

УДК 004.9

DOI <https://doi.org/10.35546/kntu2078-4481.2025.1.2.19>**В. О. МАРТОВИЦЬКИЙ**

кандидат технічних наук, доцент,
доцент кафедри електронних обчислювальних машин
Харківський національний університет радіоелектроніки
ORCID: 0000-0003-2349-0578

А. С. СВИРИДОВ

аспірант кафедри безпеки інформаційних технологій
Харківський національний університет радіоелектроніки
ORCID: 0000-0002-9830-4103

О. С. АВДЄЄВ

магістр кафедри електронних обчислювальних машин
Харківський національний університет радіоелектроніки
ORCID: 0009-0006-7249-4371

І. В. ГУДЗИНСЬКИЙ

магістр кафедри електронних обчислювальних машин
Харківський національний університет радіоелектроніки
ORCID: 0009-0008-8409-7220

О. О. КОРОТЕЦЬКИЙ

магістр кафедри електронних обчислювальних машин
Харківський національний університет радіоелектроніки
ORCID: 0009-0008-2866-2635

ДОСЛІДЖЕННЯ МЕТОДІВ ВИЯВЛЕННЯ АНОМАЛІЙ У API ЖУРНАЛАХ ДЛЯ ЗАБЕЗПЕЧЕННЯ БЕЗПЕКИ ТА НАДІЙНОСТІ ПРОГРАМНИХ СИСТЕМ

Зі збільшенням кількості журналів логів від різних API ручна перевірка та аналіз, стає все більш складнішим завданням. Методи машинного навчання дозволяють автоматизувати процес аналізу великих обсягів даних та виявляти нестандартні шаблони, які можуть свідчити про якісь аномалії чи загрози. Дослідження логів від API систем дозволяє визначити, чи є запити до системи безпечними або підозрілими. Аномалії у запитах можуть свідчити про спроби несанкціонованого доступу або зловмисних дій до комп'ютерної системи. На сьогоднішній день компанії часто залучають сторонніх фахівців із пентестингу, які періодично тестують програмне забезпечення на наявність вразливостей. Проте для підвищення автономності та безпеки доцільно впровадити систему, здатну самостійно ідентифікувати аномальний трафік, що може вказувати на потенційні загрози. Такий підхід дозволить попереджати атаки завчасно та мінімізувати ризики швидше, ніж при очікуванні результатів від зовнішніх експертів. Об'єктом дослідження є процеси виявлення аномалій у журналах API для підвищення безпеки та надійності програмних систем. Метою роботи є дослідження та оцінка ефективності методів виявлення аномалій у API журналах з метою підвищення безпеки та надійності програмних систем. Дослідження спрямоване на порівняння різних моделей неконтрольованого навчання та визначення найбільш ефективною для виявлення потенційно зловмисної активності у логах API. В роботі досліджено методи виявлення аномалій у API журналах, що є критично важливим для забезпечення безпеки та надійності програмних систем. Розглянуто сучасні підходи до аналізу журналів, зокрема методи машинного навчання, статистичного аналізу та виявлення аномалій. Встановлено, що ефективне виявлення аномалій дозволяє своєчасно ідентифікувати потенційні загрози, такі як кібератаки, помилки в роботі системи або несанкціонований доступ, що значно підвищує рівень захищеності програмного забезпечення.

Ключові слова: виявлення аномалій, кібербезпека, штучний інтелект, безпека програмних систем, виявлення загроз.

V. O. MARTOVYTSKYI

Candidate of Technical Sciences, Associate Professor,
Associate Professor at the Department of Electronic Computers
Kharkiv National University of Radio Electronics
ORCID: 0000-0003-2349-0578

A. S. SVYRYDOV

Postgraduate Student at the Department of Information Technology Security
Kharkiv National University of Radio Electronics
ORCID: 0000-0002-9830-4103

O. S. AVDEIEV

Master's Student at the Department of Electronic Computers
Kharkiv National University of Radio Electronics
ORCID: 0009-0006-7249-4371

I. V. HUDZYNSKYI

Master's Student at the Department of Electronic Computers
Kharkiv National University of Radio Electronics
ORCID: 0009-0008-8409-7220

