doi.org/10.1007/978-3-319-03584-0_26.
9. Antonyuk V.A., Segeda N.E. (2015). Recognition of graphical passwords as the means to accelerate the work with software projects. // Modern scientific and technical research in the context of language space. - T.2 – C. 5-6.
10. Draw & GO [Electronic resource]. – Access Mode: <https://drawgo.azurewebsites.net>

Received 21.11.2023.
Accepted 24.11.2023.

**Концепція асоціативного графічного інтерфейсу
у системі автоматизації робочих процесів**

Стаття акцентує увагу на важливій проблематиці створення інтуїтивно зрозумілих та ефективних інструментів для оптимізації та автоматизації робочих процедур. У контексті швидкого розвитку сучасних технологій автори статті визначили як свою мету розробку інноваційного інтерфейсу, який би забезпечував високу продуктивність з мінімальними зусиллями з боку користувачів.

Автори статті детально пропрацювали концепцію асоціативного графічного інтерфейсу, розглянувши глибину асоціативного мислення як фундаментального аспекту людської когніції. Це дозволило авторам сформулювати концепцію асоціативного графічного інтерфейсу, яка лягла в основу розробки інструменту "Draw & GO".

"Draw & GO" являє собою новаторське рішення, яке дозволяє користувачам швидко та зручно створювати автоматизації для повторюваних робочих процесів. Використовуючи метод drag-and-drop, користувачі можуть перетягувати макроси та створювати сценарії автоматизації без необхідності написання складного коду або розуміння принципів програмування.

Однією з ключових характеристик "Draw & GO" є його плагін-орієнтована архітектура. Архітектура плагінів надає велику гнучкість системі, дозволяючи легко додавати нові функціональні можливості та вдосконалюючи вже існуючі без перешкод для загальної ефективності системи. Такий підхід також забезпечує зниження часу, необхідного для впровадження нововведень та оптимізації робочих процесів завдяки швидкому розгортанню нових функцій. Draw & GO" ефективно працює на настільних та мобільних пристроях, а також у веб-браузерах, надаючи інтуїтивно зрозумілий користувацький досвід.

Користувачі можуть ознайомитися з функціональністю програмного продукту "Draw & GO" на офіційній сторінці <https://drawgo.azurewebsites.net>. Висновки дослідження подають "Draw & GO" як зразковий приклад синергії асоціативного графічного інтерфейсу та плагін-орієнтованої архітектури, що в результаті пропонує масштабоване, гнучке та користувач-орієнтоване рішення для автоматизації робочих процесів у різноманітних секторах бізнесу..

Ключові слова: Draw & GO, автоматизація, макроси, продуктивність, оптимізація, плагін-архітектура

Антонюк Владислав Андрійович - аспірант Дніпровського національного університету імені Олеся Гончара.

Сидорова Марина Геннадіївна - доцент Дніпровського національного університету імені Олеся Гончара.

Antonyuk Vladislav - PhD student Oles Honchar Dnipro National University.

Sydorova Maryna - Associate Professor Oles Honchar Dnipro National University.

В.В. Скалозуб, В.М. Горячкін, І.А. Терлецький, І.П. Дудник

ФОРМУВАННЯ МОДЕЛЕЙ КЛАСИФІКАЦІЇ НЕВИЗНАЧЕНИХ ДАНИХ ПРОЦЕДУРАМИ РЕДУКЦІЇ І КАППА СТАТИСТИКИ

Анотація. Стаття присвячена розвитку математичних моделей класифікації невизначених даних, представлених нечіткими величинами та коефіцієнтами упевненості $CF(A)$. Процедури формування шаблонів діагностування використовують модифіковані мережі Хеммінга (МХН), а також методи редукції та статистики каппа Коена. При цьому визначаються граничні розмірності та склад параметрів моделі класифікації, які забезпечують встановлені ймовірнісні вимоги достовірності результатів розрахунків. Представлена процедура редукції простору моделі діагностування невизначених даних. У статті наведено постановки, математичні моделі та реалізації завдань класифікації за недетермінованими даними. Прикладом моделі класифікації за нечіткими даними являється завдання із встановлення авторів україномовних текстів. Завдання класифікації при даних у форматі $CF(A)$ відповідає відбору кандидата. Результати чисельного моделювання дозволили встановити результативність, достовірність та ефективність запропонованих процедур формування достовірних моделей класифікації при невизначених даних.

Ключові слова: класифікація, достовірні моделі, розмірність простору, нечіткі величини, коефіцієнти $CF(A)$, модифікована мережа Хеммінга, процедура редукції, статистика каппа Коена, україномовні тексти, визначення автора, комп'ютерне моделювання.

Вступ та постановка проблеми. Завдання та процедури класифікації та діагностування за умов неповної визначеності вихідних даних (збурені, неповні ін.) являються досить поширеними на практиці [1, 2, 4, 6]. За їх результатами формуються моделі оптимального керування різноманітними технологічними процесами, вибору раціональних заходів/виконавців тощо [1, 6, 9]. Відзначається одна з головних проблем завдань класифікації, яка полягає у встановленні властивостей достовірності та повноти моделей (формати, структура та кількість зразків тощо, які забезпечують достовірний результати). Одним із актуальних нових аспектів моделей діагностування являється узгодженість ймовірнісних вимог до достовірності результатів, розмірності параметрів шаблонів, а також кількості даних (класів) за якими виконується

діагностування. Для підвищення достовірності результатів та ефективності класифікації на основі модифікованих мереж Хеммінга (МХН) актуальними являються дослідження можливостей і застосування вимог статистики каппа Коена та методу редукції розмірності [1, 3]. Важливе значення мають постановки завдань керування сервісними системами (С&С) з неповними та неточно визначеними даними і даними в природньомовній формі, які дозволяють представити рішення на основі моделей класифікації.

Аналіз останніх досліджень і публікацій. У сучасних складних системах деякі параметри станів або контрольовані характеристики процесів можуть мати значну ступінь невизначеності [1, 5, 6, 9]. Створення інтелектуальних процедур та інформаційної технології (ІІТ) для оптимізації потоків замовлень у С&С запропоновано у [1]. В ній вибір керувань виконується шляхом реалізації завдань діагностування з урахуванням умов невизначеності на основі МХН.

Завдання розвитку моделей та інтелектуальних багатопараметричних процедур діагностування за неповними і збуреними даними являються нате-пер актуальними [1, 2, 6]. У статті [6] було запропоновано підхід до формування шаблонів моделей класифікації на основі МХН [1]. При тому відзначено необхідність застосування методу редукції [3] та статистики каппа Коена [4]. Ці методи разом забезпечують отримання результатів із встановленими гарантіями щодо достовірності результатів. При тому в шаблонах МНХ використовуються дані виду коефіцієнтів упевненості, $CF(A)$ [5].

Удосконалені моделі МНХ мають суттєву відмінність від класичних моделей асоціативної пам'яті Хеммінга в завданнях класифікації за неповними та збуреними даними. В МНХ використовуються в якості моделей даних нечіткі множини ($\mu_X: X \rightarrow [0; 1]$), а також коефіцієнти впевненості $CF(A)$ із значеннями в множині $[-1; +1]$. При тому класичні моделі Хеммінга [10] представляють дані за допомогою дискретної множини $\{-1; +1\}$. У [1, 6] удосконалено математичні моделі та процедури класифікації МХН з урахуванням показників, які відображають неточність або природномовну структуру первинних даних. Також представлено програмні засоби і результати експериментальних досліджень. За ними встановлена надмірність параметрів-ознак шаблонів, які використовувались для завдань визначення авторства україномовних текстів [7].

Для оцінки граничної розмірності параметрів моделі класифікації « n_0 », простору який синтезується, використовуються результати теореми Вапника-Червоненкіса і процедури методу граничних спрощень (МСП) [3]. Величина

$$n_0 = (\varepsilon * L + \ln(h)) / \ln(m) \quad (1)$$

визначає розмірність простору, перевищення якої призводить до втрати гарантії досягнення заданих параметрів достовірності « ε , h »; $(1 - h)$ оцінка ймовірності умови безпомилкового розділення випадкової і незалежної вибірки довжини « L » при заданій граничній величині помилкової класифікації « ε ».

В цій статті розроблена нова процедура, яка забезпечує формування простору (1) багатопараметричної класифікації з указаними вимогам достовірності.

Мета дослідження – підвищення достовірності моделей класифікації при невизначених даних, представлених нечіткими величинами та коефіцієнтами упевненості $CF(A)$, шляхом застосування процедур редукції і статистики каппа Коена. Отримати постановки, моделі та реалізації завдань класифікації за нечіткими даними, які забезпечують встановлення авторів україномовних текстів, реалізувати завдання вибору кандидата за даними у форматі $CF(A)$.

