

**Шулімова Дар'я Денисівна***Державний університет інформаційно-комунікаційних технологій, Київ*

ORCID 0009-0002-9557-990X

**Бойко Анна Олександрівна***Державний університет інформаційно-комунікаційних технологій, Київ*

ORCID 0009-0001-3709-6283

**Мурзін Ігор Васильович***Державний університет інформаційно-комунікаційних технологій, Київ*

ORCID 0009-0006-5162-3875

**Довженко Тимур Павлович***Державний університет інформаційно-комунікаційних технологій, Київ*

ORCID 0000-0002-0352-8391

## АЛГОРИТМІЧНІ ПІДХОДИ ДО ВИЯВЛЕННЯ АНОМАЛІЙ НА ОСНОВІ МАШИННОГО НАВЧАННЯ

**Анотація.** Виявлення аномалій у кібербезпеці є критично важливим процесом, спрямованим на ідентифікацію нетипових шаблонів поведінки або активності, які суттєво відрізняються від усталених, нормальних операційних процедур у межах інформаційної системи або комп'ютерної мережі. Ці відхилення від норми можуть слугувати ранніми індикаторами потенційних загроз безпеці, починаючи від спроб несанкціонованого проникнення та розповсюдження шкідливого програмного забезпечення й закінчуючи експлуатацією існуючих вразливостей у програмному забезпеченні або конфігурації системи. Своєчасне та ефективне виявлення таких аномальних подій надає можливість оперативно реагувати на загрози, запобігати їхньому подальшому розвитку та мінімізувати потенційні ризики для конфіденційності, цілісності та доступності критично важливих даних і цифрової інфраструктури організації.

Сучасні системи виявлення аномалій у кібербезпеці все частіше використовують складні алгоритми машинного навчання для ефективного розпізнавання складних і ледь помітних патернів поведінки. На відміну від традиційних методів, які значною мірою базуються на заздалегідь визначених правилах і статичних порогових значеннях, алгоритми машинного навчання мають здатність до навчання на великих обсягах даних, що дозволяє їм виявляти нові та раніше невідомі типи атак, з якими статичні правила можуть не впоратися. У випадках, коли кіберзлочинці постійно вдосконалюють свої методи та інструменти, здатність алгоритмів машинного навчання адаптуватися до нових загроз, аналізуючи дані про попередні атаки та нормальну поведінку, стає надзвичайно цінною. Ці алгоритми можуть виявляти ледь помітні відхилення, які можуть бути пропущені системами, що базуються на жорстких правилах, тим самим підвищуючи загальну ефективність систем виявлення вторгнень і запобігання витоку даних. Використання машинного навчання дозволяє будувати більш інтелектуальні та проактивні системи безпеки, здатні ефективно протистояти сучасним кіберзагрозам.

**Ключові слова:** виявлення аномалій, машинне навчання, інформаційна безпека, Z-score, статистичні методи, нейронні мережі, кібератаки.

**Shulimova Daria**

*State University of Information and Communication Technologies, Kyiv*  
ORCID 0009-0002-9557-990X

**Boiko Anna**

*State University of Information and Communication Technologies, Kyiv*  
ORCID 0009-0001-3709-6283

**Murzin Ihor**

*State University of Information and Communication Technologies, Kyiv*  
ORCID 0009-0006-5162-3875

**Dovzhenko Tymur**

*State University of Information and Communication Technologies, Kyiv*  
ORCID 0000-0002-0352-8391

## ALGORITHMIC APPROACHES TO ANOMALIES DETECTION BASED ON MACHINE LEARNING

**Abstract.** *Anomaly detection in cybersecurity is a critically important process aimed at identifying atypical patterns of behavior or activity that significantly differ from established, normal operating procedures within an information system or computer network. These deviations from the norm can serve as early indicators of potential security threats, ranging from unauthorized intrusion attempts and malware distribution to exploitation of existing vulnerabilities in software or system configuration. Timely and effective detection of such anomalous events provides an opportunity to respond quickly to threats, prevent their further development, and minimize potential risks to the confidentiality, integrity, and availability of critical data and the organization's digital infrastructure.*