O. O. KOROTETSKYI

Master's Student at the Department of Electronic Computers
Kharkiv National University of Radio Electronics
ORCID: 0009-0008-2866-2635

RESEARCH ON ANOMALY DETECTION METHODS IN API LOGS TO ENSURE SECURITY AND RELIABILITY OF SOFTWARE SYSTEMS

With the increasing number of logs from various APIs, manual inspection and analysis is becoming an increasingly difficult task. Machine learning methods allow automating the process of analyzing large amounts of data and detecting unusual patterns that may indicate some anomalies or threats. Examining logs from API systems allows you to determine whether requests to the system are safe or suspicious. Anomalies in requests may indicate attempts at unauthorized access or malicious actions to a computer system. Today, companies often engage third-party pentesting specialists to periodically test software for vulnerabilities. However, to increase autonomy and security, it is advisable to implement a system that can independently identify abnormal traffic that may indicate potential threats. Such an approach will allow preventing attacks in advance and minimizing risks faster than waiting for results from external experts. The object of research is the processes of detecting anomalies in API logs to improve the security and reliability of software systems. The aim of the study is to investigate and evaluate the effectiveness of methods for detecting anomalies in API logs in order to improve the security and reliability of software systems. The study is aimed at comparing different unsupervised learning models and determining the most effective one for detecting potentially malicious activity in API logs. The paper investigates methods for detecting anomalies in API logs, which is critical to ensuring the security and reliability of software systems. Modern approaches to log analysis, including machine learning, statistical analysis, and anomaly detection methods, are considered. It is established that effective anomaly detection allows timely identification of potential threats, such as cyber attacks, system errors or unauthorized access, which significantly increases the level of software security.

Key words: anomaly detection, cybersecurity, artificial intelligence, software system security, threat detection.

Постановка проблеми

На сьогоднішній день захист програмних інтерфейсів (API) застосунків є серйозною задачею. Онлайн-спільнота OWASP регулярно оновлює список найбільш поширених загроз безпеки API, що доводить важливість їх захисту [1]. Виявлення вразливих місць у API різних сервісів – це досить складний і трудомісткий процес, який може сильно відрізнятися в залежності від масштабів або сфер застосування тієї чи іншої системи/сервісу. Більшість сучасних компаній надають доступ до своїх послуг через веб-застосунки. Тим самим клієнти зберігають та обмінюються конфіденційною інформацією між клієнтською і серверною частиною, що вимагає забезпечення високого рівня захисту таких API інтерфейсів системи. Однак при здійсненні кібернетичного впливу на такі системи компанії можуть зазнати дуже великих збитків, як фінансових, так і технічних.

На сьогодні одним із найкращих способів виявлення вразливостей API сервісів є залучення відповідних фахівців з тестування на проникнення. Вони здатні виявити слабкі місця, оцінити їхню критичність згідно зі стандартами OWASP та надати рекомендації щодо усунення відповідних вразливостей системи [2]. Хоча тестування на проникнення є важливим кроком до покращення системи безпеки, але воно включає ручний аналіз журналів безлічі log-файлів системи. Таке завдання може бути дуже трудомістким та ресурсозатратним, що може призвести до різного роду помилок. Такі помилки можуть привести до того, що не всі вразливості системи будуть знайдені. І тут на допомогу фахівцям з тестування на проникнення приходить штучний інтелект. Машинне навчання відіграє дуже важливу роль при виявленні аномалій не тільки в log-файлах від API, а й в комп'ютерних системах в цілому. Методи машинного навчання дозволяють автоматизувати процес аналізу великих обсягів даних та

виявляти нестандартні шаблони, які можуть свідчити про якісь аномалії чи загрози. Тому ключовим аспектом у забезпеченні захисту API інтерфейсів є розуміння аномальної активності запитів до API сервісів. Виявлення таких аномалій може підвищити рівень захисту системи та запобігти порушенням безпеки.

Таким чином, наше дослідження спрямоване на дослідження та порівняння різних моделей неконтрольованого навчання та визначення найбільш ефективних моделей для виявлення потенційно зловмисної активності у log-файлах від API інтерфейсів.