Результати та основний матеріал дослідження. Головним завданням, що вирішене у роботі, було формування достовірних математичних моделей класифікації та процедури аналізу при невизначених даних певних типів на основі методів редукції та каппа статистики (ПКР). ПКР складається із таких етапів:

- 1) виконати оцінку показника «каппа» ступеню подібності результатів класифікації на основі моделей шаблонів з різним числом або складом параметрів,
- 2) при забезпеченні «подібності» результатів класифікації для різних моделей-шаблонів виконати скорочення моделі, залишити один із шаблонів,
- 3) визначити і видалити найменше значимі або найбільше «подібні» між собою параметри моделі класифікації, враховуючи значення « n_0 ».

Узагальнений варіант циклу скорочення розмірності «шаблонів» моделі при вхідному «еталон»/«вимоги». Нехай $N(t) > n_0$ число параметрів шаблону на етапі (t) розмірності $N(t=0) = «m»$. Утворюється множина конкуруючих моделей шаблонів меншої розмірності. Для цього з набору змінних шаблонів $N(t)$ і «еталону» вимог вилучається, наприклад, параметр X_j ; нові вектори для цих спрощених моделей шаблонів позначимо (X/X_j) . Такі спрощені моделі формуються для кожної змінної із числа $N(t)$, а за ними виконується класифікація на основі MHX .

Перевірити «подібність» оцінок за каппа статистикою результатів класифікації для усіх пар змінних шаблонів $\{(X/X_i)\}$. Для різних наборів змінних (X/X_i) шаблонів за каппа-оцінками для кожного шаблону призначаються значення «+» або «-» в залежності від результатів класифікації. На основі порівняльного аналізу для всіх змінних утворюють таблицю розбіжностей, за якою розраховують статистичні оцінки «каппа Коена» [4]:

$$K = (P_0 - P_e) / (1 - P_e), \quad (2)$$

В (4) P_0 – ймовірнісна оцінка що показує наскільки спостережувана узгодженість краща за випадкову, а P_e – результат підрахунку максимально можливої узгодженості за винятком випадкової узгодженості конкуруючих шаблонів [4].

Для вибору параметра спрощення моделі $N(t)$ визначається пара наборів змінних виду (X/X_i) із найбільшим значенням «K» (2) та експертна оцінка «подібності» скорочених наборів (X/X_i) . Якщо величина оцінки пари (2) відповідає вимогам, то певна змінна (їх множина) може бути видалена з «шаблонів» та «еталону». Після видалення таких параметрів з моделі класифікації увесь цикл розрахунків повторюється до виконання умови ($m \leq n_0$).

З метою удосконалення мережі Хеммінга (MX) при встановленні оцінок величин відстані між нечіткими елементами (μX) зразків (w_{ik}) і вхідним вектором $X = \{x_i: i=0...n-1\}$ уведене нечітке відношення наступного виду

$$R(W, X) = \{\mu R(w_i, x_i) / \{w_i, x_i\} = (1 - \text{abs}(w_i - x_i)) / \{w_i, x_i\}, i=0, 1, \dots, n-1. \quad (3)$$

Якщо значення ступенів приналежності величин $\{w_i, x_i\}$ однакові, тоді значення $\mu R(w_i, x_i) = 1$. Коли одна з величин $\{w_i, x_i\}$ дорівнює 0, а інша 1, тоді значення $\mu R(w_i, x_i) = 0$, інакше $\mu R(w_i, x_i) \rightarrow [0; 1]$. Відношення (3) є одною із можливостей реалізації MXH для формування нечітких моделей класифікації, а також моделей представлених коефіцієнтами упевненості $CF(A)$ [5].

Результати циклу процедури редукції
при формуванні моделі класифікації

K_i	L_i	X	X/{X ₁ }	X/{X ₂ }	X/{X ₃ }	X/{X ₄ }	X/{X ₅ }	X/{X ₆ }
K_1	L_{11}		+					+
	L_{12}	+		+		+		
	L_{13}				+			+
K_2	L_{21}					+		+
	L_{22}		+	+				
	L_{23}				+			
	L_{24}	+				+		
K_3	L_{31}		+	+				
	L_{32}	+					+	
K_4	L_{41}							
	L_{42}		+		+	+		
	L_{43}							
	L_{44}	+				+		+
	L_{45}					+		
K_5	L_{51}			+				
	L_{52}						+	
	L_{53}							
K_6	L_{61}			+				
	L_{62}				+	+		
	L_{63}	+						
	L_{64}				+			
K_7	L_{71}						+	
	L_{72}				+			
	L_{73}	+						+
K_8	L_{81}					+		
	L_{82}	+	+	+				

Для пояснення процедури редукції приведемо умовний приклад реалізації наведених вище завдань на основі табл. 1 результатів класифікації для шаблонів з параметрів ($X_1, X_2, \dots, X_6$). В табл. 1 позначено : K_i – класи/зразки

моделі, які можуть містити кілька екземплярів L_i ; X – набір параметрів моделі класифікації на поточному етапі процесу редукції; $X/\{X_j\}$ – скорочені набори моделей класифікації та вхідних векторів на поточному етапі без множини параметрів $\{X_j\}$, які видалені з моделей; знаки «+» – визначення модифікованих шаблонів переможців (класів) при застосуванні моделі МХН. У стовпці X позначаються шаблони, які були визначені на основі X наборів параметрів. Саме серед скорочених моделей шаблонів $X/\{X_j\}$ визначається набір параметрів $\{X_{ij}\}$, які необхідно видалити на наступному етапі процедури редукції.

Таблиця 2

Розрахунки показників каппа статистики за таблицями розбіжностей

Розрахунок за класами (KM_1, C_2) Розрахунок за класами (KM_1, C_1)

$K(X_1/X_2)$ 0,143	Так	ні	$K(X_1/X_2)$ 0,143	Так	ні
	так	4	1	так	4
	ні	2	1	ні	1
$K(X_1/X_3)$ -0,067	Так	ні	$K(X_1/X_3)$ -0,333	Так	ні
	так	3	2	так	4
	ні	2	1	ні	2
$K(X_1/X_4)$ 0,467	Так	ні	$K(X_1/X_4)$ 0,333	Так	ні
	так	4	1	так	5
	ні	1	2	ні	1
$K(X_1/X_6)$ 0,25	Так	ні	$K(X_1/X_6)$ 0,143	Так	ні
	так	3	2	так	4
	ні	1	2	ні	2
$K(X_3/X_4)$ 0,467	Так	ні	$K(X_3/X_4)$ 0,333	Так	ні
	так	4	1	так	5
	ні	1	2	ні	1
$K(X_3/X_6)$ 0,750	Так	ні	$K(X_3/X_6)$ 0,714	Так	ні
	так	4	1	так	5
	ні	0	3	ні	0
$K(X_4/X_6)$ 0,250	Так	ні	$K(X_4/X_6)$ 0,143	Так	ні
	так	0	5	так	4
	ні	3	0	ні	1

Множини $\{X_j\}$ можуть містити не одну, а будь-яку кількість змінних, встановлену у завданні. Прийняте в табл. 1 число параметрів моделі класифікації $n=6$ – умовне, а граничне значення $n_0 = 4$. Дані табл. 1 фіксують результати класифікацій за мережею МХН, які отримані при однакових вихідних постановках завдань для всіх конкуруючих варіантів, сутність яких зараз не суттєва.

Множина варіантів моделей класифікації (КМ) містить всі комбінації пар скорочених моделей $X/\{X_j\}$. Виконуються наступні передумови аналізу результатів. КМ₁ – оцінюється класифікація хоча б одним із шаблонів L_{ij} класу K_i , КМ₂ – оцінюється класифікація кожним із L_{ij} . Також досліджуються різні способи формування таблиць розбіжностей для розрахунку показників каппа (2). Відповідно першого способу, C_1 , враховуються відомі попередні результати класифікації за набором X , відповідно способу C_2 такі дані не використовуються, контролюються лише отримані результати (шаблони/класи) за МХН. Результати формування таблиць розбіжностей і розрахунків показників каппа (2) за даними табл. 1 для моделі класифікації КМ₁, коли в якості конкуруючих варіантів розглядалися всі комбінації пар скорочених моделей $X/\{X_j\}$, показані табл. 2.