*Modern cybersecurity anomaly detection systems increasingly use sophisticated machine learning algorithms to effectively recognize complex and subtle patterns of behavior. Unlike traditional methods that rely heavily on predefined rules and static thresholds, machine learning algorithms have the ability to learn from large amounts of data, which allows them to detect new and previously unknown types of attacks that static rules may not be able to handle. In cases where cybercriminals are constantly improving their methods and tools, the ability of machine learning algorithms to adapt to new threats by analyzing data on previous attacks and normal behavior becomes extremely valuable. These algorithms can detect subtle deviations that may be missed by systems based on strict rules, thereby increasing the overall effectiveness of intrusion detection and data leakage prevention systems. The use of machine learning allows you to build more intelligent and proactive security systems that can effectively counter modern cyber threats.*

**Keywords:** *anomaly detection, machine learning, information security, Z-score, statistical methods, neural networks, cyberattacks.*

### 1. Постановка проблеми

Виявлення аномалій відіграє вирішальну роль у різних сферах, висвітлюючи потенційні загрози, помилки або важливі події. Ці відхилення можуть сигналізувати про критичні інциденти, такі як шахрайські дії, порушення кібербезпеки або збої в складних системах [1]. Враховуючи непрактичність ручної перевірки великих наборів даних, автоматизовані методи виявлення аномалій стали незамінними. Традиційні методи виявлення шкідливих програм здебільшого базуються на використанні антивірусного програмного забезпечення. Ці програми сканують файли та папки на пристрої, шукаючи відомі загрози. Якщо виявляється збіг із базою даних шкідливих програм, файл позначається як небезпечний. Такий підхід, відомий як сканування за сигнатурами, має низку обмежень. Він може виявляти лише відомі загрози, тому нове або

модифіковане шкідливе програмне забезпечення часто залишається непоміченим. Крім того, цей процес може бути тривалим і сповільнювати роботу пристрою. Інша проблема полягає в тому, що традиційні методи працюють реактивно – вони виявляють загрозу вже після того, як вона проникла у систему. Через особливості веб-протоколів засоби кібербезпеки часто не мають повного доступу до інформації, яка передається між користувачем і веб-сайтом, що додатково ускладнює виявлення загроз. Отже, хоча традиційні методи й залишаються корисними, вони не забезпечують повного захисту.

Тому дослідження та впровадження алгоритмів машинного навчання для виявлення аномалій є актуальною темою дослідження.

## **2. Аналіз останніх досліджень і публікацій**

Штучний інтелект, зокрема підгалузь машинного навчання, став трансформаційною силою в розширенні можливостей виявлення аномалій. Алгоритми машинного навчання можуть автоматично розпізнавати несподівані зміни в нормальній поведінці даних, пропонуючи значну перевагу перед традиційними статистичними підходами[1]. На відміну від статичних систем на основі правил або методів, які покладаються на попередньо визначені статистичні розподіли, виявлення аномалій на основі машинного навчання може адаптуватися до змінних шаблонів даних і ефективно обробляти складні багатовимірні набори даних. Ця адаптивність і надійність роблять машинне навчання все більш важливим інструментом для виявлення аномалій у широкому діапазоні програм.

У статті, представленій авторами [9], розглядається, як методи виявлення вторгнень на основі штучного інтелекту (ШІ), особливо машинного та глибокого навчання, підвищують ефективність аналізу трафіку та виявлення аномалій для захисту інформаційних систем організацій. Дослідження виявило значне покращення точності та ефективності при використанні методів виявлення вторгнень на основі ШІ порівняно з традиційними системами, хоча й були виявлені певні обмеження, такі як складність одночасного виявлення множинних атак, що вказує на необхідність подальших досліджень у цьому напрямку.

У дослідженні [2] автори описують, що згорткові нейронні мережі (CNN), особливо в поєднанні з шарами довгої короткочасної пам'яті (LSTM), демонструють вищу продуктивність у виявленні аномалій порівняно з рекурентними нейронними мережами (RNN). Це підкреслює важливість врахування як просторових, так і часових характеристик мережевого трафіку при розробці систем виявлення аномалій на основі машинного навчання. Результати дослідження підтверджують перспективність використання технологій глибокого навчання для кіберзахисту інформаційних систем, враховуючи їхню здатність до автоматичного вилучення ознак та адаптації до складних закономірностей у даних.

