

Міністерство освіти і науки України
Вінницький національний технічний університет
Факультет інформаційних технологій та комп'ютерної інженерії
Кафедра захисту інформації

МАГІСТЕРСЬКА КВАЛІФІКАЦІЙНА РОБОТА

на тему:

«Засіб моніторингу подій в інформаційній системі з використанням засобів штучного інтелекту»

Виконав: студент 4 курсу групи ІБС-23м
спеціальності 125 «Кібербезпека та захист
інформації»



Олександр ЯКИМОВ

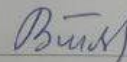
Керівник: к. т. н., доц., доцент кафедри ЗІ



Олеся ВОЙТОВИЧ

«13» 12 2024 р.

Опонент: к. т. н., доц., доцент кафедри ІІЗ



Вікторія ВОЙТКО

«13» 12 2024 р.

Допущено до захисту

Завідувач кафедри ЗІ

д. т. н., проф.

 Володимир ЛУЖЕЦЬКИЙ

«13» 12 2024 р.

Вінниця ВНТУ – 2024 року

Вінницький національний технічний університет
Факультет інформаційних технологій та комп'ютерної інженерії
Кафедра захисту інформації
Рівень вищої освіти II (магістерський)
Галузь знань – 12 Інформаційні технології
Спеціальність – 125 Кібербезпека та захист інформації
Освітньо-професійна програма – Безпека інформаційних і комунікаційних систем

ЗАТВЕРДЖУЮ

Зав. кафедри ЗІ,

д. т. н., проф.

Володимир ЛУЖЕЦЬКИЙ

«18» 09 2024 року

З А В Д А Н Н Я

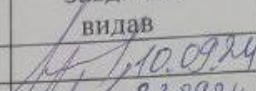
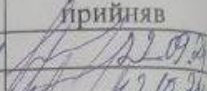
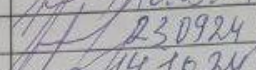
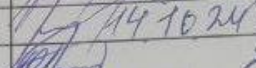
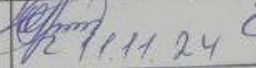
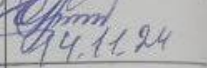
НА МАГІСТЕРСЬКУ КВАЛІФІКАЦІЙНУ РОБОТУ СТУДЕНТУ

Якімову Олександрю Павловичу

1. Тема роботи: «Засіб моніторингу подій в інформаційній системі з використанням засобів штучного інтелекту»
керівник роботи: Войтович Олеся Петрівна, к. т. н., доцент кафедри ЗІ,
затверджені наказом ректора ВНТУ від 17 вересня 2024 року №310
2. Строк подання студентом роботи 17 грудня 2024 року
3. Вихідні дані до роботи:
 - аналіз сучасних SIEM-систем та їхніх можливостей у контексті кібербезпеки;
 - робота з великими даними, які збережені у вигляді логів;
 - використання алгоритмів штучного інтелекту для аналізу мережевого трафіку та системних логів;
 - інтегрування штучного інтелекту у SIEM-систему та проведення експериментальні дослідження його ефективності;
 - оцінка економічної доцільності впровадження ШІ у SIEM-системи.
4. Зміст текстової частини: Вступ. 1. Аналіз предметної області та джерел за темою дослідження. 2. Обґрунтування обраних рішень. 3. Експериментальні дослідження. 4. Економічна частина. Висновки. Список використаних джерел. Додатки.
5. Перелік ілюстративного матеріалу: Вигляд типової SIEM-системи (плакат А4). Застосування методів машинного навчання в кібербезпеці (плакат А4). Графічна схема етапів обробки логів з штучним інтелектом (плакат А4). Топологія мережі (плакат А4). Схема алогриту впровадження штучного інтелекту в систему SIEM (плакат А4). Графічне представлення шляху обробки логів (плакат А4).

Приклади вигляду логів (плакат А4). Приклади вигляду запитів (плакат А4). Візуалізація даних отриманих після обробки штучним інтелектом (плакат А4). Рекомендації Рекомендації штучного інтелекту після обробки логів (плакат А4). Рекомендації штучного інтелекту до та після внесених змін у вигляді конкретизації запиту (плакат А4).

6. Консультанти розділів роботи

Розділ	Прізвище, ініціали та посада консультанта	Підпис, дата	
		Завдання видав	Завдання прийняв
1	Олеся ВОЙТОВИЧ к.т.н., доц. каф.ЗІ	 10.09.24	 22.09.24
2	Олеся ВОЙТОВИЧ к.т.н., доц. каф.ЗІ	 23.09.24	 22.10.24
3	Олеся ВОЙТОВИЧ к.т.н., доц. каф.ЗІ	 14.10.24	 07.11.24
4	Ольга РАТУШНЯК, к.т.н., доц. каф.ЕПВМ	 11.11.24	 14.11.24

7. Дата видачі завдання 1 вересня 2024 року.

КАЛЕНДАРНИЙ ПЛАН

№ з/п	Назва етапів магістерської кваліфікаційної роботи	Строк виконання етапів роботи	Приміт
1	Аналіз завдання. Вступ	01.09.2024 – 10.09.2024	
2	Аналіз літературних джерел за напрямком магістерської кваліфікаційної роботи	10.09.2024 – 15.09.2024	
3	Науково-технічне обґрунтування	16.09.2024 – 22.09.2024	
4	Аналіз предметної області та формування вимог до програмного засобу	23.09.2024 – 29.09.2024	
5	Розробка архітектури програмного засобу	30.09.2024 – 12.10.2024	
6	Програмна реалізація засобу	14.10.2024 – 07.11.2024	
7	Тестування розробленого засобу	07.11.2024 – 10.11.2024	
8	Розробка розділу економічного обґрунтування доцільності розробки	11.11.2024 – 14.11.2024	
9	Аналіз виконання ТЗ, висновки	15.11.2024 – 24.11.2024	
10	Оформлення пояснювальної записки	25.11.2024 – 30.11.2024	
11	Попередній захист та доопрацювання МКР	28.11.2024 – 01.12.2024	
12	Перевірка магістерської роботи на наявність плагіату	02.12.2024 – 14.12.2024	
13	Представлення МКР до захисту	13.12.2024 – 16.12.2024	
14	Захист МКР	17.12.2024 – 20.12.2024	

Студент  Олександр ЯКИМО
Керівник роботи  Олеся ВОЙТОВИЧ

АНОТАЦІЯ

Магістерська кваліфікаційна робота складається з 91 сторінки формату А4, на яких є 31 рисуноків, 12 таблиць, список використаних джерел містить 31 найменування.

Магістерська кваліфікаційна робота присвячена впровадженню та налаштуванню штучного інтелекту в систему управління інформаційною безпекою та подіями (SIEM). У роботі проведено огляд сучасних методів обробки даних у системах SIEM, включно з традиційними алгоритмами обробки логів і сучасними підходами на основі штучного інтелекту та машинного навчання. Виявлено основні обмеження традиційних підходів, зокрема недостатню ефективність роботи з великими та неструктурованими обсягами даних. Запропоновано архітектуру, яка передбачає використання моделей глибокого навчання, таких як GPT, для автоматизації аналізу логів різних форматів. Випробувано різні сценарії роботи, включно з аналізом мережевих, системних і поштових логів, та дійсно налаштування алгоритмів для адаптації до специфіки реальних інцидентів кібербезпеки. У результаті реалізовано модуль інтеграції GPT у середовище SIEM. Також було розраховано економічну частину для проведення дослідження.

Ключові слова: штучний інтелект, обробка логів, автоматизація аналізу, GPT, системи SIEM, кіберзагрози, інтеграція ШІ, неструктуровані журнали подій, аналіз мережевого трафіку, кіберполігон, безпека інформаційних систем..

ABSTRACT

The master's qualification work consists of 91 pages of A4 format, which include 31 figures, 12 tables, and the list of used sources contains 31 names.

The master's qualification work is devoted to the implementation and configuration of artificial intelligence in the information security and event management system (SIEM) to increase its efficiency. The work reviews modern data processing methods in SIEM systems, including traditional log processing algorithms and modern approaches based on artificial intelligence and machine learning. The main limitations of traditional approaches are identified, in particular, insufficient efficiency of working with large and unstructured volumes of data. An architecture is proposed that involves the use of deep learning models, such as GPT, to automate the analysis of logs of various formats. Various scenarios were tested, including analysis of network, system and email logs, and algorithms were adjusted to adapt to the specifics of real cybersecurity incidents. As a result, a module for integrating GPT into the SIEM environment was implemented. The economic part for conducting the research was also calculated.

Keywords: artificial intelligence, log processing, analysis automation, GPT, SIEM systems, cyber threats, AI integration, unstructured event logs, network traffic analysis, cyber polygon, information systems security.