У табл. 2 приведена частина результатів розрахунків оцінок «каппа», яка дозволяє визначити першу змінну для скорочення моделі класифікації – X_6 , $K(X_3/X_6) = 0,75$ (значний рівень узгодженості). Для інших не приведених у табл. 2 пар оцінки величин «каппа» були несуттєвими, подібність відсутня, як для $K(X_1/X_3)=-0,067$. Можна сказати, що «вплив» фактору X_6 на модель класифікації ураховують та опосередковано представляють фактори X_1 , X_3 і X_4 . Відзначається, що оцінки подібності за каппа статистикою для таблиць розбіжностей без урахування класифікації за набором X , (спосіб C_2) перевищують чи дорівнюють результати за способом C_1 . У табл. 3 приведено основні результати розрахунків оцінок показників «каппа» для скороченого складу параметрів класифікації, $(X_1, X_2, \dots, X_5)$, представлені за схемою табл. 2.

Найвище значення «каппа» Коена $K(X_1/X_3) = 0,53$ (помірний рівень узгодженості), тому видаляється X_3 (подібність до X_1, X_4, X_5). Відзначається також відмінність результатів стосовно етапів табл. 2 та табл. 3 за способами формування таблиць розбіжностей C_1 і C_2 . А саме – щодо врахування попередньо відомих результати класифікації за набором X . У табл. 3 більш високі оцінки показників були отримані за способом C_1 , а не у табл. 2. Тобто необхідно розглядати обидва способи C_1 та C_2 при виконанні скорочення числа параметрів моделі, обираючи рішення за більшими значеннями «каппа». Також

необхідна перевірка тотожності та можливості корегування набору шаблонів моделі класифікації після видалення параметрів $\{X_j\}$. У разі утворення кількох однакових шаблонів скороченої моделі $X/\{X_j\}$ у одному або у кількох класах K_i , залишається лише один. Якщо такі шаблони виникли в одному класі, залишають тільки один, а в разі належності зазначених зразків до різних класів K_i видаляються шаблони із класу з більшим числом екземплярів. Наведемо результати процедури редукції для KM_2 , коли оцінюється класифікація табл. 1 кожним із 26 зразків L_{ij} .

Таблиця 3

$K(X_1/X_2)$ -0,600		Так	ні	$K(X_1/X_3)$ 0,529		Так	ні
	так	2	3		так	3	2
	ні	3	0		ні	0	3
$K(X_2/X_5)$ -0,412		Так	ні	$K(X_3/X_4)$ 0,467		Так	ні
	так	1	4		так	2	1
	ні	2	1		ні	1	4
$K(X_3/X_5)$ 0,143		Так	ні	$K(X_4/X_5)$ -0,429		Так	ні
	так	1	2		так	0	3
	ні	1	4		ні	2	3

Параметр X_1 табл. 4 має найвищий показник «каппа» $K(X_1/X_2)= 0,494$ (помірний рівень узгодженості) і разом з тим високі оцінки подібності для $K(X_1/X_3) = 0,299$ та $K(X_1/X_5) = 0,325$, які перевищують значення інших. X_1 представлений через зазначені змінні, тому X_1 видаляється з моделі. У табл. 4 також більш високі оцінки показників «каппа» отримані за способом C_1 .

В якості моделі класифікації за нечіткими даними розглядається завдання із визначення авторства творів україномовних авторів (ЗАТ) [1, 7]. У ЗАТ на підставі сукупності ознак, які характеризують набори текстів і утворюють класи або зразки моделі, необхідно визначити клас, до якого належить новий текст (закодований вектор ознак). Система ознак у авторів однакова, а число наборів (творів) може бути різним. Також певні дані кодів ознак профілю можуть бути «збуреними» або навіть відсутніми. Встановлено, що формування одного окремого шаблону для кожного автору представляє певну проблему. Тому кожному автору відповідає кілька зразків/шаблонів у формі векторів нечітких величин. Відповідно [7] натеper існують системи «вимірювання» властивостей текстів із понад 60 параметрів/ознак (розмірність класифікації $m>60$). Вимоги до моде-

лей класифікації ЗАТ виконанні при представленні даних у формі таблиці табл. 1.

Таблиця 4

Розрахунки $m=6$ за зразками (KM_2, C_2) Розрахунки $m=6$ за зразками (KM_2, C_1)

$K(X_1/X_2)$ 0,425		Так	ні	$K(X_1/X_2)$ 0,494		Так	ні
	так	3	2		так	14	4
	ні	3	18		ні	2	6
$K(X_1/X_3)$ -0,035		Так	ні	$K(X_1/X_3)$ 0,299		Так	ні
	так	1	4		так	10	5
	ні	5	16		ні	4	7
$K(X_1/X_5)$ -0,169		Так	ні	$K(X_1/X_5)$ 0,325		Так	ні
	так	0	5		так	13	5
	ні	3	18		ні	3	5
$K(X_3/X_5)$ -0,182		Так	ні	$K(X_3/X_5)$ 0,308		Так	ні
	так	0	6		так	11	2
	ні	3	17		ні	7	6
$K(X_5/X_6)$ -0,169		Так	ні	$K(X_5/X_6)$ 0,278		Так	ні
	так	0	3		так	14	4
	ні	5	18		ні	4	4

У [1, 7] кожному автору україномовного тексту відповідав один узагальнений шаблон, вибір творів до формування проблематичний. Вхідний вектор (твір невідомого автора) кодується встановленим набором нечітких величин, всі зразки творів моделі мають єдину форму відображення властивостей текстів. Для прикладу дослідження авторства україномовних текстів використано дані [1]. Ознаками текстів україномовних авторів (K_1 – І. Багрянний, K_2 – О. Довженко, K_3 – М. Жадан, K_4 – М. Коцюбинський, K_5 – Л. Українка, K_6 – П. Мирний, K_7 – І. Франко. (K_8 – інші)) були такі: X_1 – математичне очікування; X_2 – середнє квадратичне відхилення; X_3 – рекурентність; X_4 – детермінізм; X_5 – середня довжина діагональних ліній; X_6 – дивергенції; X_7 – ентропія; X_8 – завмирання; X_9 – затримки; X_{10} – середня кількість слів, X_{11} – середня кількість складів, X_{12} – середня кількість літер у реченні; X_{13} – середня кількість складів та X_{14} – літер у словах. Для авторів ($K_1, \dots, K_7$) на основі десяти творів формувалися шаблони на основі розрахунку областей можливих діапазонів та середніх значень. За творів інших авторів (Т. Шевченко, М. Стельмах ін.) формувався шаблон K_8 – інші. Величини $X_1 - X_{14}$ нормувалися

до інтервалу $[0; 1]$. Для «визначення авторства» на мережу МХН подавалися нормовані значення характеристик твору невідомого автора; Мережа МХН, як правило, вірно встановлювала клас автора твору.

Разом з тим було встановлено, що для певних вхідних творів були визначені спрощені шаблони із 4-х параметрів, які за результатами класифікації відповідали сукупності параметрів-ознак $X_1 - X_{14}$. Тож показана необхідність формування шаблонів моделей ЗАТ з урахуванням вимог методу спрощень [4].

Модель класифікації при даних у форматі CF(A) змістовно відповідає завданню менеджменту – відбору кандидата із зазначеної множини (ЗК). Завдання призначення ЗК [8] відоме, має багато моделей і процедур реалізації, в залежності від типу та ознак параметрів, способів їх отримання і оцінювання ін. Особливість ЗК цієї роботи наступна. Вважається, що реалізується процедура відбору «кандидата» за даних його портфоліо (переліку робіт у певних проєктах) або резюме. Кожний з кандидатів K_i має набір завдань Z_{ik} портфоліо P_i , які описані без попередніх вимог до структури опису Z_{ik} . Параметри та оцінки за Z_{ik} можуть бути представлені і текстом/мовна-форма, відрізняються для різних кандидатів P_i . У тому числі мають оцінки у формі неточних показників ін. Остаточну структуру і значення оцінок щодо представлення Z_{ik} визначає менеджмент, який встановлює оцінки всіх ознак у формі коефіцієнтів упевненості CF(A) У підсумку шаблони різних кандидатів K_i у моделі ЗК суттєво неоднорідні, можуть мати «пропуски» в ознаках, число шаблонів кандидатів K_i різне. Вхідний «еталон» містить всі характеристики, ураховані та означені «менеджментом» при формуванні моделі ЗК за портфоліо P_i . Метою являється визначення шаблону та K_i , який в найбільше відповідає «еталону», ким виконувалося подібне завдання.