Практичний досвід підтверджує ефективність використання штучного інтелекту (ШІ) для виявлення веб-шкідливих програм. Наприклад, компанія DEKENEAS NEXT-GEN TECHNOLOGIES SRL застосовує алгоритми машинного навчання для аналізу поведінки веб-елементів. Система була навчена розпізнавати шкідливі патерни, що дозволяє їй виявляти загрози ще до того, як вони завдадуть шкоди. Результати впровадження показали високу точність у виявленні та нейтралізації шкідливих програм. Це не лише підвищило рівень безпеки клієнтів, але й значно зменшило витрати часу та ресурсів на ручне виявлення загроз. Крім того, система має можливість самостійно оновлюватися на основі нових загроз, що підвищує її ефективність у довгостроковій перспективі. Такий досвід демонструє, що ШІ відіграє критичну роль у зміцненні захисту від кіберзагроз.

### 3. Мета і задачі дослідження

Метою даного дослідження є аналіз сучасних методів машинного навчання, які застосовуються у виявленні аномалій в інформаційних системах організацій та аналіз їх ефективності у виявленні різноманітних типів кібератак, включаючи як відомі, так і нові, невідомі раніше загрози, а також оцінка їхньої здатності до адаптації до кіберзагроз. Крім того, дослідження спрямоване на виявлення потенційних напрямків для подальшого розвитку та вдосконалення методів машинного навчання з метою підвищення ефективності, точності та стійкості систем виявлення аномалій у кібербезпеці майбутнього.

### 4. Виклад основного матеріалу

Виявлення аномалій є критично важливим процесом у багатьох галузях, спрямованим на ідентифікацію незвичайних або рідкісних елементів даних, які відрізняються від більшості спостережень. Ці аномалії можуть свідчити про помилки, шахрайство, незвичайні події або нові тенденції, тому їхнє своєчасне виявлення є важливим для прийняття обґрунтованих рішень та запобігання потенційним проблемам.

Одним із фундаментальних підходів до виявлення аномалій є використання статистичних методів, зокрема, аналіз на основі нормального розподілу та Z-score.

Нормальний розподіл є однією з найпоширеніших моделей розподілу даних у природі та різних сферах діяльності [1]. Він часто описує розподіл випадкових величин, на які впливає велика кількість незалежних випадкових факторів. Графічно нормальний розподіл має форму симетричного дзвона, де більшість даних зосереджена навколо середнього значення. Цей розподіл повністю визначається двома параметрами: середнім значенням ( $\mu$ ) та стандартним відхиленням ( $\sigma$ ) [3]. Середнє значення вказує на центр розподілу, а стандартне відхилення характеризує розкид даних навколо цього центру. У програмуванні, зокрема при використанні бібліотеки NumPy у Python, функція *np.random.normal()* використовує саме ці два параметри для генерації випадкових чисел, що підпорядковуються нормальному розподілу [1].

Середнє значення ( $\mu$ ) є мірою центральної тенденції, що представляє собою середнє арифметичне набору даних [4]. Воно розраховується як сума всіх значень, поділена на їхню кількість. Стандартне відхилення ( $\sigma$ ), з іншого боку, є мірою розсіювання, яка показує, наскільки далеко в середньому окремі точки даних відхиляються від середнього значення [5]. Чим більше стандартне відхилення, тим більший розкид даних. Для кількісної оцінки віддаленості окремої точки даних ( $x$ ) від середнього значення у контексті стандартного відхилення використовується Z-score, формула якого виглядає наступним чином [6]:

$$Z = (x - \mu) / \sigma$$

У програмному коді, наприклад, у Python з використанням NumPy, середнє значення та стандартне відхилення набору даних можуть бути легко обчислені за допомогою функцій *np.mean()* та *np.std()* відповідно [1].

Z-score, також відомий як стандартне відхилення, є безрозмірною величиною, яка показує, на скільки стандартних відхилень конкретна точка даних відхиляється від середнього значення [6]. Позитивний Z-score вказує на те, що значення знаходиться вище середнього, а негативний Z-score – нижче. Формула для розрахунку Z-score залишається незмінною [7]:

$$Z = (x - \mu) / \sigma$$

У кодї, після обчислення середнього значення та стандартного відхилення набору даних, Z-score для кожної точки даних розраховується шляхом віднімання середнього значення від точки даних і ділення отриманої різниці на стандартне відхилення.