ЗМІСТ

ВСТУП	3
1 АНАЛІЗ ПРЕДМЕТНОЇ ОБЛАСТІ ТА ДЖЕРЕЛ ЗА ТЕМОЮ ДОСЛІДЖЕННЯ...	5
1.1 Огляд SIEM-систем та їх роль у кібербезпеці.....	5
1.2 Проблеми існуючих SIEM-систем при роботі з великими даними	11
1.3 Штучний інтелект і машинне навчання в кібербезпеці	16
1.4 Постановка завдання	22
1.5 Висновок до розділу	23
2 ОБГРУНТУВАННЯ ОБРАНИХ РІШЕНЬ	24
2.1 Джерела логів, особливості їх структури та алгоритми обробки	24
2.2 Моделювання мережевого трафіку для виявлення аномалій за допомогою методів штучного інтелекту	29
2.3 Модель обробки логів на основі ШІ	32
2.4 Інтеграція ШІ у традиційну архітектуру SIEM.....	38
2.5 Вибір GPT для вирішення задач обробки логів	43
2.6 Висновок до розділу	47
3 ЕКСПЕРЕМЕНТАЛЬНІ ДОСЛІДЖЕННЯ	49
3.1 Опис програмного середовища та інструментів	49
3.2 Розробка прототипу системи	53
3.3 Аналіз результатів та оцінка ефективності	59
3.4 Виявлені проблеми та пропозиції щодо вдосконалення.....	69
3.5 Висновок до розділу	72
4 ЕКОНОМІЧНА ЧАСТИНА.....	73
4.1 Оцінювання наукового ефекту	73
4.2 Розрахунок витрат на здійснення науково-дослідної роботи.....	75
4.3 Оцінювання важливості та наукової значимості науково-дослідної роботи	83
4.4 Висновки до розділу	84
ВИСНОВКИ	85
ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ.....	86
ДОДАТКИ	90
Додаток А. ПРОТОКОЛ ПЕРЕВІРКИ НАВЧАЛЬНОЇ (КВАЛІФІКАЦІЙНОЇ) РОБОТИ.....	91

ВСТУП

В умовах сучасного розвитку інформаційних технологій і глобалізації комп'ютерних мереж безпека інформаційних систем стає надзвичайно важливим аспектом діяльності будь-якої організації. Кількість кібератак щороку зростає, а методи атак стають усе складнішими та витонченішими. Одним із ключових інструментів для протидії кіберзагрозам є SIEM-системи (Security Information and Event Management), які забезпечують централізований моніторинг, збір та аналіз системних логів і мережевого трафіку з метою виявлення інцидентів безпеки [1].

Однак, зростання обсягів даних, складність сучасних кіберзагроз та вимоги до швидкості обробки інформації створюють суттєві виклики для традиційних SIEM-систем. Часто такі системи не можуть оперативно виявляти складні аномалії або атаки через обмеження в аналізі великих обсягів даних та залежність від заздалегідь визначених правил і сигнатур, що призводить до великої кількості хибних тривог, зниження ефективності реагування та збільшення навантаження на фахівців з кібербезпеки.

З огляду на ці виклики, з'являється необхідність інтеграції сучасних технологій, таких як штучний інтелект (ШІ) та машинне навчання (ML), для підвищення ефективності SIEM-систем. Використання ШІ дозволяє значно покращити можливості аналізу даних, зменшити кількість хибних тривог, автоматизувати процеси виявлення загроз та зробити систему більш адаптивною до нових і невідомих типів атак.

Об'єктом дослідження є система управління інформаційною безпекою (SIEM), яка виконує моніторинг і аналіз мережевого трафіку та системних логів з метою виявлення кіберзагроз.

Предметом дослідження є методи інтеграції штучного інтелекту та машинного навчання в SIEM-системи для підвищення ефективності виявлення та аналізу загроз у великих обсягах даних.

Метою роботи є покращення безпеки, за можливості аналізу більшої частини логів шляхом інтеграції штучного інтелекту у SIEM-систему.

Використовуючи цей підхід можна досягти підвищення здатності аналізувати великий обсяг мережевого трафіку та системних логів, зменшення кількості хибних тривог і покращення точності виявлення кіберзагроз.

Для досягнення мети необхідно вирішити такі задачі:

- провести огляд сучасних SIEM-систем та їхніх можливостей у контексті кібербезпеки;
- проаналізувати проблеми традиційних SIEM-систем при роботі з великими даними;
- розробка методів обробки даних SIEM-систем за допомогою штучного інтелекту для аналізу мережевого трафіку та системних логів;
- інтегрувати штучний інтелект у SIEM-систему та провести експериментальні дослідження його ефективності;
- оцінити економічну доцільність впровадження ШІ у SIEM-системи з точки зору ресурсних затрат та ефективності роботи.

Удосконалено підхід до аналізу даних у системах SIEM, за рахунок використання адаптивного штучного інтелекту, що дозволяє враховувати особливості обробки неструктурованих даних та адаптації до різних форматів логів та корелювати події з різних джерел, що приводить до мінімізації кількості хибнопозитивних спрацювань

Проміжні результати дослідження були представлені на конференції «Молодь в науці: дослідження, проблеми, перспективи (МН-2025)» [2].

1 АНАЛІЗ ПРЕДМЕТНОЇ ОБЛАСТІ ТА ДЖЕРЕЛ ЗА ТЕМОЮ ДОСЛІДЖЕННЯ

1.1 Огляд SIEM-систем та їх роль у кібербезпеці

Сучасні інформаційні технології та мережева інфраструктура відіграють ключову роль у функціонуванні більшості організацій. Від ефективності та надійності цих технологій залежить не лише продуктивність бізнесу, але й його здатність забезпечувати безперервність операцій, захищати конфіденційні дані та відповідати вимогам нормативно-правових актів. Технологічні платформи, бази даних, хмарні сервіси, системи електронної комерції та інші елементи інформаційної інфраструктури створюють складну екосистему, яка дозволяє компаніям взаємодіяти з клієнтами, партнерами та співробітниками в режимі реального часу [3].

З цього випливає, що такі інфраструктури також потребують моніторингу та перевірки їх компонентів на безпеку, де на допомогу прийде один із важливих інструментів забезпечення безпеки – система SIEM. Функціонал таких систем дозволяє збирати, аналізувати та співвідносити дані про події у інформаційній інфраструктурі, дозволяючи виявляти потенційні загрози та реагувати на них у реальному часі [4].

Система SIEM представляє собою інтегроване рішення, що поєднує два ключових підходи: керування інформацією про безпеку (Security Information Management, SIM) та керування подіями безпеки (Security Event Management, SEM) [5]. SIM відповідає за централізоване збирання, зберігання та аналіз інформації про безпеку, зокрема журналів подій, записів аудиту та даних із різних джерел, що дозволяє виявляти довгострокові тенденції та відхилення в системі. SEM, у свою чергу, зосереджується на моніторингу подій у реальному часі, автоматичному виявленні інцидентів та їхньому оперативному аналізі для забезпечення швидкої реакції на загрози.

Об'єднання цих двох підходів у межах єдиного рішення дає змогу значно розширити функціональні можливості SIEM. Завдяки інтеграції даних із різних

джерел, система забезпечує комплексний підхід до управління безпекою інформаційної інфраструктури, що дозволяє одночасно відстежувати поточний стан подій, реагувати на інциденти та здійснювати ретроспективний аналіз для виявлення прихованих загроз або покращення політик безпеки. Всі описані особливості роблять SIEM ефективним інструментом для вирішення широкого спектра завдань, зокрема ідентифікації загроз, реагування на інциденти, забезпечення відповідності нормативним вимогам та стратегічного планування захисту.

Однією з основних функцій системи є збір і централізоване зберігання даних. Процес збору інформації реалізується через інтеграцію з різноманітними джерелами, до яких належать системні журнали, мережеві пристрої, служби безпеки, програми та бази даних. Зібрані дані перетворюються в уніфікований формат, що забезпечує їх сумісність і готовність до подальшого використання.

Усі зібрані дані акумулюються в єдиному централізованому сховищі, що виступає основою для проведення різних аналітичних операцій, таких як виявлення аномалій, кореляція подій, створення звітів та оцінка довгострокових тенденцій у сфері інформаційної безпеки. Централізований підхід до зберігання сприяє підвищенню ефективності аналізу завдяки доступності всієї інформації в одному місці, що спрощує її обробку та прискорює прийняття рішень у разі інцидентів.

Щодо захисту інформаційних систем, то SIEM забезпечує високий рівень захисту завдяки можливості моніторингу у режимі реального часу, адже атаки відбуваються щосекунди, потрібно правильно та вчасно реагувати на подібні загрози, що може бути критично вирішальним для підприємства. Система SIEM дозволяє адміністраторам безпеки постійно відстежувати події, що відбуваються у мережі та інших компонентах інфраструктури, виявляючи будь-які підозрілі дії чи аномалії. Миттєвий доступ до інформації про події безпеки дає змогу швидко оцінювати загрози та визначати їх критичність.