Постановка ЗК структури табл. 1 (для $K_1, ..., K_4, ..., K_{ij}$) з даними, закодованими коефіцієнтами CF(A), відзначається таким. Кандидати мають різну кількість зразків, певні зразки не мають деяких ознак (наприклад, K_{11} і K_{12} не мають даних для X_5 і X_{11} , а K_{41} не містить X_5, X_6, X_9 та X_{11}). У «еталон» включені наступні ознаки: 1) Пріоритет задачі що виконувалась. 2) Складність задачі. 3) Оцінки навичок категорії H_1 . 4) Оцінки навичок категорії H_2 . 5) Оцінка рівню певного фаху 6) Оцінка досвіду виконавця таких завдань. 7) Оцінка рівня визначених знань Z_1 . 8) Завантаженість у проєкті. 9) Навичка $H_9, ..., 11)$ Оцінка навичок контролю. Засобами МХН [1, 8], а також процедурами редукції та кап-па Коена, необхідно встановити шаблон індивідуальних ознак Z_{ik} «виконавця,

клас K_i » у форматі $CF(A)$ для $(X_1, \dots, X_{11})$, що найкраще відповідає «еталону», наприклад, $(X_1=0.9, X_2=0.75, X_3= - 0.5, \dots, X_6= 1, \dots, X_{11}= - 0.3)$. Структура простору моделі класифікації формується окремо для кожного «еталону вимог», щоб забезпечити встановлену достовірність результату класифікації.

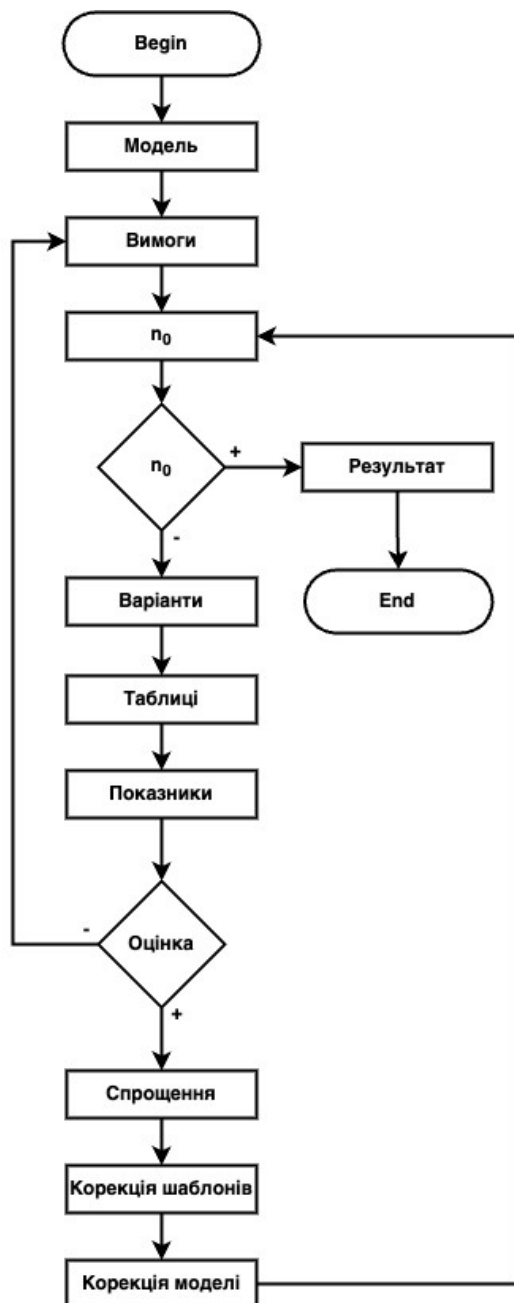


Рисунок 1 - Блок-схема алгоритму процедури редукції розмірності моделі

За схемою рис. 1 узагальнено алгоритм редукції на основі каппа статистики виконується таким чином: отримується Модель завдання класифікації за умов невизначеності даних; визначаються Вимоги щодо забезпечення точності та достовірності очікуваних результатів класифікації; за отриманими даними

розраховується гранична розмірність моделі n_0 ; перевіряється відповідність поточної моделі класифікації вимогам до моделі за значенням n_0 ; в разі виконання вимог – маємо достовірний Результат; у разі невиконання вимог починається аналіз простору і формування конкуруючих Варіантів спрощених моделей класифікації; для всіх сформованих варіантів реалізації конкуруючих моделей формуються таблиці розбіжностей; за таблицями розраховуються показники статистики каппа Коена та визначається загальна оцінка всіх варіантів; якщо така оцінка схожості конкуруючих моделей не відповідає вимогам достовірності, необхідно змінити вимоги; інакше виконується процедура спрощення (видалення певної структури параметрів моделі класифікації); при тому перевіряється необхідність і виконується корегування систем шаблонів і моделі класифікації.

Розроблено програмний комплекс, призначений для побудови математичної моделі процесів нечіткої класифікації – шаблонів баз знань для реалізації класифікації, розрахунків процесів класифікації на основі мережі МХН. Комплекс забезпечує автоматизацію обробки завдань нечіткої та CF(A) класифікації, щодо визначення класів вхідних елементів на основі їх неточно визначених ознак, прогнозування категорії «виконавця» та автору україномовного твору.

Висновки. У статті отримані удосконалені математичні моделі, алгоритми та програмні засоби, призначені для підвищення достовірності результатів класифікації при невизначених даних, представлених нечіткими величинами та коефіцієнтами упевненості CF(A). Удосконалення моделей класифікації з наведеними властивостями даних забезпечується шляхом модифікації нейронних мереж Хеммінга, а також застосування процедур редукції і статистики каппа Коена, які дозволяють сформувати математичні моделі зі встановленими ймовірнісними вимогами до результатів. В роботі запропоновані та реалізовані нові постановки, математичні моделі, а також виконано реалізацію моделей класифікації за нечіткими даними, що вирішують завдання із встановлення авторів україномовних текстів, а також завдання вибору кандидата за моделлю даних у форматі коефіцієнтів CF(A). Розроблено програмні засоби для формування моделей класифікації при невизначених даних процедурами редукції і каппа статистики.

Числові експерименти підтвердили достовірність результатів класифікації, а також ефективність запропонованої процедури редукції розмірності моделей на основі статистики каппа Коена.

ЛІТЕРАТУРА

1. Скалозуб В. В., Горячкін В. М., Клименко І. В., Терлецький І. А., Терленко А. П. Дослідження процедур мережі Хеммінга для управління сервісними системами при неточно визначених і природомовних даних. Наука та прогрес транспорту. 2022. № 3-4 (99-100). С. 33–47. DOI: <https://doi.org/10.15802/stp2022/276411> (in Ukrainian)
2. Великоіваненко Г. І. Оцінювання рівня економічної безпеки на підґрунті відстані Хеммінга. 2018. URL: <https://core.ac.uk/download/pdf/197269753.pdf>
3. Васильев В. И. Индукция и редукция в проблемах экстраполяции. Кибернетика и вычислительная техника. 1998. Вып. 116. С. 65–81.
4. Колесник А. С., Хайрова Н. Ф. Обґрунтування використання статистики каппа Коена в експериментальних дослідженнях NLP Text Mining. Кібернетика та системний аналіз. Т. 58, № 2. 2022. С. 143–153.
5. Li Min Fu, Shortliffe E. H. The application of certainty factors to neural computing for rule discovery. IEEE Transactions on Neural Networks. 2000. Vol. 11. Iss. 3. P. 647–657. DOI: <https://doi.org/10.1109/72.846736>
6. Скалозуб В. В., Горячкін В. М., Терлецький І. А. Багатопараметричні інтелектуальні процедури діагностування за неповними і збуреними даними // Логістика і транспортна безпека: Проблеми та перспективи розвитку в контексті аналізу сучасних викликів і загроз: матеріали доповідей II Міжнародної науково-практичної конференції. — Дніпро: Середняк Т.К., 2023. С. 42 – 47.
7. Шинкаренко В. І., Демидович І. М. Визначення ознак авторства природномовних текстів. Штучний інтелект. 2018. № 3. С. 27–35.
8. Richard A. Brualdi. Combinatorial matrix classes. — Cambridge: Cambridge University Press, 2006. — (Encyclopedia of Mathematics and Its Applications). — ISBN 0-521-86565-4
9. Leszek Rutkowski Metody i techniki sztucznej inteligencji. Naukowe PWN, Warszawa, 2005. – 520 p.
10. Haykin S. Neural networks: A Comprehensive Foundation. Prentice hall: New Jersey, 1999. 1103 p.