Для визначення того, чи є певна точка даних аномалією на основі її Z-score, використовується поріг (threshold) [8]. Поріг є заздалегідь визначеним значенням, яке використовується для порівняння з абсолютним значенням Z-score. Якщо абсолютне значення Z-score перевищує встановлений поріг, відповідна точка даних вважається аномальною, оскільки вона знаходиться на значній відстані від середнього значення. У кодї це часто реалізується за допомогою логічного виразу:

$$\text{np.abs}(Z - \text{score}) > \text{threshold},$$

де *np.abs()* обчислює абсолютне значення Z-score, а результат порівняння є булевим масивом, що вказує, які точки даних є аномаліями.

У випадках, коли дані мають кілька змінних або ознак, для позначення аномалій може використовуватися логічний оператор OR. Цей оператор дозволяє об'єднувати кілька умов для виявлення аномалій, які можуть бути помітними лише за однією або кількома змінними [6]. Наприклад, якщо для кожної змінної розраховано Z-score і встановлено поріг, точка даних може бути позначена як аномалія, якщо її Z-score перевищує поріг хоча б для однієї з розглянутих змінних.

Виявлення аномалій за допомогою Z-score є ефективним методом, що базується на розумінні розподілу даних через його середнє значення та стандартне відхилення. Розраховуючи Z-score для кожної точки даних, можна визначити її віддаленість від середнього значення в одиницях стандартного відхилення. Встановлення порогу для абсолютного значення Z-score дозволяє ідентифікувати точки даних, які є статистично значущими відхиленнями від норми і, отже, можуть вважатися аномаліями.

Таблиця 1

## Інтерпретація Z-score

| Діапазон Z-score      | Інтерпретація                                                                                                                           |
|-----------------------|-----------------------------------------------------------------------------------------------------------------------------------------|
| Від -1 до 1           | Більшість даних потрапляє в цей діапазон (приблизно 68% у нормальному розподілі). Вважається типовим.                                   |
| Від -2 до 2           | Значна частина даних потрапляє сюди (приблизно 95%). Зазвичай вважається нормальним.                                                    |
| Від -3 до 3           | Майже всі дані знаходяться в цьому діапазоні (приблизно 99.7%). Значення, що наближаються або перевищують ці межі, стають незвичайними. |
| Більше 3 або менше -3 | Ці значення є досить рідкісними в нормальному розподілі і часто вважаються значними аномаліями.                                         |

Виявлення аномалій можна реалізувати на основі математичного аналізу розподілу даних. Один із простих підходів – визначення статистичних відхилень у вибірці даних. У цьому прикладі використовується генерація синтетичних даних та їх подальший аналіз для виявлення потенційних аномалій.

```
import numpy as np
import matplotlib.pyplot as plt
```

```
np.random.seed(2)
# Генерація синтетичних даних
x = np.random.normal(3, 1, 130)
y = np.random.normal(180, 40, 130) / x
# Візуалізація даних
plt.scatter(x, y)
plt.xlabel('X')
plt.ylabel('Y')
plt.title('Scatter plot of the data')
plt.show()
```

Після генерації даних аналізуються їх статистичні характеристики, включаючи середнє значення та стандартне відхилення. Для виявлення аномалій встановлюється поріг, після перевищення якого точки класифікуються як потенційно небезпечні.

```
# Обчислення середнього та стандартного відхилення
mean_x = np.mean(x)
std_x = np.std(x)
mean_y = np.mean(y)
std_y = np.std(y)
# Встановлення порогу аномалій
threshold = 2.5
# Виявлення аномалій
anomalies = np.where((np.abs((x - mean_x) / std_x) > threshold) | (np.abs((y - mean_y) / std_y) > threshold))
# Візуалізація аномалій
plt.scatter(x, y)
plt.scatter(x[anomalies], y[anomalies], color='red', label='Anomalies')
plt.xlabel('X')
plt.ylabel('Y')
plt.title('Anomaly Detection')
plt.legend()
plt.show()
```

Аномальні точки виділені на графіку, що свідчить про їх значне відхилення від загального розподілу даних. У реальних системах подібний аналіз може використовуватися для виявлення нетипових мережевих активностей, підозрілих транзакцій або поведінкових відхилень користувачів.