Аналіз останніх досліджень і публікацій

Сучасні інформаційні технології створюють великі обсяги log-файлів, що генеруються в результаті взаємодії користувачів із різними компонентами системами через різні API інтерфейси. Однією з ключових задач у сфері кібербезпеки є своєчасне виявлення зловмисної активності, які можуть приховано здійснюватися через API-запити, або іншими словами аномалії. log-файли API містять велику кількість інформації, яка може бути використана для аналізу та виявлення аномальної активності. Однак, через відсутність розмічених даних та необхідності обробки великих обсягів інформації, традиційні методи контрольованого навчання часто виявляються не ефективними. У цьому випадку методи неконтрольованого навчання стають перспективним напрямком розвитку, оскільки вони дозволяють виявляти шаблони нормальної та аномалії активності без попередньої розмітки даних.

Для більш детального розуміння сучасного стану останніх досліджень і публікацій в цьому напрямку розглянемо різні моделі неконтрольованого навчання для виявлення зловмисної активності у log-файлів від API.

У статті [3] автори використовуючи адаптивне навчання та еволюційні обчислення, запропонований підхід, що покращує виявлення загроз, контроль доступу та виявлення аномалій у реальному часі. У дослідженні обговорюються ключові проблеми безпеки, стратегії впровадження та потенційні досягнення в галузі захисту API для сектору FinTech.

Стаття [4] підкреслює важливість інтеграції захисту API в стратегію Zero Trust, яка базується на принципі «нікому не довіряй, перевіряй завжди». Автори наголошують, що API, будучи ключовим компонентом сучасних інформаційних технологій, часто стають слабкою ланкою в кібербезпеці через недостатню увагу до їх захисту. У статті розглядаються основні загрози, такі як несанкціонований доступ, ін'єкції та витоки даних, а також пропонуються рекомендації щодо впровадження чітких механізмів аутентифікації, авторизації, моніторингу та шифрування для API. Наголошується, що без надійного захисту API реалізація стратегії Zero Trust буде неповною, що може призвести до серйозних наслідків для безпеки організацій.

У роботі [5] автори розглядають атаки типу ін'єкція в API, які стосуються несанкціонованого або зловмисного використання API і часто використовуються для отримання незаконного доступу до конфіденційних даних або маніпулювання онлайн-системами з незаконними цілями. Автори запропонували нову неконтрольовану модель виявлення аномалій з малим набором даних, яка складається з двох основних частин:

- по-перше, навчання спеціалізованої мовної моделі для API на основі FastText-векторизації;
- по-друге, використання методу наближений пошук найближчих сусідів (Approximate Nearest Neighbor) у підході до класифікації даних.

Запропонований підхід дозволяє навчати швидко та легку модель класифікації, використовуючи лише кілька прикладів нормальних запитів до API. Автори оцінили продуктивність своєї моделі на наборах даних CSIC 2010 та ATRDF 2023. Результати демонструють, що запропонована авторами модель покращує точність виявлення атак на API порівняно з сучасними неконтрольованими методами виявлення аномалій.

У статті [6] розглядається ключова роль аутентифікації у забезпеченні безпеки веб API. Обговорюються основні принципи безпеки, популярні методи аутентифікації (такі як OAuth 2.0, OpenID Connect та робочі процеси на основі JWT), а також поширені помилки, які призводять до порушень. Автори також досліджують найкращі практики та новітні тенденції, такі як архітектура нульової довіри (zero trust) та системи децентралізованої ідентифікації.

Таким чином, актуальним стає дослідження спрямоване на вивчення методів та підходів захисту API сервісів.

Формулювання мети дослідження

Метою роботи є розробка та оцінка ефективності методів виявлення аномалій у API журналах з метою підвищення безпеки та надійності програмних систем.

Викладення основного матеріалу дослідження

В цій роботі пропонується підхід який бере за основу результати роботи [7]. Наше дослідження включає в себе чотири основні етапи, такі як: підготовка даних, виявлення аномалій, отримання результатів та оцінка результатів дослідження. На рисунку 1 представлено схематично взаємодію цих етапів. Суцільні лінії відображають кроки, які виконуються під час оцінки, тоді як пунктирні лінії позначають кроки, які відсутні в роботі [7] і по суті є нашою модифікацією.