REFERENCES

1. Skalozub V.V., Goryachkin V.M., Klymenko I.V., Terletsky I.A., Terlenko A.P. Doslidzhennia protsedur merezhi khemminha dlia upravlinnia servisnymy systemamy pry netochno vyznachenyykh i pryrodomovnykh danykh. 2022. No. 3-4 (99-100). P. 33–47. DOI: <https://doi.org/10.15802/stp2022/276411>

2. Velykoivanenko, H.I. (2018). Otsiniuvannia rivnia ekonomichnoi bezpeky na pidgrunti vidstani Khemminha. Retrieved from <https://core.ac.uk/download/pdf/197269753.pdf> (in Ukrainian)
3. Vasilev, V.I. (1998). Induktsiya i reduktsiya v problemakh ekstrapolyatsii. Cybernetics and Computer Engineering, 116, 65-81. (in Russian).
4. Kolesnyk A. S., Khairova N. F. . Obgruntuvannia vykorystannia statystyky kappa Koena v eksperymentalnykh doslidzhenniakh NLP Text Mining Cybernetics and system analysis. Vol. 58, No. 2. 2022. P. 143–153.
5. Li Min Fu, Shortliffe E. H. The application of certainty factors to neural computing for rule discovery. IEEE Transactions on Neural Networks. 2000. Vol. 11. Iss. 3. P. 647–657. DOI: <https://doi.org/10.1109/72.846736>
6. Skalozub V.V., Goryachkin V.M., Terletskyi I.A. Bahatoparmetrychni intelektualni protsedury diahnostuvannia za nepovnymy i zburenymy danymy// Logistics and transport safety: Problems and prospects of development in the context of analysis of modern challenges and threats: materials of reports II International scientific and practical conference. — Dnipro: Serednyak T.K., 2023. P. 42-47.
7. Shinkarenko V. I., Demidovych I. M. Determination of signs of authorship of natural language texts. Artificial Intelligence. 2018. No. 3. P. 27–35.
8. Richard A. Brualdi. Combinatorial matrix classes. — Cambridge: Cambridge University Press, 2006. — (Encyclopedia of Mathematics and Its Applications). — ISBN 0-521-86565-4
9. Leszek Rutkowski Metody i techniki sztucznej inteligencji. Naukowe PWN, Warsaw, 2005. — 520 p.
10. Haykin S. Neural networks: A Comprehensive Foundation. Prentice hall: New Jersey, 1999. 1103 p.

Received 23.11.2023.

Accepted 24.11.2023.

***Formation classification models of undetermined data
by means of procedures reduction and kappa statistic***

The article is devoted to the development of mathematical models for the classification of uncertain data represented by fuzzy values and certainty factors $CF(A)$. Diagnostic pattern formation procedures use modified Hamming networks (MHN), as well as reduction methods and Cohen's kappa statistics. At the same time, the limiting dimensions and composition of the parameters of the classification model are determined, which ensure the established probabilistic requirements for the reliability of the calculation re-

sults. The model space reduction procedure and the structure of the software complex for diagnosing uncertain data are presented. An example of a classification model based on fuzzy data is the task of identifying the authors of Ukrainian-language texts. The classification task for data in CF(A) format corresponds to candidate selection. The results of the numerical modeling made it possible to establish the effectiveness, reliability and efficiency of the proposed procedures for the formation of reliable classification models with uncertain data.

Keywords: classification, reliable models, dimensionality of space, fuzzy values, CF(A) certainty factors, modified Hamming network, reduction procedure, Cohen's kappa statistic, Ukrainian-language texts, author's definition, computer simulation.

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

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

Терлецький Ігор Андрійович – аспірант каф. «Комп'ютерні інформаційні технології», Український державний університет науки і технологій, УДУНТ.

Дудник Ілля Петрович – магістрант каф. «Комп'ютерні інформаційні технології», Український державний університет науки і технологій, УДУНТ.

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

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

Terlitskyi Ihor – post-graduate student, Dep. “Computer and Information Technology”, Ukrainian State University of Science and Technology, USUST.

Dudnyk Ilya – graduate student, Dep. “Computer and Information Technology”, Ukrainian State University of Science and Technology, USUST

ЗМІСТ

CONTENTS

Жульковський О.О., Жульковська І.І., Костенко В.В., Булгакова О.Ф. Дослідження проблем синхронізації та захисту даних при реалізації багатопоточних програм	3	Zhulkovskyi O., Zhulkovska I., Kostenko V., Bulhakova O. Research on synchronization and data protection problems in implementing multithreaded programs	3
Герасимов В.В., Карпенко Н.В., Скуратовский И.А. Доступ к элементам структуры и неопределенное поведение кода на языке C	12	Gerasimov V.V., Karpenko N.V., Skuratovskyi I.A. Access to struct members and undefined behavior of C code	12
Кошель Є.В. Адаптація фреймворку WORLD для пофреймового аналізу мовлення в реальному часі	21	Koshel E. Adaptation of the WORLD frame- work for frame-by-frame real-time speech analysis	21
Гулий Т.О., Білозьоров В.Є. Використання методу нелінійного рекурентного аналізу до пошуку DDOS аномалій часових рядів мережевого трафіку	37	Hulyi T.A., Belozyorov V.Ye. Using the method of nonlinear recursive analysis for detecting DDOS anomalies in time series da- ta	37
Лук'янець М.О., Сулема Є.С. Ві- зуалізація даних у режимі реального часу для систем мереж ІоТ: виклики та стратегії оптимізації продуктивності	52	Lukianets M.O., Sulema Y.S. Re- al-time data visualization for IoT network systems: challenges and strategies for performance optimization	52
Песчанський В.Ю., Сулема Є.С. Роектування архітектури програмної системи для створення цифрових двійників медико-біологічних об'єктів	62	Peschanskii V., Sulema Y. Design of software system architecture for medical-biological objects digital twins creating	62
Ус С.А., Сергєєв О.С. Алгоритм розв'язання двоетапної неперервно-дискретної задачі ро- зміщення на прикладі оптимізації медичної логістики	71	Us S.A., Serhieiev O.S. An algorithm for solving a two- stage continuous-discrete location problem for medical logistics optimization	71

Іванов О.П., Скалозуб В.В., Горячкін В.М., Шинкаренко В.І. Визначення здатності штучного інтелекту до встановлення авторства художнього україномовного тексту за значними фрагментами	86	Ivanov O.P., Skalozub V.V., Horiachkin V.M., Shynkarenko V.I. Determining the ability of artificial intelligence to establish authorship of artistic ukrainian texts using significant fragments	86
Щербаков А.О., Ліхоузова Т.А. Програмне забезпечення розум- ного годинника для самостійних занять спортом та фітнесом	99	Shcherbakov A., Likhousova T. Smart watch software for independent sports and fitness	99
Кіріченко Л.О., Хацько Д.С., Зінченко П.П. Кластеризація траєкторій броунівського руху за допомогою машинного навчання	109	Kirichenko L., Khatsko D., Zinchenko P. Clustering Brownian motion tra- jectories using machine learning	109
Шедловський І.А., Гнатушенко В.В., Шедловська Я.І., Горєв В.М. Імітаційна модель плаского повітряного сонячного колектора	120	Shedlovsky I.A., Hnatushenko V.V., Shedlovska Y.I., Gorev V.M. Sim- ulation model of a flat air solar col- lector	120
Антонюк В.А., Сидорова М.Г. Концепція асоціативного графічного інтерфейсу у системі автоматизації робочих процесів	133	Antonyuk V.A., Sydorova M.G. Concept of associative graphic interface in work process automation system	133
Скалозуб В.В., Горячкін В.М., Терлецький І.А., Дудник І.П. Формування моделей класифікації невизначених даних процедурами редукції і каппа статистики	141	Skalozub V.V., Horiachkin V.M., Terlitskyi I.A., Dudny I.P. Formation classification models of undetermined data by means of procedures reduction and kappa statistic	141

РЕФЕРАТИ

УДК 004.4'6

Жульковський О.О., Жульковська І.І., Костенко В.В., Булгакова О.Ф. **Дослідження про блем синхронізації та захисту даних при реалізації багатопоточних програм** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 5(148). – Дніпро, 2023. – С.3 – 11.

Проблематика спільного використання даних потоками особливо актуальна в сучасних багатоядерних та багатопроцесорних системах. Основними проблемами реалізації багатопоточних програм є стан гонки (race conditions), взаємне блокування (deadlocks), голодування потоків (resource starvation). Метою роботи є вирішення проблеми гонки потоків при багатопоточних обчисленнях ресурсномістких задач із паралельним доступом до спільних даних за допомогою використання відповідних механізмів синхронізації, таких як м'ютекси. Розроблено та досліджено багатопоточний алгоритм реалізації типової задачі обробки надвеликих масивів даних із захистом критичної області у конкурентних програмах, що виконуються на багатопроцесорних та багатоядерних системах.