Особливо важливим аспектом функціонування SIEM є аналіз подій та їх кореляція. Завдяки здатності системи співвідносити події, що надходять з різних

джерел, стає можливим виявлення прихованих загроз. Наприклад, на підприємстві спостерігається незвична активність в мережевому трафіку, що поєднується з підозрілими записами у журналах серверів. Система SIEM реєструє ці події одночасно і проводить кореляцію між ними, аналізуючи дані з різних джерел у реальному часі. Завдяки цій кореляції система може виявити взаємозв'язки між окремими подіями, що допомагає виявити складні атаки, які інакше могли б залишитися непоміченими. Потім команда, відповідальна за інформаційну безпеку, може використати отримані дані для подальшої ідентифікації загроз і запобігання інцидентам у майбутньому, забезпечуючи таким чином більш надійний захист інформаційних ресурсів підприємства [5].

Окрім постійного моніторингу, система SIEM автоматично генерує сповіщення про можливі інциденти безпеки, що базуються на безперервному аналізі даних, які надходять із різних джерел, і спрямовані на оперативне інформування ІТ-фахівців про підозрілі дії або небезпечні ситуації. Завдяки цим сповіщенням спеціалісти з інформаційної безпеки можуть швидко вжити відповідних заходів, мінімізуючи ризики для організації.

Додатково, система SIEM забезпечує автоматичне формування звітів про стан безпеки інформаційної системи. Звіти надають інформацію про виявлені загрози, заходи, які були вжиті для їх нейтралізації, а також загальний рівень відповідності системи нормативним стандартам безпеки. Звіти є ключовими для оцінки поточного рівня безпеки та можуть використовуватися не лише фахівцями з кібербезпеки, але й адміністраторами та аудитором, що сприяє не лише підтримці прозорості у процесах безпеки, але й допомагає організаціям своєчасно вносити корективи в свої стратегії захисту, удосконалюючи їх відповідно до нових викликів та вимог.

Типова SIEM-система складається з кількох ключових компонентів:

- агенти збору даних – відповідають за збір інформації з різних джерел: серверів, мережевих пристроїв, додатків тощо. Агенти можуть працювати як локально на пристроях, так і віддалено через мережу.

- центральне сховище – дані, зібрані з агентів, зберігаються у центральному сховищі для подальшого аналізу. Залежно від архітектури, це може бути реляційна база даних або сховище великих даних (Big Data).
- аналітичний модуль – компонент, що відповідає за аналіз та кореляцію зібраних подій. Він використовує різні алгоритми для пошуку аномалій та виявлення загроз. Найсучасніші рішення в цій сфері активно використовують методи машинного навчання та штучного інтелекту для підвищення ефективності аналізу.
- модуль сповіщень та інтерфейс користувача – інформація про потенційні загрози передається адміністраторам через спеціальні панелі управління (dashboards) або у вигляді автоматизованих сповіщень (alerts).

Разом усі ці компоненти формують єдину інтегровану систему, яка забезпечує збирання, обробку, аналіз і візуалізацію даних для підвищення рівня інформаційної безпеки. Завдяки такій інтеграції з'являється можливість не лише виявляти загрози у реальному часі, а й ефективно реагувати на інциденти, мінімізуючи їхній вплив на інформаційну інфраструктуру. На рисунку 1.1 наведено схематичне зображення взаємодії всіх вищезазначених модулів у межах типового рішення SIEM [6].

Що стосується переваг використання SIEM-систем, вони надають значні можливості для вдосконалення кібербезпеки. Оперативність систем SIEM дозволяє зменшити час реакції на безпекові події та знизити ймовірність поширення загроз у межах інфраструктури завдяки здатності аналізувати велику кількість даних з різноманітних джерел у режимі реального часу. Використання SIEM-систем сприяють швидшому виявленню загроз і потенційних інцидентів.

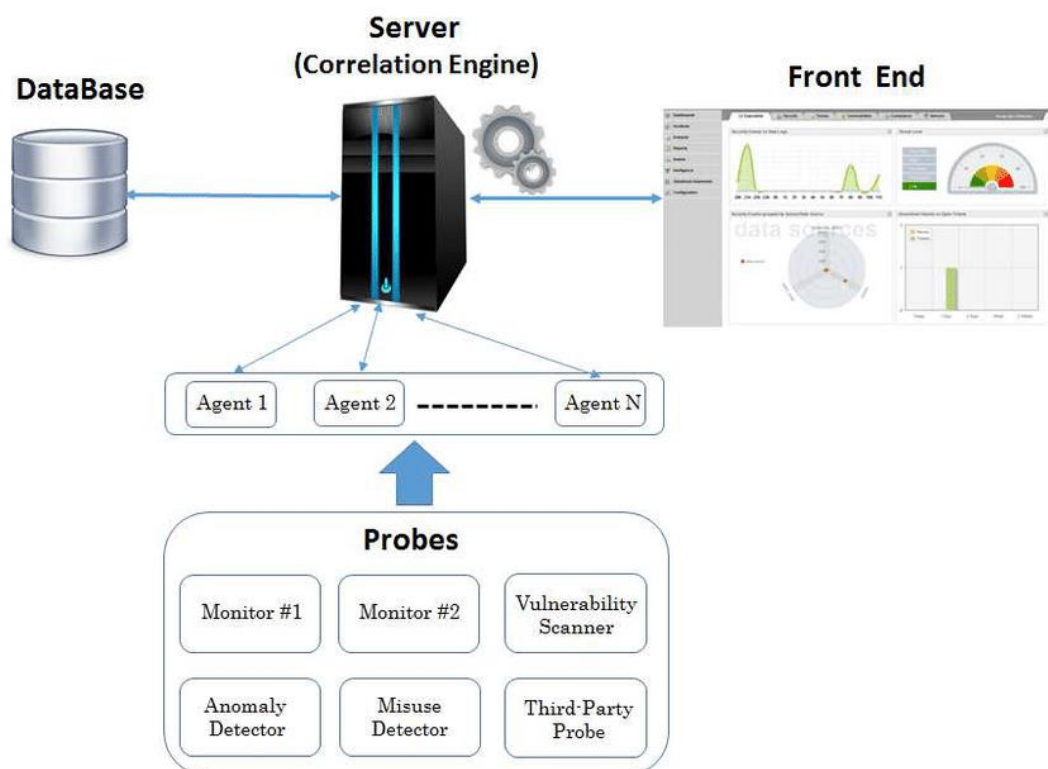


Рисунок 1.1 – Вигляд типової системи SIEM

Крім того, SIEM надає централізоване управління безпекою, забезпечуючи єдину точку доступу до всієї інформації про інфраструктуру організації, що дає змогу адміністраторам отримати повну картину про стан системи, полегшуючи як моніторинг, так і контроль за безпековими процесами. На рисунку 1.2 представлено приклад інтерфейсу системи SIEM, який демонструє, як відображаються дані про безпекові події в режимі реального часу. На цьому рисунку можна побачити корельовані події, графіки мережевої активності, сповіщення про підозрілі дії та журнали системних змін. Централізація також підвищує ефективність виявлення й усунення загроз, оскільки вся інформація доступна в одному місці.

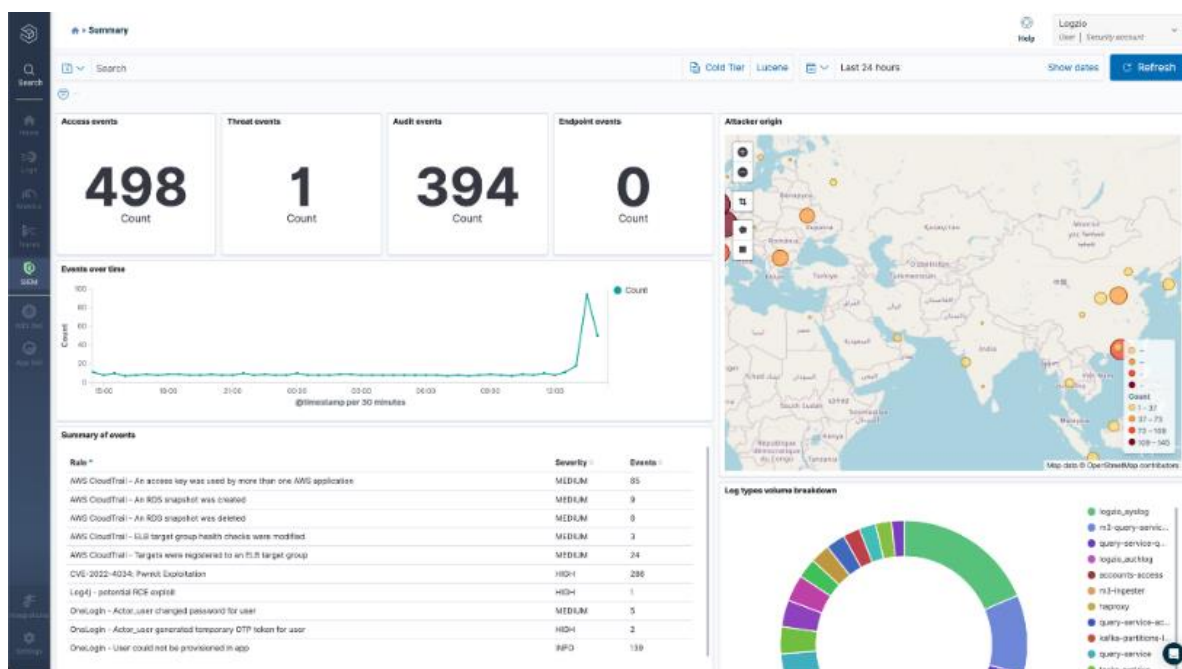


Рисунок 1.2 – Приклад вигляду системи SIEM

Ще однією важливою перевагою є підтримка відповідності регуляторним вимогам. SIEM-системи допомагають організаціям дотримуватися стандартів і нормативних актів у галузі інформаційної безпеки. Автоматизуючи створення звітів та спостереження за подіями, SIEM спрощує процес документування відповідності та допомагає знизити ризики порушень нормативних вимог, що є критичним для багатьох організацій, які мають високі вимоги до безпеки.