З самого початку необхідно підготувати відповідні дані для навчання та тестування моделей. Зазвичай такими вхідними даними є набір даних, який дозволяє відрізнити безпечну активність, тобто нормальну, від зловмисної, або ж необроблені журнали логів. Ці вхідні дані часто потребують попередньої обробки з використанням методів,

які залежать від особливостей набору даних. Наприклад, IP-адреси мають ієрархічну структуру, номери портів подаються у числовому форматі, а протоколи класифікуються на різні типи.

Методи обробки залежать від обраних технік виявлення аномалій, які поділяються на контрольовані та неконтрольовані. В рамках нашого дослідження розглядалися виключно методи неконтрольованого виявлення аномалій, як було зазначено раніше. Після попередньої обробки до даних застосовуються вибрані методи виявлення аномалій, що генерують результати, які можна оцінити за допомогою кількісної оцінки (scores) або міток (labels).

У нашій роботі для оцінювання використовуються лише кількісні показники, оскільки мітки (labels) зазвичай застосовуються у задачах класифікації, що не відповідає меті цього дослідження. Отримані показники використовуються для оцінювання всього процесу. Якщо результати оцінювання виявляються незадовільними або буде можливість їх покращити, застосовується ітеративний підхід, що дозволяє налаштовувати моделі для досягнення оптимальної продуктивності на вибраному наборі даних.