Бібл. 10, іл. 2.

УДК 004.052

Герасимов В.В., Карпенко Н.В., Скуратовский И.А. **Доступ к элементам структуры и не определенное поведение кода на языке С** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 5(148). – Дніпро, 2023. – С.12 – 20.

Исследован доступ к переменным типа данных структура на языке программирования С. Проанализированы особенности неопределенного поведения кода после компилятора Visual Studio 2019 под управлением Windows 10 и компилятора gcc под управлением ОС Linux. Показано результаты пренебрежения сообщениями при компиляции программ, при которых последние работают некорректно. Даны рекомендации по избежанию этой проблемы.

Бібл. 7, ил. 2.

УДК 519.85:004.42

Кошель Є. В. **Адаптація фреймворку WORLD для пофреймового аналізу мовлення в реальному часі** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 5(148). – Дніпро, 2023. – С.21 – 36.

WORLD – це система для синтезу мовлення на основі вокодера, яка була розроблена М. Морісом та ін. і реалізована на C++. Було продемонстровано, дана система має високу ефективність та точність у порівнянні з аналогічними системами. Однак вона виявилася не придатною для використання у певних сценаріях, наприклад, при потоковій обробці аудіо фрейм за фреймом. Ця стаття розглядає недоліки C++ імплементації системи WORLD та пропонує модифіковані версії її складових алгоритмів для вирішення виявлених проблем. Результуючий фреймворк було протестовано на синтетичних сигналах та на реальних записах мовлення.

Бібл. 17, іл. 7, табл. 1.

УДК 519.2:004.9

Гулий Т.О., Білозьоров В.Є. **Використання методу нелінійного рекурентного аналізу до пошуку DDOS аномалій часових рядів мережевого трафіку** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 5(148). – Дніпро, 2023. – С.37 – 51.

У статті розглядається питання використання методу нелінійного рекурентного аналізу до проблеми пошуку аномалій у мережевому трафіку, що надана у вигляді даних часових рядів знімку трафіку за певний період часу. Описано методика визначення якісної характеристики для вхідних даних та її використання для побудови відповідної рекурентної діаграми (recurrence plots, RP). В якості яких було запропоновано узяти статистику кількості отриманих байтів у секунду. Для отримання чисельних значень показників RP пропонується використувати середовище Матлаб та розроблений для цього пакет crptools. Наведені обраховані показники RP дозволили здійснити типізацію отриманих даних та визначити тип якій отримав назву «DDOS RP», що надає змогу виділити деякі види атак типу DoS/DDoS.

Бібл. 9, іл.8, табл. 3.

УДК 004.51

Лук'янець М.О., Сулема Є.С. **Візуалізація даних у режимі реального часу для систем мереж IoT: виклики та стратегії оптимізації продуктивності** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 5(148). – Дніпро, 2023. – С.52 – 61.

У цій статті досліджуються актуальні проблеми візуалізації даних у режимі реального часу та пропонуються стратегії оптимізації продуктивності. Автори аналізують нещодавні дослідження та публікації, щоб ідентифікувати проблеми продуктивності, які впливають на візуалізацію даних у режимі реального часу, та обговорюють потенційні рішення цих проблем. У статті пропонуються такі підходи, як агрегація даних, попереднє оброблення та стиснення, щоб зменшити кількість оброблюваних даних, а також кешування, індексацію та паралельне оброблення, щоб покращити швидкість та точність візуалізацій. Застосовуючи ці стратегії, користувачі можуть приймати більш обґрунтовані рішення в режимі реального часу, поліпшуючи ефективність процесу прийняття рішень.

Бібл. 9, іл. 2.

УДК 004.4 : 004.94

Песчанський В.Ю., Сулема Є.С. **Роектування архітектури програмної системи для створення цифрових двійників медико біологічних об'єктів** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 5(148). – Дніпро, 2023. – С.62 – 70.

У статті запропоновано універсальну об'єктно орієнтовану архітектуру програмної системи, призначеної для створення двійників медико біологічних об'єктів на прикладі отоларингології. Архітектура ґрунтується на використанні поліморфізму для мінімізації роботи розробників програмного забезпечення у разі виникнення необхідності модифікації системи для використання в інших галузях медицини. Розглянуто групи компонентів системи та надано рекомендації, щодо їхнього розроблення.

Бібл. 8.

УДК 519.85:004.42+004.02

Ус С.А., Сергєєв О.С. **Алгоритм розв'язання двоетапної неперервно дискретної задачі розміщення на прикладі оптимізації медичної логістики** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 5(148). – Дніпро, 2023. – С.71 – 85.

У роботі розглядається оптимізація логістики з акцентом на медичну сферу у критичних випадках. Запропоновано математичну модель та алгоритм розв'язання двоетапної неперервно дискретної задачі розміщення, специфічної для зазначеної галузі. Розроблений підхід поєднує застосування генетичних алгоритмів з теорією оптимального розбиття множин та реалізований з використанням мов програмування Python3 та C++. Роботу алгоритму проілюстровано на репрезентативній модельній задачі.

Бібл. 20, іл. 3, табл. 0.

УДК 004.8: 821

Іванов О.П., Скалозуб В.В., Горячкін В.М., Шинкаренко В.І. **Визначення здатності штучного інтелекту до встановлення авторства художнього україномовного тексту за значними фрагментами** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 5(148). – Дніпро, 2023. – С.86 – 98.

Інтелектуальна пошукова система Bing, може служити інструментом для визначення авторства художнього тексту українською мовою. Ця робота має на меті дослідити ефективність визначення авторства художніх текстів за допомогою сучасних інструментів штучного інтелекту на основі фрагментів творів значної довжини. Встановлено, що довгі фрагменти дозволяють з високою точністю визначити автора художнього україномовного тексту. Іван Франко має найвищий відсоток відповідей, де було згадано ім'я автора та назва твору 87%.

Бібл. 11, іл. 4, табл. 3.

УДК 004.89

Щербаков А.О., Ліхоузова Т.А. **Програмне забезпечення розумного годинника для самостійних занять спортом та фітнесом** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 5(148). – Дніпро, 2023. – С.99 – 108.

За наявною статистикою, найбільш популярним розумним годинником серед користувачів є Apple Watch. На сьогодні існує більше 100 мільйонів унікальних користувачів цього девайсу, 75% з яких використовують його для занять спортом. Маючи широкий набір різноманітних датчиків для відстеження фізичних параметрів користувача, ані компанія Apple, ані сторонні розробники поки не розробили програмного забезпечення для систематизації усіх зібраних даних задля покращення фізичних параметрів спортсмена та досягнення особистих спортивних цілей.

Метою дослідження є пошук можливості покращення фізичних параметрів починаючого спортсмена за рахунок комплексного аналізу зібраних розумним годинником даних його активності та створенні на основі цих даних більш персоналізованих рекомендацій.

Запропоновано програмне забезпечення – застосунок для розумних годинників, який аналізує зібрану статистику тренувань та на її основі надає рекомендації по проведенню поточного тренування.

Бібл. 10.

УДК 519.2:004.9

Кіріченко Л.О., Хацько Д.С., Зінченко П.П. **Кластеризація траєкторій броунівського руху за допомогою машинного навчання** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 5(148). – Дніпро, 2023. – С.109 – 119.

Стаття присвячена виявленню пасток, які захоплюють броунівську частинку, за допомогою машинного навчання. Проведено моделювання траєкторії броунівської частинки за допомогою моделі броунівського руху із дрейфом, що охоплює як вільну дифузію, так і рух у пастці. Для кластеризації траєкторії руху використовується метод кластеризації на основі щільності DBSCAN. Універсальність цього метода дозволяє ідентифікувати кластери без попереднього знання їхньої кількості або форми, що робить його придатним для виявлення пасток. Проведене дослідження показує, що застосування методу DBSCAN дозволяє досягнути середньої точності 95.0%..

Бібл. 23, рис.3, табл.0.

УДК 004.9

Шедловський І.А., Гнатушенко В.В., Шедловська Я.І., Горєв В.М. **Імітаційна модель плаского повітряного сонячного колектора** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 5(148). – Дніпро, 2023. – С.120 – 132.