Незважаючи на численні переваги, використання SIEM-систем супроводжується певними викликами. Однією з основних проблем є складність налаштування і підтримки таких систем, особливо у великих організаціях з великою кількістю джерел даних. Ще одним викликом є велика кількість хибнопозитивних спрацьовувань, які можуть перевантажувати ІТ-фахівців і знижувати ефективність реагування на реальні загрози.

З огляду на це, все більш актуальним стає питання інтеграції штучного інтелекту у SIEM-системи, що дозволяє пришвидшити аналіз даних, які надходять від багатьох джерел, підвищити ефективність та час аналізу даних, а також та знизити кількість хибних тривог.

У висновку, з вищесказаного можна виділити те, що SIEM-системи є важливим інструментом у сучасній кібербезпеці, забезпечуючи збір, аналіз і кореляцію подій з різних джерел інформації для своєчасного виявлення загроз. Вони надають можливість централізованого моніторингу та управління інформаційною безпекою в реальному часі, що значно зменшує ризики успішних кібератак. Незважаючи на деякі виклики, такі як складність налаштування та висока кількість хибнопозитивних сповіщень, SIEM-системи залишаються незамінними у забезпеченні відповідності регуляторним вимогам і стандартизації процесів безпеки. Інтеграція штучного інтелекту у SIEM має потенціал суттєво підвищити ефективність таких систем, мінімізуючи недоліки та розширюючи їх можливості в аналізі великих обсягів даних та виявленні складних загроз.

1.2 Проблеми існуючих SIEM-систем при роботі з великими даними

З розвитком цифрових технологій та стрімким поширенням інтернету, обсяг даних, які генеруються у сучасних організаціях, зростає експоненційно., а дані включають широкий спектр інформації: від стандартного мережевого трафіку до складних логів, створених такими компонентами, як сервери, програмні додатки, бази даних, хмарні сервіси та мережеве обладнання. Логи є невід'ємною частиною функціонування будь-якої цифрової інфраструктури, відображаючи всі події, транзакції та операції, що відбуваються в системі. У сфері кібербезпеки ці дані набувають критичного значення, оскільки вони містять інформацію, яка може вказувати на наявність загроз, включаючи спроби несанкціонованого доступу, аномалії у функціонуванні мережі або ознаки кібератак.

Попри очевидну цінність логів для моніторингу та захисту інформаційних систем, обробка таких великих масивів даних є складним завданням, особливо для традиційних SIEM-систем. Кількість генерованих даних постійно зростає через збільшення кількості пристроїв у мережах, підвищення складності інфраструктури, активний перехід до віддаленої роботи, що призводить до того,

що SIEM-системи нерідко стикаються з низкою технічних і функціональних проблем, які обмежують їх ефективність.



Рисунок 1.3 – Виклики з якими стикаються сучасні SIEM системи

Однією з головних проблем є зниження продуктивності під час обробки великих обсягів даних. Подібні системи призначені для збору даних із багатьох джерел, таких як мережеві пристрої, сервери, додатки, що призводить до масивних потоків інформації, яку необхідно обробляти в режимі реального часу. Однак із зростанням обсягів логів виникає серйозна проблема з продуктивністю: обробка величезної кількості даних потребує значних обчислювальних ресурсів, що призводить до затримок у реагуванні на загрози. У результаті атаки можуть залишатися непоміченими, оскільки своєчасне виявлення загрози вимагає оперативної обробки даних, а будь-яка затримка може мати фатальні наслідки [6].

Водночас, велика кількість зібраних логів містить різноманітну інформацію, яка може бути як корисною, так і надмірною для ефективного аналізу. Для виявлення реальних загроз потрібно виділити саме ту інформацію, що має відношення до потенційних атак, і ігнорувати несуттєві або застарілі дані. Особливо коли лог-файли не структуровані або мають низьку якість, завдання ускладнюється у кілька разів і як результат, багато систем стикаються з

проблемами оптимізації алгоритмів обробки, оскільки звичайна методика фільтрації не завжди дає потрібний результат у масштабах великих обсягів даних. Крім того, важливою складовою є не лише швидкість обробки, а й здатність системи правильно інтерпретувати контекст даних та оперативно реагувати на зміни в мережевому середовищі.

Масштабованість є однією з найбільших проблем, з якими стикаються традиційні SIEM-системи, особливо в умовах великих організацій. У таких організаціях, де використовуються тисячі пристроїв, обсяг зібраних логів може зрости до величезних масштабів, що створює значні труднощі як для обробки, так і для зберігання цих даних, оскільки більшість SIEM-систем базуються на централізованих архітектурах зберігання, які мають обмеження щодо обсягу даних, що можуть бути оброблені та збережені на даний момент.

Якщо говорити про зберігання, то централізовані сховища не можуть забезпечити достатній рівень ефективності при масштабуванні на великі обсяги даних. Зі збільшенням кількості пристроїв і, відповідно, логів, організації потребують додаткових ресурсів для збереження та обробки нових даних, що часто вимагає інвестицій у нове обладнання та збільшення обчислювальних потужностей. Звідси як результат, поява серйозних фінансових витрат, оскільки організаціям потрібно постійно оновлювати і збільшувати обсяги зберігання, щоб відповідати вимогам системи.

Більш того, централізовані архітектури, які використовуються в традиційних SIEM-системах, створюють проблеми зі швидкістю обробки даних. Оскільки система повинна обробляти величезні обсяги логів в реальному часі, вона може стати «тягарем» для обчислювальних ресурсів і призвести до затримок у виявленні загроз, а це особливо критично для забезпечення безпеки, оскільки навіть незначні затримки можуть призвести до того, що загроза залишиться непоміченою або не буде миттєво нейтралізована.

Ще одна проблема пов'язана з великою кількістю хибнопозитивних сповіщень, які генерують SIEM-системи, адже кількість цих сповіщень зростає разом із кількістю даних, але велика частина цих попереджень є

хибнопозитивними, тобто вони сигналізують про небезпеки, яких насправді немає. Сповіднення такого типу можуть створювати інформаційний шум і перевантажувати ІТ-фахівців, які змушені витратити час на аналіз таких тривог замість того, щоб зосередитися на дійсно небезпечних інцидентах. У результаті критичні загрози можуть залишитися непоміченими або виявленими із запізненням, що підвищує ризики для безпеки організації.

Інтеграція з різними джерелами даних є ще однією складністю для SIEM-систем, проблемою є те, що у великих організаціях, де використовуються різні технології, пристрої та платформи, можуть бути труднощі зі збором та агрегацією логів, адже вони надходять із різних джерел (мережевих пристроїв, серверів, додатків, хмарних сервісів тощо). За цим стоять додаткові складнощі, що створюють проблему уніфікації форматів даних, оскільки різні джерела можуть генерувати дані в різних форматах, а відсутність єдиного стандарту для логів ускладнює їхню кореляцію та аналіз, що знижує ефективність роботи SIEM. Неповна або обмежена інтеграція з різними джерелами може призвести до втрати важливих даних і зниження точності аналізу загроз.

Окрім технічних проблем, впровадження та підтримка SIEM-систем потребує значних фінансових вкладень. Високі витрати на закупівлю програмного забезпечення, його налаштування та регулярне обслуговування можуть стати обтяжливими для організацій. Зі збільшенням обсягів даних зростають і витрати на зберігання та обробку інформації. Окрім того, SIEM-системи потребують постійного оновлення для адаптації до нових загроз, що вимагає кваліфікованого персоналу, здатного працювати з такими складними інструментами, що може додати додаткового навантаження на бюджет організації, особливо для середніх та малих компаній.

Ще однією критичною проблемою, з якою стикаються традиційні SIEM-системи, є їх недостатня адаптивність до нових кіберзагроз. Багато з цих систем спираються на підхід, заснований на сигнатурах і правилах, що ефективно працює для виявлення відомих загроз, які вже мають чіткі ознаки або сигнатури, що дозволяють швидко їх ідентифікувати. Проте цей метод не є достатньо

ефективним у випадку нових або невідомих загроз, зокрема атак нульового дня, які ще не мають відомих сигнатур. Традиційні SIEM-системи, орієнтуючись на заздалегідь визначені правила і сигнатури, не здатні виявляти нові типи атак, що робить їх вразливими до більш складних і інноваційних кіберзагроз.