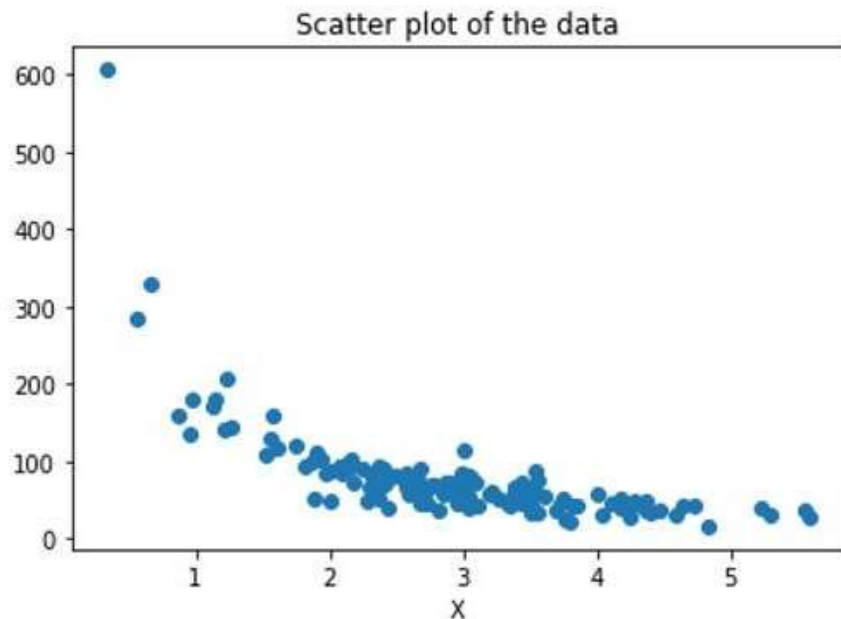


Рис.1. Загальний розподіл даних

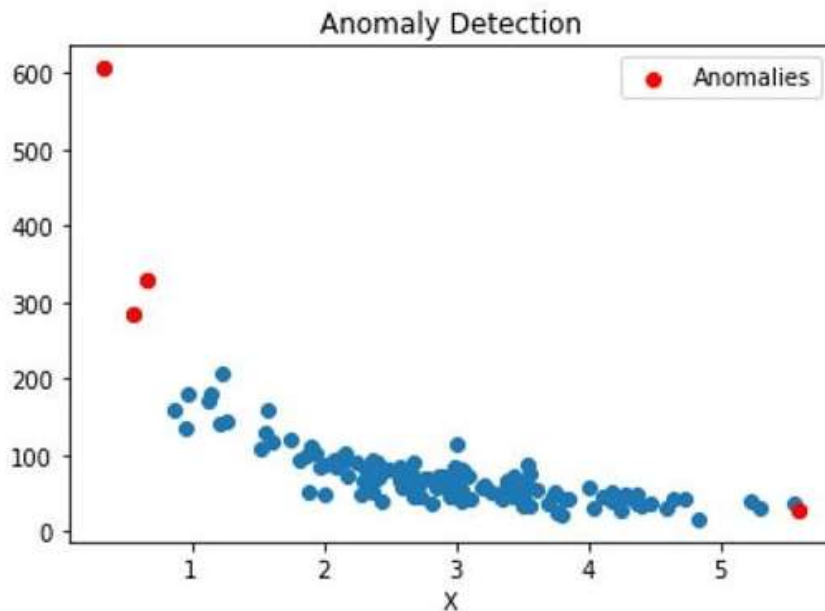


Рис.2. Виявлення аномалій

Більш складні методи виявлення аномалій включають алгоритми машинного навчання, такі як Isolation Forest, Autoencoders або DBSCAN. Вони дозволяють обробляти великі обсяги даних і виявляти складні нелінійні закономірності, що недоступні для простих статистичних методів. Завдяки цьому використання ШІ у сфері кібербезпеки стає ключовим фактором у захисті цифрових систем від нових та еволюційних загроз.