Проведено аналіз використання повітряних сонячних колекторів у системах опалення. На основі аналізу існуючих досліджень зроблено висновок щодо необхідності створення імітаційної моделі сонячного колектора з урахуванням його використання у складі сучасних комп'ютерних інформаційних та керуючих систем. З використанням відомих залежностей, що описують процеси теплопередачі у сонячних колекторах, показано, що імітаційна модель складається з послідовно з'єднаних компонентів. Один із яких визначає нелінійність об'єкта. Можливість була використана для обґрунтування послідовного включення нелінійної частини моделі та її інерційної складової. Це дозволяє експериментально визначити параметри динаміки нагріву повітря. Практичне використання розробленої моделі в інформаційно керуючій системі потребує постійного контролю температури повітря на вході та виході колектора, а також контролю за потужністю сонячної радіації.

Бібл. 9, рис. 5.

УДК 004.42

Антонюк В.А., Сидорова М.Г. **Концепція асоціативного графічного інтерфейсу у системі автоматизації робочих процесів** // Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 5(148). – Дніпро, 2023. – С.133 – 140.

Стаття розглядає проблему розробки асоціативного графічного інтерфейсу для систем автоматизації робочих процесів. Автори виклали нову концепцію інтерфейсу, що забезпечує оптимальну ефективність та зручність використання, на основі аналізу сучасних технологій та методів дослідження. Було використано плагінну архітектуру, яка

ISSN 1562-9945 (Print) 161
ISSN 2707-7977 (Online)

дозволяє легко інтегрувати нові функції та оптимізувати роботу системи. Також було розроблено Draw & GO – новий інструмент для автоматизації робочих процесів.

Бібл. 10, іл. 3

УДК 004.6: 007.5: 004.94

Скалозуб В.В., Горячкін В.М., Терлецький І.А., Дудник І.П. **Формування моделей класифікації невизначених даних процедурами редукції і Каппа статистики //** Системні технології. Регіональний міжвузівський збірник наукових праць. Випуск 5(148). – Дніпро, 2023. – С.141 – 155.

Стаття присвячена розвитку математичних моделей класифікації невизначених даних, представлених нечіткими величинами та коефіцієнтами упевненості $CF(A)$. Процедури формування шаблонів діагностування використовують модифіковані мережі Хеммінга (МХН), а також методи редукції та статистики каппа Коена. При цьому визначаються граничні розмірності та склад параметрів моделі класифікації, які забезпечують встановлені ймовірнісні вимоги достовірності результатів розрахунків. Представлена процедура редукції простору моделі діагностування невизначених даних. У статті наведено постановки, математичні моделі та реалізації завдань класифікації за недетермінованими даними. Прикладом моделі класифікації за нечіткими даними являється завдання із встановлення авторів україномовних текстів. Завдання класифікації при даних у форматі $CF(A)$ відповідає відбору кандидата. Результати числового моделювання дозволили встановити результативність, достовірність та ефективність запропонованих процедур формування достовірних моделей класифікації при невизначених даних.

Бібл. 10.

UDC 004.4'6

Zhulkovskyi O., Zhulkovska I., Kostenko V., Bulhakova O. **Research on synchronization and data protection problems in implementing multithreaded programs** // System technologies. N 5(148) Dnipro, 2023. P.3 – 11.

The issue of shared data usage by threads is especially relevant in modern multi core and multiprocessor systems. The main problems of implementing multithreaded programs are race conditions, deadlocks, and thread starvation. The aim of the work is to solve the problem of thread racing in multithreaded calculations of resource intensive tasks with parallel access to shared data using appropriate synchronization mechanisms, such as mutexes. A multithreaded algorithm for implementing a typical task of processing large data arrays with protection of the critical area in concurrent programs running on multiprocessor and multi core systems has been developed and researched.

Ref. 10, im. 2.

UDC 004.052

Gerasimov V.V., Karpenko N.V., Skuratovskyi I.A. **Access to struct members and undefined behavior of C code** // System technologies. N 5(148) Dnipro, 2023. P.12 – 20.

During software development, novice developers usually receive a lot of error messages and just warnings of various kinds. And if the code simply won't run when there are errors, then the program usually starts when there are various warnings. And here it is important to understand what consequences the presence of warnings of various kinds can lead to.

This work aims to study the code's undefined behavior when working with structures in the C programming language when issuing a corresponding compiler warning about returning the address of a local or temporary variable.

In the procedural programming language C, there is an ancestor of the OOP class – the structure struct, which encapsulates only the state of the entity. And the question arises – is it possible to work with separate components fields of such a structure analogously to OOP languages?

For research, a simple structure was taken, which contains information about the person's name, surname, and phone number. To access parts of the structure, pseudogetters were used – functions that returned a pointer to the corresponding part of the structure.

The research was conducted in the Visual Studio 2019 environment under the control of the Windows 10 operating system when the default C language standard – MSVC and the more modern ISO standard C17 (2018) was selected in the project settings.

As a result, a truly undefined behavior of the code was obtained, when the result of the work of the code fragment (function call) depends on many factors: the length of the array, the standard of the C language, the position of a certain part in the structure.

An attempt to conduct similar research under the control of the Linux Mint operating system using the gcc compiler version 5.4 was unsuccessful. When compiling the code, a similar warning about returning the address of a local variable was also issued, as in the case of Visual Studio. But when the program was launched, it simply crashed with a message about a segmentation error.

Thus, both the Visual Studio 2019 compiler and the gcc compiler warned us about undefined code behavior. But this uncertain behavior was radically different for operating systems and compilers. If after gcc under the Linux OS, the code simply does not work at all and the program stops its work with a segmentation error message, then after Visual Studio under Windows, inexperienced developers with improper testing and verification of the code can miss the code that "does not always work", which can lead to unexpected results, not always pleasant, to say the least. And that's why software developers, especially beginners, should pay attention not only to compilation errors but also to warnings, even if the code works.

Bibl. 7, ill. 2.

UDC 519.85:004.42

Koshel E. **Adaptation of the WORLD framework for frame by frame real time speech analysis** // System technologies. N 5(148) Dnipro, 2023. P.21 – 36.

WORLD is a vocoder based speech synthesis system developed by M. Morise et al. and implemented in C++. It was demonstrated to have improved performance and accuracy when compared to other algorithms. However, it turned out to not perform well in certain scenarios, particularly, when applying the framework to very short waveforms on a frame by frame basis. This paper reviews the issues of the C++ implementation of WORLD and proposes modified versions of its constituting algorithms that attempt to mitigate those issues. The resulting framework is tested on both synthetic signals and on real recorded speech.

Bible 17, ill. 7, tab. 1.

UDC 519.2:004.9

Hulyi T.A., Belozyorov V.Ye. **Using the method of nonlinear recursive analysis for detecting DDOS anomalies in time series data** // System technologies. N 5(148) Dnipro, 2023. P.37 – 51.

The paper considers the issue of using the method of nonlinear recurrent analysis to the problem of finding anomalies in network traffic, which is provided in the form of time series data of a traffic snapshot for a certain period of time. The method of determining a qualitative characteristic for input data and its use for constructing the corresponding recurrence plot (RP) is described. As such, it was suggested to take statistics of the number of received bytes per second. To obtain numerical values of RP indicators, it is suggested to use the Matlab environment and the crptools package developed for this purpose. The given calculated RP indicators made it possible to typify the received data and determine the type, which was named "DDOS RP", which makes it possible to distinguish some types of DoS/DDoS type attacks.

Bible 9, ill.8, tab. 2.

UDC 004.51

Lukianets M.O., Sulema Y.S. **Real time data visualization for IoT network systems: challenges and strategies for performance optimization** // System technologies. N 5(148) Dnipro, 2023. P.52 – 61.

This paper explores the challenges of real time data visualization and proposes strategies to optimize performance. The authors analyze recent research and publications to identify performance issues that affect real time data visualization and discuss potential solutions to these issues. The paper suggests techniques such as data aggregation, pre processing, and compression to reduce the amount of data being processed, and caching, indexing, and parallel processing to improve the speed and accuracy of visualizations. By implementing these strategies, decision makers can make more informed decisions in real time, improving the overall efficiency and effectiveness of the decision making process.

Lib. 9, fig. 2.

UDC 004.4 : 004.94

Peschanskii V., Sulema Y. **Design of software system architecture for medical biological objects digital twins creating** // System technologies. N 5(148) Dnipro, 2023. P.62 – 70.

The paper presents a universal object oriented architecture of a software system designed to create digital twins of medical biological objects using the example of otolaryngology. The architecture is based on the used polymorphism to minimize the work of software developers in the case of the need to modify systems for use in other fields of medicine. The groups of system components are considered and recommendations for their development are provided.

Bible 8.

UDC 519.85:004.42+004.02

Us S.A., Serhieiev O.S. **An algorithm for solving a two stage continuous discrete location problem for medical logistics optimization** // System technologies. N 5(148) Dnipro, 2023. P.71 – 85.