Наприклад, виявлення атаки нульового дня, яка експлуатує невідому вразливість програмного забезпечення, не може бути здійснене за допомогою методів, заснованих на сигнатурах. Замість цього хакери можуть використовувати нові техніки, такі як обхід традиційних засобів захисту або модифікацію існуючого коду для реалізації атак, які не будуть виявлені стандартними сигнатурними методами. Для таких випадків SIEM-системи потребують додаткових механізмів для виявлення аномалій та поведінкових шаблонів, що допоможе виявити атаки на основі невідомих патернів, навіть якщо вони ще не мають чітко визначених сигнатур.

Як приклад, можна згадати атаку, що відбулася в 2017 році, коли було виявлено вірус Petya, який на той момент не мав чіткої сигнатури для антивірусних систем. Хоча інші інструменти безпеки могли виявити поведінкові аномалії, традиційні SIEM-системи, що покладалися на сигнатури, не змогли оперативно відреагувати. У такому випадку відсутність адаптивності призвела до серйозних наслідків для компаній, де атака залишалася непоміченою до того часу, як вона завдала значної шкоди.

Нарешті, ще одним фінальним викликом є недостатній розвиток інструментів для аналізу аномалій. Хоча аналіз поведінкових аномалій може суттєво підвищити ефективність SIEM, більшість систем мають обмежені можливості в цьому напрямку. Виявлення аномальних дій користувачів або відхилень у мережевому трафіку є ключовим для виявлення складних атак, які не відповідають стандартним шаблонам. Проте, через слабку здатність до аналізу аномалій, багато SIEM-систем не можуть виявляти подібні загрози, що робить їх менш ефективними у боротьбі з сучасними кібератаками.

Сучасні SIEM-системи стикаються з низкою серйозних проблем при роботі з великими обсягами даних, що суттєво обмежує їхню ефективність у

забезпеченні кібербезпеки. Від зниження продуктивності та масштабованості до великої кількості хибнопозитивних сповіщень і високих витрат на впровадження – виклики вимагають нових підходів до розробки та функціонування систем управління інформаційною безпекою. Особливо актуальною є потреба у вдосконаленні механізмів інтеграції та аналізу аномалій, а також у використанні інноваційних технологій, таких як штучний інтелект, для підвищення ефективності SIEM-систем при роботі з великими даними.

1.3 Штучний інтелект і машинне навчання в кібербезпеці

З кожним роком технології штучного інтелекту та машинного навчання відіграють все більшу роль у сфері кібербезпеки. В умовах стрімкого зростання обсягів даних, що генеруються в інформаційних системах, та зростання кількості кіберзагроз традиційні методи аналізу виявляються недостатньо ефективними. ШІ та ML відкривають нові можливості для автоматизації процесів виявлення загроз, аналізу логів і мережевого трафіку, а також зменшують навантаження на фахівців з кібербезпеки.

У сучасній кібербезпеці ШІ та ML стали ключовими технологіями, що допомагають у виявленні, аналізі та реагуванні на загрози. Їхнє впровадження дозволяє автоматизувати процеси обробки великих обсягів даних, що істотно підвищує ефективність систем безпеки. Використання ШІ та ML у кібербезпеці охоплює різноманітні методи та підходи, які допомагають фахівцям швидше виявляти аномалії, класифікувати загрози та реагувати на інциденти.

Нижче представлена таблиця 1.1, в якій узагальнено основні концепції та методи машинного навчання, що застосовуються в сфері кібербезпеки, їхні описи та конкретні приклади використання, що ілюструють важливість інтеграції цих технологій у сучасні системи захисту інформації.

Таблиця 1.1 – Застосування методів машинного навчання в кібербезпеці

Концепція/ Метод	Опис	Застосування в кібербезпеці
Машинне навчання (ML)	Підрозділ штучного інтелекту, що фокусується на створенні алгоритмів, які можуть навчатися на даних і робити прогнози або рішення без прямого програмування.	Використовується для виявлення аномалій у мережевому трафіку, ідентифікації шкідливих програм.
Глибинне навчання (DL)	Підмножина машинного навчання, що використовує багаторівневі нейронні мережі для обробки великих обсягів даних і виявлення складних патернів.	Використовується для класифікації зображень, аудіо, а також у виявленні шкідливого ПЗ.
Аналіз поведінки користувачів (UBA)	Метод виявлення аномалій на основі вивчення звичок і поведінки користувачів у системах.	Використовується для виявлення несанкціонованого доступу або порушень політики безпеки.
Класифікація	Процес, у якому алгоритм машинного навчання призначає категорію або клас об'єкту на основі його характеристик.	Застосовується для класифікації вхідних логів за типами загроз.
Регресія	Метод, який використовує статистичні техніки для прогнозування числових значень на основі вхідних даних.	Може бути використана для прогнозування ймовірності атаки або фінансових втрат від інцидентів.
Кластери	Техніка, яка групує схожі дані в єдині категорії, що дозволяє виявляти аномалії шляхом вивчення об'єктів, що не підходять до жодної категорії.	Застосовується для виявлення нових типів загроз або атак, які раніше не були зафіксовані.
Адаптивні алгоритми	Алгоритми, які можуть змінювати свої параметри в залежності від зміни умов або вхідних даних.	Використовуються для динамічного налаштування порогів виявлення загроз у реальному часі.

З інформації, поданої в таблиці, можемо зробити висновок, що інтеграція штучного інтелекту та машинного навчання в кібербезпеку суттєво підвищує ефективність виявлення загроз, знижує кількість хибних спрацювань і дозволяє швидко адаптуватися до нових атак. Використання цих методів не лише стане у

нагоді при виявленні вже відомих атак, але й дозволяють оперативно реагувати на нові загрози, підвищуючи адаптивність систем до змінюваного ландшафту кіберзагроз. Важливим прикладом є використання машинного навчання для автоматичного виявлення шкідливих програм, а також для прогнозування потенційних атак, що допомагає своєчасно виявляти та мінімізувати ризики [8].

Машинне навчання є підмножиною штучного інтелекту і охоплює алгоритми, які дозволяють комп'ютерним системам "навчатися" на основі даних без явного програмування. Основною метою ML є створення моделей, які можуть прогнозувати або класифікувати дані на основі тренувальних наборів.

Існує декілька основних підходів у машинному навчанні:

- Контрольоване навчання – метод машинного навчання, який використовує алгоритми для створення функції, що перетворює вхідні дані у бажані вихідні значення. В цьому процесі модель навчається на попередньо позначених даних, де кожному набору вхідних даних відповідає чітко визначений вихід.
- Неконтрольоване навчання націлене на моделювання набору вхідних даних без використання маркованих прикладів. У цьому випадку алгоритм самостійно виявляє структуру в даних і встановлює взаємозв'язки між елементами вхідного набору.
- Напівконтрольоване навчання комбінує марковані та немарковані приклади для формування відповідної функції або класифікатора. Використання немаркованих даних поряд із маркованими підвищує точність моделі, особливо коли маркованих даних небагато.
- Навчання з підкріпленням передбачає, що алгоритм вивчає оптимальну політику дій на основі спостереження за навколишнім середовищем, де кожна дія алгоритму викликає певні зміни у середовищі, яке, у свою чергу, надає зворотний зв'язок для алгоритму, що спрямовує подальше навчання.
- Трансдукція є методом, подібним до контрольованого навчання, проте не передбачає явного побудування функції. У трансдуктивному навчанні алгоритм прагне передбачити нові вихідні значення, ґрунтуючись на

наявних вхідних даних, вихідних даних навчального набору та нових вхідних значеннях.

- Навчання навчатися (learning to learn) – процес, за якого алгоритм розвиває власне індуктивне упередження, орієнтуючись на попередній досвід. Використання такого підходу дозволяє моделі адаптуватися до нових завдань, використовуючи узагальнену інформацію про те, як ефективніше виконувати навчання в подібних умовах [9].

Штучний інтелект у широкому значенні охоплює не лише машинне навчання, а й низку інших підходів і технологій, які дозволяють комп'ютерам виконувати завдання, що зазвичай асоціюються з людською розумовою діяльністю. До цього відносяться алгоритми пошуку оптимальних рішень, системи розпізнавання мови і зображень, інструменти для прогнозування поведінкових патернів, а також технології, які підвищують автономність і адаптивність систем. Наприклад, розпізнавання мови знаходить застосування в голосових асистентах, таких як Siri або Alexa, а алгоритми пошуку використовуються у маршрутизації даних та логістичних оптимізаціях [10].

Інтеграція ШІ та ML у сферу аналізу логів і мережевого трафіку суттєво збільшує ефективність цих процесів. Завдяки можливості автоматизованої обробки великих обсягів даних, ці технології дозволяють швидко ідентифікувати аномальні патерни, які можуть вказувати на загрози, зокрема атаки нульового дня, що залишаються непоміченими при використанні традиційних методів аналізу. Нейронні мережі можуть виявити приховані тенденції у мережевому трафіку, такі як підозріла активність бот-мереж, тоді як класичні алгоритми часто ігнорують подібні деталі через брак контекстної інформації.