##### 5. Висновки та перспективи подальших досліджень.

Машинне навчання відіграє вирішальну роль у сучасній кібербезпеці. Його здатність аналізувати величезні обсяги даних, розпізнавати патерни та адаптуватися до нових загроз робить

його незамінним інструментом. Зростання обсягів даних та складність інформаційних систем зумовлюють потребу в ефективних інструментах для виявлення аномалій — відхилень, які можуть вказувати на помилки, збої, загрози безпеці або нетипову поведінку систем. У сферах, де точність і швидкість реагування критично важливі (зокрема, в інформаційних технологіях, фінансах, енергетиці, медицині), здатність ідентифікувати аномальні події відіграє ключову роль у забезпеченні надійності та стабільності процесів.

Машинне навчання надає нові можливості для автоматизованого виявлення аномалій, дозволяючи будувати адаптивні моделі, здатні аналізувати великі масиви даних і виявляти нетипові патерни без попередньо заданих правил. Це визначає практичну й наукову цінність дослідження методів машинного навчання у задачах детекції аномалій.

Подальші напрямки досліджень включають розробку більш інтерпретованих моделей машинного навчання для підвищення довіри до результатів виявлення аномалій, адаптацію моделей до змін у мережевому трафіку та атаках у режимі реального часу, а також інтеграцію методів машинного навчання з існуючими системами безпеки для створення більш інтелектуальних та автономних рішень для захисту мереж. У підсумку, машинне навчання є важливим інструментом для виявлення мережевих аномалій, що забезпечує значні переваги порівняно з традиційними підходами та відіграє ключову роль у підвищенні рівня кібербезпеки в сучасному цифровому світі.

#### Список використаних джерел

1. Numeracy, Maths and Statistics - Academic Skills Kit. The things we do here make a difference out there | Newcastle University. URL: <https://www.ncl.ac.uk/webtemplate/ask-assets/external/maths-resources/statistics/distributions/normal-distribution.html> (дата звернення: 27.03.2025).
2. Гайдур, Г. І., Гахов, С. О., & Гамза, Д. Є. (2024). Модель виявлення шкідливої активності в інформаційній системі організації на основі гібридної класифікації. Сучасний захист інформації, (4), 30-38. DOI: 10.31673/2409-7292.2024.040003 (дата звернення: 27.03.2025).
3. Admin. Normal Distribution Formula. BYJUS. URL: <https://byjus.com/normal-distribution-formula/> (дата звернення: 27.03.2025).
4. numpy.random.standard\_normal – NumPy v2.1 Manual. NumPy. URL: [https://numpy.org/doc/2.1/reference/random/generated/numpy.random.standard\\_normal.html](https://numpy.org/doc/2.1/reference/random/generated/numpy.random.standard_normal.html) (дата звернення: 28.03.2025).
5. numpy.random.normal – NumPy v2.3.dev0 Manual. NumPy. URL: <https://numpy.org/devdocs/reference/random/generated/numpy.random.normal.html> (дата звернення: 28.03.2025).
6. Numeracy, Maths and Statistics - Academic Skills Kit. The things we do here make a difference out there | Newcastle University. URL: <https://www.ncl.ac.uk/webtemplate/ask-assets/external/maths-resources/statistics/descriptive-statistics/mean-median-and-mode.html> (дата звернення: 28.03.2025).
7. How to Find the Mean | Definition, Examples & Calculator. Scribbr. URL: <https://www.scribbr.com/statistics/mean/> (дата звернення: 29.03.2025).
8. Khan Academy. Mean, median, and mode. Khan Academy. URL: <https://www.khanacademy.org/math/statistics-probability/summarizing-quantitative-data/mean-median-basics/a/mean-median-and-mode-review> (дата звернення: 29.03.2025).
9. Sobchuk V., Gakhov S., Smoliev Y., Haidur H. Enhancing Intrusion Detection in Organizational Information Systems through AI-Powered Traffic Analysis (2024) CEUR Workshop Proceedings, 3933, pp. 190 - 203. <https://www.scopus.com/inward/record.uri?eid=2-s2.0-86000713893&partnerID=40&md5=f59cab3a67f02a57321b8172f4b75ed6> (дата звернення: 27.04.2025).