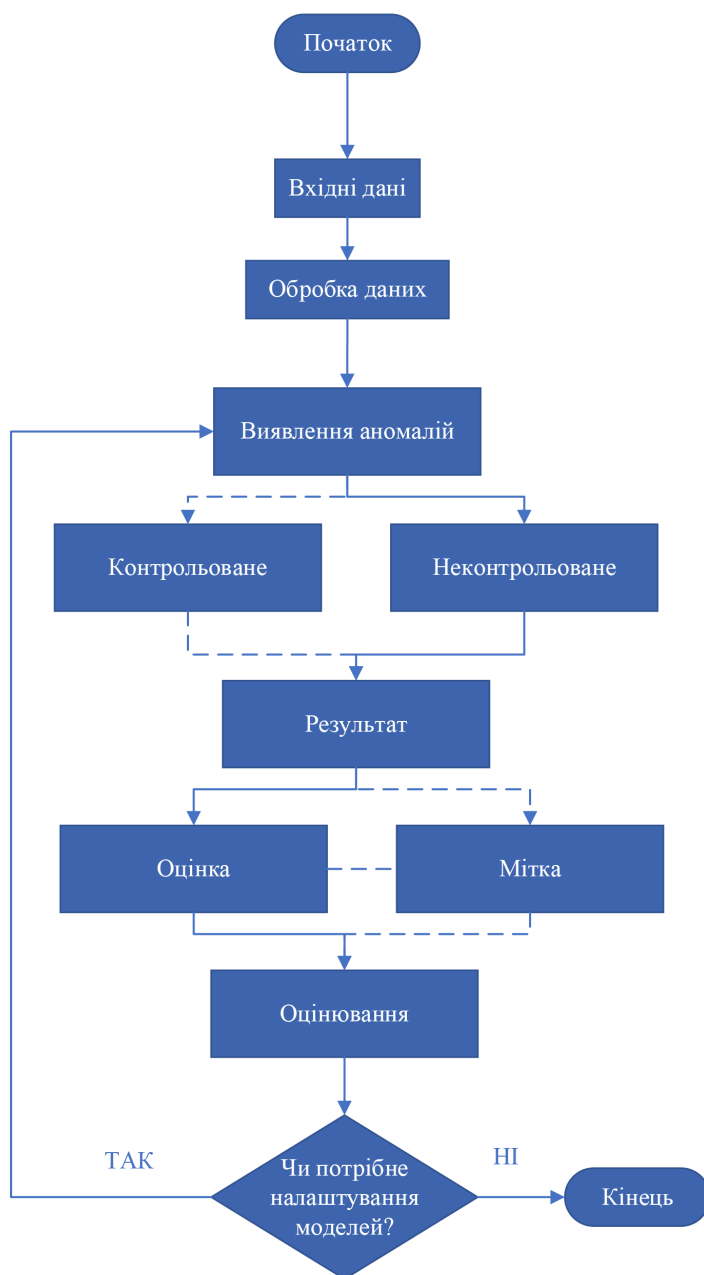


Рис. 1. Схема запропонованого підходу до виявлення аномалій

Етап підготовки може значно відрізнятися залежно від необхідної попередньої обробки даних. Підготовча фаза включає такі завдання, як збір даних, виділення ознак і іноді нормалізація даних або стандартизація. На рисунку 2 проілюстровано відповідні завдання підготовчого етапу, які складаються зі збору журналів, вилучення

ознак, нормалізації даних і машинного навчання. Ці кроки було вжито під час неконтрольованого виявлення аномалій. Однак нормалізація і стандартизація не завжди потрібні. Але нормалізація була виконана в цій роботі, щоб потенційно підвищити точність передбачення та візуалізувати деякі моделі для подальшого дослідження.



Рис. 2. Етапи підготовки

Фаза виявлення аномалій починається після завершення етапу підготовки, коли дані готові для подачі в моделі виявлення аномалій. Це найпростіший етап, оскільки єдина ручна дія – це передача матриці ознак у вибрані моделі для навчання та тестування. Основні завдання цього етапу включають вибір моделі та застосування обраних моделей до попередньо оброблених даних для отримання результатів.

З готовими до використання попередньо обробленими даними було обрано моделі для процесу неконтрольованого виявлення аномалій. Вибір моделей залежав від кількох факторів, зокрема характеристик використаних наборів даних, результатів, описаних у розділі аналіз літератури, а також інших досліджень, що можуть бути цікавими для оцінювання, хоча й не розглядаються безпосередньо в літературному огляді. Обрані моделі включають K-means, GMM (Gaussian Mixture Model), IF (Isolation Forest) та OCSVM (One-Class Support Vector Machine).

Моделі K-means і GMM є методами кластеризації. Використання раніше застосованих моделей дозволяє оцінити, наскільки добре вони працюють із журналами API. Ці моделі також використовувалися в дослідженні [7] для оцінки неконтрольованого виявлення аномалій, оскільки вони є одними з найбільш відомих алгоритмів машинного навчання, що дозволяють групувати дані у кластери.

Також для зменшення часу обчислень моделі OCSVM через високу розмірність набору даних, який використовувався для експериментів було використано метод аналіз головних компонентів (PCA) як метод для зменшення розмірності.

На основі проаналізованих досліджень для оцінювання ефективності методів виявлення аномалій також були обрані моделі IF (Isolation Forest) і OCSVM. Обидві моделі є методами неконтрольованого виявлення викидів: IF – це ансамблева модель, а OCSVM належить до сімейства алгоритмів SVM.

Модель IF була обрана для оцінки неконтрольованого виявлення аномалій, оскільки вона є методом пошуку викидів у неконтрольованому режимі і вже демонструвала перспективні результати у попередніх дослідженнях. OCSVM була обрана завдяки своїм обнадійливим результатам у попередніх дослідженнях.

У нашому дослідженні було використано набір даних HTTP Dataset CSIC. Цей набір містить автоматично згенеровані веб-запити до вебсайту електронної комерції та був розроблений в Інституті інформаційної безпеки CSIC (Іспанська національна дослідницька рада).

Ключові характеристики набору даних наведені в таблиці 1.

У веб застосунку, який використовувався для створення набору даних, користувачі можуть, наприклад, купувати товари, додавати їх до кошика та вводити особисту інформацію.

Набір даних містить 61 000 записів, серед яких 36 000 є нормальними запитами, а 25 000 – аномальними. Кожен запит має мітку, що визначає його як нормальний або аномальний.