Аналіз аномалій є одним з головних напрямів використання ШІ у кібербезпеці, адже дозволяє виявляти відхилення від типової поведінки користувачів або систем. Використовуючи алгоритми машинного навчання, система може визначати потенційні загрози або порушення безпеки на основі таких аномалій, що особливо важливо для виявлення атак нульового дня, коли ще не створено сигнатур або правил для їхньої ідентифікації [10].

Штучний інтелект відкриває нові горизонти у виявленні та аналізі складних кіберзагроз, забезпечуючи можливість кореляції подій із різноманітних джерел даних, що дозволяє об'єднувати інформацію з мережевих пристроїв, серверів, застосунків та інших компонентів інфраструктури для створення цілісної картини потенційних атак. Особливо це актуально у випадках складних багатовекторних атак або так званих "просунутих постійних загроз" (APT). Наприклад, якщо зловмисник використовує кілька точок входу – такі як злам облікового запису через фішинговий електронний лист, викрадення даних через вразливий сервер та подальше їхнє відправлення за допомогою шкідливого програмного забезпечення, – ШІ здатний об'єднати розрізнені події та виявити загальну закономірність.

Однією з найцінніших характеристик ШІ у цьому контексті є його здатність значно зменшувати кількість хибнопозитивних сповіщень [10]. У традиційних системах кібербезпеки надмірна кількість попереджень може перевантажувати фахівців, що призводить до втрати уваги до реальних загроз. Застосування методів машинного навчання дозволяє враховувати контекст кожної події, наприклад, відрізнити легітимний доступ адміністратора до системи від потенційної загрози, а також підвищує точність ідентифікації загроз і оптимізує роботу команд безпеки [8].

Крім того, алгоритми глибокого навчання, такі як рекурентні нейронні мережі (RNN) та згорткові нейронні мережі (CNN), активно застосовуються для аналізу мережевого трафіку в режимі реального часу [10]. Наприклад, ці алгоритми можуть виявляти аномалії у потоці даних, які свідчать про небезпечну активність, таку як DDoS-атаки або несанкціоноване завантаження конфіденційних файлів. Завдяки можливості глибокого аналізу, ШІ здатний не лише розпізнавати підозрілі пакети даних, але й визначати їхнє джерело та мету, що може бути використано для швидкого нейтралізування атаки.

Ще однією критично важливою функцією ШІ є розпізнавання прихованих команд, які можуть використовуватися зловмисниками для управління шкідливим програмним забезпеченням. Наприклад, під час атак із

використанням команд і контролю (C2) зловмисники можуть передавати інструкції своїм вірусам, замасковані під звичайний мережевий трафік. Алгоритми ШІ, такі як BERT чи GPT, можуть аналізувати текстові або двійкові структури команд і виявляти невідповідності або ознаки маніпуляції, які можуть свідчити про шкідливу активність.

Штучний інтелект відкриває нові перспективи у прогнозуванні та аналізі кіберзагроз завдяки своїй здатності обробляти великі обсяги історичних даних і знаходити в них приховані закономірності. На основі попередніх атак ШІ може аналізувати тенденції, визначати потенційні вектори атак або моделі поведінки, які сигналізують про підготовку до нових загроз. Наприклад, якщо атака певного типу часто передуює іншим діям зловмисників, ШІ може передбачити наступні кроки й повідомити команду кібербезпеки про можливі загрози до їхньої реалізації, що дозволяє організаціям проактивно вживати заходів і захистити свої системи від майбутніх атак.

Автоматизація аналізу логів є ще однією важливою перевагою використання ШІ [10]. У сучасних системах фіксуються мільйони записів щодня, і обробка такого масиву даних вручну є надзвичайно трудомістким процесом. ШІ здатний обробляти ці записи за лічені секунди, аналізуючи велику кількість подій і знаходячи приховані зв'язки, які залишаються непомітними при традиційному аналізі. Наприклад, зловмисники можуть використовувати декілька точок атаки для маскування своїх дій, таких як атаки на мережевий периметр, потім проникнення в систему через фішинговий електронний лист. ШІ може ідентифікувати ці події як частину єдиного сценарію атаки, навіть якщо вони здаються не пов'язаними.

Завдяки таким можливостям ШІ значно скорочується час на виявлення проблем і реагування на них. Системи стають не лише швидшими, але й точнішими, що дозволяє уникати помилкових сповіщень і зосереджуватися на реальних загрозах. А також, це особливо важливо в умовах зростання складності кібератак і збільшення обсягу даних, які необхідно обробляти для забезпечення безпеки інформаційних систем.

Штучний інтелект і машинне навчання відкривають нові горизонти у сфері кібербезпеки. Вони значно підвищують ефективність SIEM-систем, забезпечуючи швидший, точніший і адаптивніший аналіз великих обсягів даних. Завдяки можливостям автоматизованого виявлення загроз та прогнозування потенційних атак, інтеграція ШІ у системи SIEM забезпечує суттєве підвищення рівня захисту інформаційних систем, дозволяючи організаціям ефективно боротися з дедалі складнішими кіберзагрозами.

1.4 Постановка завдання

В умовах постійного зростання кіберзагроз і обсягів даних, які необхідно обробляти для забезпечення інформаційної безпеки, традиційні SIEM-системи стикаються з низкою обмежень. Основні проблеми, такі як зниження продуктивності при обробці великих обсягів даних, велика кількість хибнопозитивних сповіщень і складність виявлення складних загроз, потребують інноваційних підходів для вирішення. У цьому контексті інтеграція ШІ та ML у SIEM-системи є перспективним шляхом для підвищення їхньої ефективності та продуктивності.

Основною задачею, яку необхідно вирішити в рамках даного дослідження, є покращення існуючих SIEM-систем за рахунок інтеграції ШІ, що дозволить покращити виявлення загроз у реальному часі, знизити кількість хибнопозитивних сповіщень, оптимізувати процес аналізу логів і мережевого трафіку, а також підвищити рівень автоматизації процесів виявлення та реагування на загрози.

Інтеграція ШІ має вирішити такі конкретні задачі:

- автоматизувати процес обробки великих обсягів логів та мережевого трафіку;
- підвищити точність виявлення кіберзагроз;
- знизити кількість хибнопозитивних сповіщень;
- забезпечити своєчасне виявлення складних атак.

Постановка завдання інтеграції штучного інтелекту в SIEM-систему має на меті вирішення низки критичних проблем, з якими стикаються сучасні системи інформаційної безпеки. Виконання цих завдань дозволить створити систему, яка не тільки ефективно працюватиме в умовах великих обсягів даних, але й забезпечить високий рівень автоматизації та масштабованості для сучасних організацій.

1.5 Висновок до розділу

Отже, у цьому розділі проведено огляд сучасних SIEM-систем, розглянуто їхню роль у забезпеченні кібербезпеки, а також проаналізовано ключові проблеми, з якими стикаються такі системи при роботі з великими обсягами даних. SIEM-системи є важливими інструментами для централізованого збору та аналізу даних про безпеку, однак їхні обмеження пов'язані зі зростанням складності мережевого трафіку та кібератак.

Основними викликами для SIEM-систем є низька масштабованість, обмежена здатність обробляти великі дані в реальному часі, а також велика кількість хибних тривог, що ускладнює роботу аналітиків. Використання застарілих правил та сигнатурних методів для виявлення загроз вимагає значних зусиль з постійного оновлення й налаштування системи, що робить її менш ефективною в умовах швидкозмінного середовища кіберзагроз.

Інтеграція штучного інтелекту та методів машинного навчання у SIEM-системи відкриває нові можливості для вдосконалення процесу моніторингу та аналізу логів і мережевого трафіку. ШІ дозволяє не тільки автоматизувати процес виявлення загроз, але й робить систему здатною адаптуватися до нових типів атак, що зменшує кількість хибних спрацьовувань і підвищує загальну ефективність захисту.

Таким чином, висвітлені в цьому розділі питання створюють основу для подальшого дослідження й розробки методів інтеграції ШІ у SIEM-системи з метою покращення їхньої ефективності та надійності при роботі з великими обсягами даних.

2 ОБГРУНТУВАННЯ ОБРАНИХ РІШЕНЬ

2.1 Джерела логів, особливості їх структури та алгоритми обробки

Для забезпечення всебічного моніторингу та аналізу інформаційної безпеки, SIEM-системи збирають і обробляють великий обсяг даних, які надходять у вигляді логів від різних пристроїв та додатків в інфраструктурі підприємства. Логи – структуровані або неструктуровані записи подій, які надають інформацію про роботу системи та окремих її компонентів [11]. Оскільки джерела логів та їхній формат є неоднорідними, обробка та аналіз цих даних стає складним завданням, яке потребує ефективних методів обробки та алгоритмів аналізу, включаючи штучний інтелект та машинне навчання. Нижче наведено докладний огляд джерел логів, їх структурних характеристик та сучасних підходів до обробки, що дозволяють максимально використовувати цю інформацію для кіберзахисту підприємства.

Логи генеруються практично всіма елементами ІТ-інфраструктури, і їхня роль у системах моніторингу безпеки полягає у фіксації подій, які допомагають оцінити стан і безпеку роботи системи [11]. Розглянемо основні джерела логів та їх особливості.