The research paper examines optimizing logistics in the healthcare sector, focusing on medical logistics for public health, especially during emergencies. It introduces a mathematical model and algorithm for a two stage continuous discrete location problem specific to medical logistics. The algorithm integrates genetic methods with optimal partition of sets theory and is validated through a software application on a representative model problem.

Bible 20, ill. 3, tab. 0.

UDC 004.8: 821

Ivanov O.P., Skalozub V.V., Horiachkin V.M., Shynkarenko V.I. **Determining the ability of artificial intelligence to establish authorship of artistic ukrainian texts using significant fragments** // System technologies. N 5(148) Dnipro, 2023. P.86 – 98.

The Bing intellectual search system can serve as a tool for determining the authorship of artistic texts in Ukrainian. This work aims to explore the effectiveness of determining the authorship of artistic texts using modern artificial intelligence tools based on fragments of significant length. It has been established that long fragments allow for high precision identification of the author of an artistic Ukrainian text. Ivan Franko has the highest percentage of responses where the author's name and title of the work were mentioned, at 87%.

Bible 11, ill. 4, tab. 3.

UDC 004.89

Shcherbakov A., Likhousova T. **Smart watch software for independent sports and fitness** // System technologies. N 5(148) Dnipro, 2023. P.99 – 108.

The purpose of the research is to find the possibility of improving the physical parameters of a novice athlete by means of a comprehensive analysis of his activity data collected by a smart watch and creating more personalized recommendations during training based on this data.

According to available statistics, the most popular smart watch among users is the Apple Watch. With a wide array of different sensors to track a user's physical parameters, neither Apple nor third party developers have yet developed software to systematize all the collected data to improve an athlete's physical parameters and achieve personal athletic goals. There are many fitness apps available for the Apple Watch, each with its own unique features and features. Unfortunately, none of the analyzed applications provide sufficient information regarding the correct execution of training and the collection of indicators, which is a disadvantage for users who want to do sports without risks to their health.

Bible 10.

UDC 519.2:004.9

Kirichenko L., Khatsko D., Zinchenko P. **Clustering Brownian motion trajectories using machine learning** // System technologies. N 5(148) Dnipro, 2023. P.109 – 119.

The article is dedicated to detecting traps encountered by a Brownian particle based on machine learning methods. The trajectory of the Brownian particle was modeled using a drift extended Brownian motion model, encompassing both free diffusion and particle movement within a trap. The density based spatial clustering of applications with noise (DBSCAN) method was employed for clustering the motion trajectory. The versatility of this method allows the identification of clusters without prior knowledge of their quantity or shape, making it suitable for trap detection. The conducted research demonstrates that the application of the DBSCAN method achieves an average accuracy of 95.0%

Ref.23, fig.3, tab.0.

UDC 004.9

Shedlovsky I.A., Hnatushenko V.V., Shedlovska Y.I., Gorev V.M. **Simulation model of a flat air solar collector** // System technologies. N 5(148) Dnipro, 2023. P.120 – 132.

An analysis of the use of air solar collectors in heating systems was carried out. Based on the analysis of existing research, a conclusion was made about the need for a simulation model of the solar collector, taking into account its use as part of modern computer information and control systems. With the use of known dependencies describing heat transfer processes in solar collectors, it is shown that the simulation model consists of sequentially connected components. One of which determines the non linearity of the object. The opportunity was used to justify the sequential inclusion of the nonlinear part of the model and its inertial component. This provides an opportunity to experimentally determine the parameters

of air heating dynamics. The practical use of the developed model in the information and control system requires continuous monitoring of the air temperature at the inlet and outlet of the collector, as well as monitoring of the power of solar radiation.

Bibl. 9, Fig. 5.

UDK 004.42

Antonyuk V.A., Sydorova M.G. **Concept of associative graphic interface in work process automation system** // System technologies. N 5(148) Dnipro, 2023. P.133 – 140.

The article deals with the problem of developing an associative graphical interface for work process automation systems. The authors presented a new interface concept that ensures optimal efficiency and ease of use, based on the analysis of modern technologies and research methods. A plugin architecture was used, which allows for easy integration of new features and optimization of system operation. Also, Draw & GO was developed a new tool for work process automation.

Ref. 10, Ill. 3.

UDK 004.6: 007.5: 004.94

Skalozub V.V., Horiachkin V.M., Terlitskyi I.A., Dudny I.P. **Formation classification models of undetermined data by means of procedures reduction and kappa statistic** // System technologies. N 5(148) Dnipro, 2023. P.141 – 155.

The article is devoted to the development of mathematical models for the classification of uncertain data represented by fuzzy values and certainty factors $CF(A)$. Diagnostic pattern formation procedures use modified Hamming networks (MHN), as well as reduction methods and Cohen's kappa statistics. At the same time, the limiting dimensions and composition of the parameters of the classification model are determined, which ensure the established probabilistic requirements for the reliability of the calculation results. The model space reduction procedure and the structure of the software complex for diagnosing uncertain data are presented. An example of a classification model based on fuzzy data is the task of identifying the authors of Ukrainian language texts. The classification task for data in $CF(A)$ format corresponds to candidate selection. The results of the numerical modeling made it possible to establish the effectiveness, reliability and efficiency of the proposed procedures for the formation of reliable classification models with uncertain data.

Ref. 10.

Системні технології

ЗБІРНИК НАУКОВИХ ПРАЦЬ

Випуск 5 (148)

Головний редактор: к.т.н., доц. Т.В. Селівьорстова

Технічний редактор та секретар збірки: к.т.н., доц. К.Ю. Островська

Здано до набору 27.11.2023. Підписано до друку 30.11.2023.

Формат 60x84 1/16. Друк - різнограф. Папір типограф.

Умов. друк арк. – 11,93. Обл.–видавн. арк. – 10,438.

Тираж 300 прим. Замовл. – 05/23

Український державний університет науки і технологій,
ННІ «Інститут промислових та бізнес технологій»,
кафедра Інформаційних технологій та систем: ІВК «Системні технології»
49600, Дніпро, а/с 493

<http://journals.nmetau.edu.ua/index.php/st>

Свідоцтво про державну реєстрацію друкованого засобу масової інформації:

Серія КВ № 8684 від 23 квітня 2004 рік

Редакційна колегія

Селівьорстова Тетяна Віталіївна
(головний редактор)

доцент, кандидат технічних наук

Алпатов Анатолій Петрович

Член-кореспондент НАН України,
професор, доктор технічних наук

Архипов Олександр Євгенійович

професор, доктор технічних наук

Бабічев Сергій Анатолійович

доцент, доктор технічних наук

Білозьоров Василь Євгенович

професор,

доктор фізико-математичних наук

Гече Федір Елемирович

професор, доктор технічних наук

Гуда Антон Ігорович

(заст. головного редактора)

професор, доктор технічних наук

Гнатушенко Вікторія Володимирівна

(вчений секретар)

професор, доктор технічних наук

Гнатушенко Володимир Володимирович

професор, доктор технічних наук

Гожий Олександр Петрович

професор, доктор технічних наук

Єрьомін Олександр Олегович

професор, доктор технічних наук

Кіріченко Людмила Олегівна

професор, доктор технічних наук

Світличний Дмитро Святозарович

професор, доктор технічних наук

Скалозуб Владислав Васильович

професор, доктор технічних наук

Хандецький Володимир Сергійович

професор, доктор технічних наук

Український державний університет науки і технологій, ННІ «Інститут промислових та бізнес технологій», Україна

Інститут технічної механіки

НАНУ і ДКАУ, Україна

Національний технічний університет

України «Київський політехнічний інститут» імені Ігоря Сікорського», Україна

Jan Evangelista Purkyně University
in Ústí nad Labem

Університет імені Яна Євангеліста Пуркіне,
Усті над Лабем, Чеська Республіка

Дніпровський національний університет
імені Олеся Гончара, Україна

Ужгородський національний університет,
Україна

Український державний університет науки і технологій, ННІ «Інститут промислових та бізнес технологій», Україна

Український державний університет науки і технологій, ННІ «Інститут промислових та бізнес технологій», Україна

Національний технічний університет
«Дніпровська політехніка», Україна

Чорноморський національний університет
імені П.Могили, Україна

Український державний університет науки і технологій, ННІ «Інститут промислових та бізнес технологій», Україна

Харківський національний університет
радіоелектроніки, Україна

Akademia Górniczo-Hutnicza

Краківська гірничо-металургійна академія
ім. С. Сташіца, Польща

Український державний університет науки і технологій, ННІ «Дніпровський інститут інфраструктури і транспорту» Україна

Дніпровський національний університет
імені Олеся Гончара, Україна