Дані були згенеровані за допомогою процесу, який спочатку включав збір реальних даних для всіх параметрів вебзастосунку. Автентичні дані, включаючи імена, прізвища, адреси тощо, були отримані з реальних баз даних. Ці значення потім зберігалися у двох окремих базах даних: одна містила нормальні значення, а інша – аномальні

Таблиця 1

Характеристики набору даних

Характеристики	Значення
Розмір набору даних	61 000
Кількість записів про нормальну поведінку	36 000
Кількість записів про аномальну поведінку	25 000
Розподіл класів	Нормальні: 59 %, Аномалії: 41 %
Типи даних	Змішані (категоріальні, числові, текстові)

значення, призначені для імітації шаблонів атак. Крім того, був складений список усіх загальнодоступних сторінок веб застосунку. Потім для кожної веб сторінки генерувалися нормальні та аномальні HTTP-запити. Нормальні запити мали параметри, заповнені даними, випадково вибраними з бази нормальних значень, а аномальні запити формувалися зі значеннями, вибраними з бази аномальних значень, щоб імітувати різні сценарії атак.

Для запитів, класифікованих як аномальні, було створено три категорії. По-перше, статичні атаки, що включають запити до прихованих або неіснуючих ресурсів, таких як застарілі файли, ідентифікатори сесій у змінених URL-адресах, конфігураційні файли, файли за замовчуванням тощо. По-друге, динамічні атаки передбачають зміну дійсних аргументів запиту. Вони включають, але не обмежуються, SQL-ін'єкціями, Carriage Return Line Feed (CRLF) ін'єкціями, міжсайтовим скриптингом (XSS) та переповненням буфера. Нарешті, також були змодельовані ненавмисні нелегальні запити. Хоча ці запити не є зловмисними за своєю природою, вони відхиляються від нормальної роботи веб застосунку та мають структуру, відмінну від нормальних параметрів, наприклад, номер телефону, що складається з літер.

Базові моделі налаштовані за замовчуванням. Показники продуктивності цих моделей наведені в таблиці 2.

Таблиця 2

Показники продуктивності для базових моделей

Модель	Точність	Повнота	F1-міра	AUC	Точність	Час (с)
GMM	0.61	0.74	0.65	0.78	0.70	0.35
K-means	0.53	0.61	0.55	0.56	0.61	0.12
Isolation Forest	0.79	0.27	0.40	0.84	0.67	0.18
OCSVM	0.57	0.66	0.60	0.63	0.63	32.35

Розглядаючи базову модель GMM, можна помітити, що вона досягла збалансованого результату з коротким часом навчання. Точність 0.61 вказує на те, що модель здатна досить добре знаходити справжні аномалії. Крім того, повнота 0.74 свідчить про те, що GMM успішно виявляє реальні аномалії в даних. F1-міра 0.65 відображає баланс між точністю та повнотою, що вказує на помірну ефективність моделі у визначенні аномалій.

Модель GMM також має хороший AUC-індекс 0.78, що означає її здатність добре розрізняти нормальні запити та аномалії. Точність 0.7 показує, що модель здатна правильно ідентифікувати різні аномалії.

Модель K-means демонструє змішані результати, але з коротшим часом навчання порівняно з GMM. Її точність, яка становить 0.53, свідчить про те, що лише приблизно половина передбачених аномалій була правильною. Повнота моделі становить 0.61, що означає, що вона виявила значну частку справжніх аномалій. Однак 38 % аномалій були визначені неправильно, що підкреслює необхідність подальшого вдосконалення. Значення F1-міри 0.55 додатково вказує на необхідність покращення балансу між точністю та повнотою. AUC, що дорівнює 0.56, є близьким до випадкового вгадування, що свідчить про низьку здатність моделі розрізняти аномалії. Незважаючи на те, що модель досягла точності 0.61, базова модель K-means може бути недостатньо надійною для критичних додатків, що вимагають точного виявлення аномалій. Це додатково вказує на необхідність покращення моделі або дослідження альтернативних алгоритмів для підвищення точності.

Модель Isolation Forest швидко обробляє вибраний набір даних, виконуючи обчислення всього за 0.18 секунди, що свідчить про її високу ефективність і потенційну придатність для застосувань, які вимагають швидкої обробки даних. Хоча базова модель показує високий AUC, що вказує на хорошу здатність до розпізнавання, її F1-міра 0.40 свідчить про можливий дисбаланс між точністю та повнотою. Це підтверджується високою точністю 0.78 і низькою повнотою 0.27. Висока точність означає, що модель має багато істинно позитивних і мало хибнопозитивних спрацьовувань. У той же час низька повнота вказує на те, що модель не виявляє всі справжні аномалії. Ці спостереження свідчать про те, що точність 0.67 не дає повної картини через дисбаланс у наборі даних.