- 1) Мережеве обладнання – один із найбільших джерел логів у будь-якій організації. У цей список входять роутери, комутатори, фаєрволи та інші мережеві пристрої генерують логи, що містять інформацію про мережеві з'єднання, маршрутизацію, налаштування, а також про підозрілі або незвичні дії. Вони допомагають виявляти аномальну мережеву активність і можливі спроби несанкціонованого доступу.
- 2) Сервери – створюють записи про події, які стосуються роботи операційної системи, виконуваних процесів, взаємодії з користувачами та стану безпеки. Серверні логи містять інформацію про запуск і завершення роботи, можливі збої, системні помилки та інші події. Наприклад, серверні логи можуть вказувати на раптове завершення

важливого процесу, що може свідчити про спробу злому або збої в роботі сервера.

- 3) Базы даних – їхні логи містять записи про транзакції, запити та будь-які зміни даних. Такі дані є критично важливими, оскільки більшість конфіденційної інформації зберігається у базах даних, і ці логи можуть фіксувати можливі спроби несанкціонованого доступу, SQL-ін'єкції або маніпуляції з даними.
- 4) Додатки – генерують логи, які фіксують взаємодію користувачів із програмним забезпеченням, а також внутрішні операції та помилки в роботі додатків. Вони можуть включати відомості про запити до серверів, зміни в налаштуваннях користувачів, виникнення помилок або збоїв в роботі. Наприклад, логи веб-додатків містять інформацію про IP-адреси користувачів, які підключаються до системи, параметри запитів і тривалість їх обробки. Це дозволяє виявляти спроби злому, DDoS-атаки, а також оцінювати загальне навантаження на додаток.
- 5) Антивірусні системи та інші засоби безпеки – зберігають записи про виявлені загрози, такі як віруси, шкідливе ПЗ або підозрілі файли. Такі логи містять дані про ідентифіковані загрози, дії, які було вжито для їхнього нейтралізування, та інші події, що стосуються безпеки.

Щодня IT-інфраструктура середнього та великого підприємства може генерувати значні обсяги логів. За даними звіту SolarWinds, великі компанії можуть створювати в середньому 100 ГБ логів щодня, що може перевищувати 3 ТБ на місяць [12]. Такі великі обсяги даних потребують ефективних засобів зберігання та аналізу, щоб підтримувати швидке виявлення потенційних загроз та забезпечувати належний рівень кібербезпеки.

Логи можуть бути представлені у форматах JSON, XML, текстових файлів (наприклад, Syslog) або в спеціалізованих форматах, розроблених конкретними виробниками пристроїв чи програмного забезпечення.

Нерідко різні джерела логів мають неузгоджену структуру, що значно ускладнює їхню кореляцію та уніфікований аналіз.

Однією з основних проблем у роботі з логами є їхня неструктурованість, адже кожне джерело фіксує дані у своєму форматі, що може змінюватися від системи до системи.

Мережеві логи можуть включати IP-адреси, час підключення, тип протоколу, розмір переданих пакетів тощо. Проблема в тому, що таких записів за кілька хвилин виникає чимало, і всі потрібно аналізувати. Приклад вигляду списку мережевих логів представлено на рисунку 2.1.

timestamp
2024-11-07 16:36:40.001
NetFlowV5 [192.168.137.1]:5353 <> [224.0.0.251]:5353 proto:17 pkts:14 bytes:2203
2024-11-07 16:36:30.002
NetFlowV5 [192.168.137.1]:137 <> [192.168.137.255]:137 proto:17 pkts:3 bytes:234
2024-11-07 16:36:25.000
NetFlowV5 [192.168.137.1]:62234 <> [224.0.0.252]:5355 proto:17 pkts:2 bytes:100
2024-11-07 16:36:25.000
NetFlowV5 [192.168.137.1]:57521 <> [224.0.0.252]:5355 proto:17 pkts:2 bytes:100
2024-11-07 16:36:25.000

Рисунок 2.1 – Список мережевих логів

Системні логи можуть містити стандартні повідомлення про стан системи (успішне завершення процесів, критичні помилки), а також детальну інформацію про процеси та використання ресурсів. Наприклад, у Linux-системах стандартні логи зберігаються у форматах, характерних для Syslog, що також можуть змінюватися залежно від налаштувань адміністратора, та приклад вигляду таких логів зображено на рисунку 2.2.

```

2024-11-07 16:35:55.000
user-VirtualBox dbus-daemon[2513]: [system] Activating via systemd: service name='org.freedesktop.nm_dispatcher'
no-daemon " label="unconfined")

2024-11-07 16:35:55.000
user-VirtualBox dbus-daemon[2513]: [system] Successfully activated service 'org.freedesktop.nm_dispatcher'

2024-11-07 16:35:55.000
user-VirtualBox systemd[1]: Started Network Manager Script Dispatcher Service.

2024-11-07 16:35:55.000
user-VirtualBox avahi-daemon[2511]: Withdrawing address record for fe80::2827:f122:e092:1a9e on ens33.

2024-11-07 16:35:55.000
user-VirtualBox NetworkManager[9832]: <info> [1730986555.7279] dhcp4 (ens33): state changed no lease

```

Рисунок 2.2 – Список системних логів

Логи додатків включають дані про запити, помилки, тривалість виконання операцій, а також іншу інформацію, специфічну для кожного конкретного додатка. Наприклад, веб-додатки можуть фіксувати спроби входу в систему, HTTP-запити і статуси відповідей.

```

203.0.113.1 - - [07/Nov/2024:10:15:30 +0000] "GET /home HTTP/1.1" 200 1024
198.51.100.23 - - [07/Nov/2024:10:16:02 +0000] "POST /login HTTP/1.1" 401 512
192.168.1.15 - - [07/Nov/2024:10:16:45 +0000] "GET /dashboard HTTP/1.1" 200 2048
203.0.113.4 - - [07/Nov/2024:10:17:11 +0000] "DELETE /api/item/1234 HTTP/1.1" 204 0
198.51.100.27 - - [07/Nov/2024:10:17:56 +0000] "PUT /profile/update HTTP/1.1" 500 1024

```

Рисунок 2.3 – Список логів web-сервера

Обробка логів у SIEM-системах виконується у кілька етапів, кожен з яких спрямований на впорядкування даних, виявлення кореляцій та ідентифікацію потенційних загроз.

1) Збір та нормалізація даних.

На цьому етапі всі логи надходять до системи для попередньої обробки. Нормалізація включає переведення даних у єдиний формат, який дозволяє структурувати різнотипні дані для подальшого аналізу. Системи на зразок Logstash і Graylog використовують різні механізми для збору та нормалізації даних. Наприклад, Logstash може отримувати дані з бази даних, мережних

пристроїв або додатків, фільтрувати їх і переводити у формат, сумісний з Elasticsearch, що використовується для подальшого аналізу.

2) Фільтрація та відсіювання надлишкових даних.

Часто обсяги зібраних логів є надзвичайно великими, і серед них значна частина є надлишковою або повторюваною. Алгоритми фільтрації можуть автоматично відсіювати логи, які не мають цінності для моніторингу безпеки, і таким чином зменшувати обсяг даних, що підлягають аналізу.

3) Кореляція подій.

Кореляція – це процес співставлення подій із різних джерел для виявлення аномальних чи небезпечних дій. Кореляційні алгоритми SIEM-систем допомагають об'єднувати записи з різних логів на основі певних параметрів, таких як IP-адреса, час події або користувач. Кореляція дозволяє визначати, наприклад, чи є незвичне підключення до сервера, за яким слідує велика кількість запитів до бази даних, потенційною загрозою.

4) Аналіз за допомогою алгоритмів машинного навчання та штучного інтелекту.

Сучасні методи аналізу логів включають використання ШІ та машинного навчання для ідентифікації загроз. Нейронні мережі, алгоритми кластеризації та дерева рішень дозволяють знаходити приховані закономірності в даних, автоматично класифікувати події та виявляти аномалії. Наприклад, алгоритми класифікації можуть навчитися розрізняти нормальну та аномальну активність, виходячи з історичних даних, і підвищувати точність детекції загроз без необхідності постійного людського втручання.

5) Реалізація в SIEM-системах.

Більшість сучасних SIEM-рішень, таких як Splunk, IBM QRadar та Elastic Security, підтримують модулі штучного інтелекту та машинного навчання для автоматизації процесів аналізу логів. Такі системи використовують алгоритми для автоматичної обробки та аналізу великих обсягів логів, що дозволяє знизити навантаження на аналітиків безпеки, зменшити кількість хибнопозитивних сповіщень і підвищити точність виявлення загроз.

Ефективна обробка логів забезпечує підприємствам надійний захист, адже дозволяє оперативно реагувати на потенційні загрози та автоматично блокувати підозрілу активність. Проте робота з великими обсягами даних має свої виклики – необхідність у масштабованих обчислювальних ресурсах, налаштуванні точних алгоритмів фільтрації та кореляції, а також управлінні складними моделями машинного навчання.