Модель OCSVM показала перспективні результати, але її навчання зайняло набагато більше часу порівняно з іншими моделями. Точність 0.55, представлена в таблиці 2, означає, що трохи більше половини виявлених аномалій були правильними. Повнота 0.66 свідчить про те, що модель змогла ідентифікувати більшість справжніх аномалій. Збалансована F1-міра 0.60 вказує на компроміс між точністю та повнотою. Значення AUC 0.63 свідчить про помірну здатність до розрізнення аномалій, а точність 0.63 означає, що модель правильно класифікувала 63 % усіх запитів. Однак є можливість покращення моделі, щоб зменшити кількість хибнонегативних випадків і підвищити кількість істиннопозитивних.

Загалом таблиця відображає початкову продуктивність K-means і GMM у кластеризації та виявленні аномалій, вказуючи на необхідність подальшої оптимізації та вдосконалення для підвищення точності кластеризації та здатності до виявлення аномалій.

Висновки

В роботі досліджено методи виявлення аномалій у API журналах, що є критично важливим для забезпечення безпеки та надійності програмних систем. Розглянуто сучасні підходи до аналізу журналів, зокрема методи

машинного навчання, статистичного аналізу та виявлення аномалій. Встановлено, що ефективне виявлення аномалій дозволяє своєчасно ідентифікувати потенційні загрози, такі як кібератаки, помилки в роботі системи або несанкціонований доступ, що значно підвищує рівень захищеності програмного забезпечення.

Крім того, представлені результати дослідження підтверджують важливість інтеграції методів виявлення аномалій у процеси моніторингу та логування, що дозволяє автоматизувати процес аналізу та оперативного реагування на потенційні загрози. Результати дослідження підтверджують, що використання сучасних алгоритмів, таких як кластеризація є ефективним інструментом для підвищення надійності та безпеки програмних систем.

Подальші дослідження можуть бути спрямовані на вдосконалення методів обробки великих обсягів даних, оптимізацію алгоритмів для роботи в реальному часі та інтеграцію з іншими системами кібербезпеки для створення комплексного захисту програмних систем.

Список використаної літератури

1. OWASP, OWASP API Security Project. Available at: <https://owasp.org/www-project-api-security/> (Accessed: 22 February 2025).
2. Lala, S.K., Kumar, A., & Subbulakshmi, T. (2021, May). Secure web development using owasp guidelines. In 2021 5th International Conference on Intelligent Computing and Control Systems (ICICCS) (pp. 323-332). IEEE.
3. Catherine, A., Anastasia, D., Olga, S., & Adebayo, F.R. (2025). Enhancing API Security in FinTech with Genetic Algorithm-Based Machine Learning Models.
4. Costa, Fabio, and Akamai Solutions Engineer Principal. "Sécurité des APIs: un pilier oublié de la stratégie Zéro Trust—Global Security Mag Online". (2025).
5. Aharon, U., Dubin, R., Dvir, A., & Hajaj, C. (2025). A classification-by-retrieval framework for few-shot anomaly detection to detect API injection. *Computers & Security*, 150, 104249.
6. Mohiuddin Ahmed, Abdun Naser, and Jiankun Hu. "A survey of network anomaly detection techniques". In: *Journal of Network and Computer Applications* 60 (2016), pp. 19–31. ISSN: 1084-8045. DOI: 10.1016/j.jnca.2015.11.016. URL: <https://www.sciencedirect.com/science/article/pii/S1084804515002891>.
7. Carmen Sánchez-Zas, Xavier Larriva-Novo, Víctor A Villagrà, Mario Sanz Rodrigo, and José Ignacio Moreno. "Design and Evaluation of Unsupervised Machine Learning Models for Anomaly Detection in Streaming Cybersecurity Logs". In: *Mathematics* 10.21 (2022), p. 4043.