Загалом, впровадження алгоритмів обробки логів із використанням ШІ дозволяє SIEM-системам швидко адаптуватися до нових типів загроз, оптимізуючи процеси моніторингу та аналізу інформації в реальному часі, що критично важливо для підтримки високого рівня кібербезпеки на сучасних підприємствах. Інформацію про це буде описано в наступному розділі.

2.2 Моделювання мережевого трафіку для виявлення аномалій за допомогою методів штучного інтелекту

Аналіз мережевого трафіку з метою виявлення аномалій є одним із основних завдань забезпечення кібербезпеки в сучасних інформаційних системах. Використання ШІ для моделювання мережевого трафіку дозволяє створювати складні моделі, які можуть ефективно виявляти аномальні події, навіть у випадках нових або невідомих атак. Виявлення аномалій у трафіку дозволяє не лише вчасно виявляти атаки, але й оптимізувати роботу систем моніторингу та захисту, таких як SIEM.

Математичне моделювання виявлення аномалій у мережевому трафіку є важливим етапом у створенні ефективних систем для автоматизованого моніторингу та реагування на загрози. В основі таких моделей лежить опис мережевих даних, їх аналіз і порівняння з нормальними патернами. Аномалія в контексті мережевого трафіку визначається як відхилення від цих патернів, що може вказувати на несанкціоновану активність, атаки або системні збої.

Основними етапами математичного моделювання є:

- опис трафіку;
- побудова нормальної моделі;
- виявлення аномалій.

Етап опису трафіку відбувається наступним чином – мережевий трафік описується за допомогою числових характеристик, таких як частота пакетів, обсяг даних, час між пакетами, типи протоколів і портів тощо.

Наступним етапом йде створення математичних моделей, які описують типову, «нормальну» поведінку мережі, для яких можуть бути використані статистичні методи, методи машинного навчання чи глибоке навчання.

Останнім етапом йде порівняння поточного стану трафіку з моделлю «нормальної» поведінки дозволяє виявити відхилення. Якщо показники поточного трафіку значно відрізняються від тих, що вважаються нормальними, це є сигналом про можливу аномалію.

Для ефективного виявлення аномалій у мережевому трафіку широко застосовуються й статистичні методи та алгоритми машинного навчання. Статистичні методи зазвичай базуються на припущенні, що нормальний трафік має певний статистичний розподіл, відхилення від якого є аномалією [13].

- Методи оцінки ймовірності, для прикладу максимізація правдоподібності, дозволяють побудувати ймовірнісні моделі трафіку, які потім використовуються для визначення аномалій. Ефективність цих методів є високою в тих випадках, коли розподіл трафіку можна описати певною функцією ймовірності.
- Методи кластеризації, такі як k-середніх (k-means) або DBSCAN, дозволяють групувати схожі за характеристиками потоки трафіку. Аномалії можуть бути виявлені як потоки, що не належать до жодного з кластерів.

Машинне навчання дозволяє автоматизувати процес навчання моделей на основі великих масивів даних, що значно підвищує точність виявлення аномалій. Одним із основних підходів є використання алгоритмів без нагляду, таких як

автокодери та класифікатори, що дозволяють працювати з даними без попередніх міток.

Нейронні мережі, зокрема рекурентні нейронні мережі (RNN) та їхні розширення, такі як LSTM (довго- та короткочасна пам'ять), є потужним інструментом для моделювання часових рядів, таких як мережевий трафік. Вони дозволяють не лише враховувати поточний стан мережі, а й історичні дані про її поведінку, що є важливим для виявлення складних аномалій, які не можуть бути зловлені статичними методами [13].

Автокодери також широко застосовуються в цій сфері. Вони навчаються відновлювати вхідні дані, і будь-які відхилення в процесі відновлення можуть вказувати на аномалії в трафіку. Такі моделі особливо корисні, коли немає чіткої мітки для даних або коли необхідно виявити невідомі раніше типи атак.

Ефективність моделей виявлення аномалій оцінюється через різні метрики, такі як точність, чутливість, специфічність та F1-міра. Для порівняння результатів різних моделей також використовуються ROC-криві та матриця плутанини, що дозволяє побудувати більш точну картину роботи системи.

Крім того, важливим етапом є перевірка узгодженості моделі з реальними умовами експлуатації. Моделі повинні бути адаптовані до змінних умов мережевого трафіку, тому часте перенавчання та тюнінг моделей є необхідними для досягнення високої ефективності в реальному часі.

Отже, зібравши всю цю інформацію, її було проаналізовано та графічно зображено схему збору та подальшої обробки логів за допомогою методів штучного інтелекту, що зображена на рисунку 2.4. Основна суть процесу складається з двох модулів – збору логів та обробки їх за допомогою ШІ, збір логів відбувається з різних систем та пристроїв, ці дані агрегуються і передаються далі для перевірки на наявні загрози або компрометації.

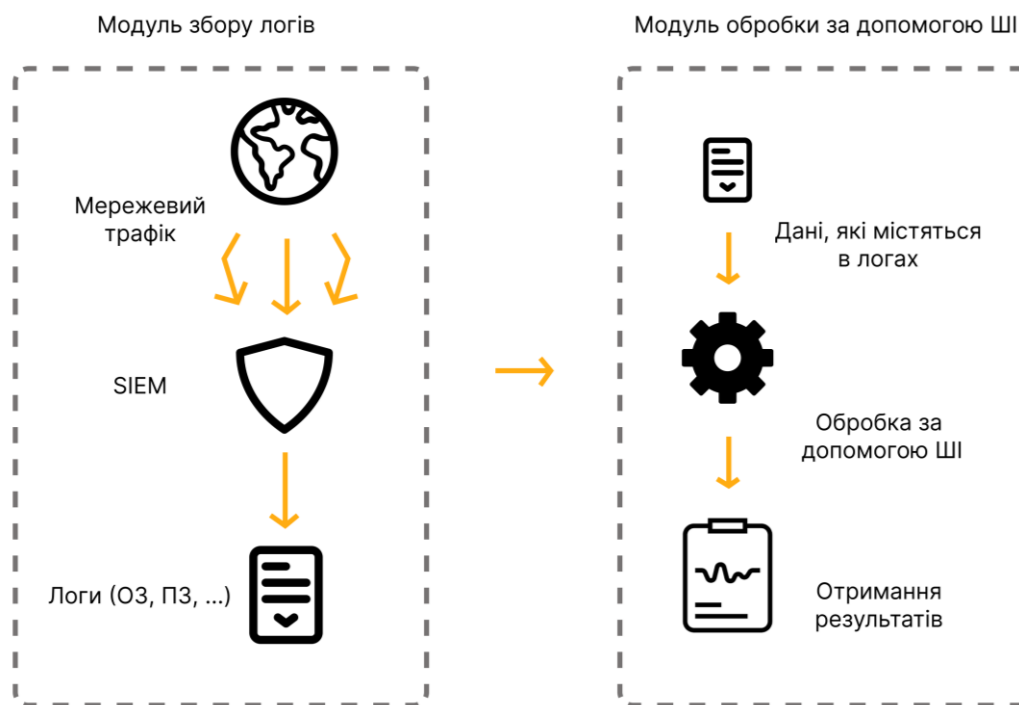


Рисунок 2.4 – Графічна схема етапів обробки логів з ШІ

Моделювання мережевого трафіку для виявлення аномалій за допомогою методів штучного інтелекту є ключовим елементом у підвищенні ефективності сучасних систем моніторингу безпеки. Завдяки використанню математичних моделей та алгоритмів машинного навчання можна автоматизувати процес виявлення атак, зменшити кількість хибних спрацьовувань і знизити навантаження на аналітиків. Водночас важливо розуміти, що ці моделі повинні постійно адаптуватися до нових умов і вдосконалюватися в процесі експлуатації.

2.3 Модель обробки логів на основі ШІ

Збір логів та їх подальша обробка є критично важливим етапом для моніторингу та аналізу діяльності комп'ютерних систем, мереж і серверів, так як вони містять інформацію про виконувані операції, статуси сервісів, помилки та інші важливі події, що можуть бути індикаторами несанкціонованої активності або збоїв у системі. Використання ШІ для обробки цих логів відкриває нові можливості для автоматизованого виявлення загроз та аномалій у реальному часі, що значно підвищує ефективність систем моніторингу та захисту.

Системи обробки логів (Log Management Systems) збирають, зберігають, індексують і аналізують логи з різних джерел, таких як сервери, мережеві пристрої, бази даних та інші компоненти інформаційної інфраструктури. Такі системи можуть включати в себе різноманітні інструменти для збору, зберігання та аналізу логів, як-от Graylog, ELK Stack (Elasticsearch, Logstash, Kibana), Splunk і інші. Вони використовують різноманітні методи для фільтрації та агрегації даних, пошуку по індексах і створення аналітичних звітів.

Завдяки логам можна отримати інформацію про аномальну поведінку, невідомі уразливості або навіть загрози на етапі їх появи. Однак через величезний обсяг даних та необхідність аналізувати численні події, вручну визначити всі критичні інциденти є надзвичайно складним завданням. Тут і вступає в гру штучний інтелект.

Використання штучного інтелекту для обробки логів дозволяє автоматизувати виявлення аномалій, що значно зменшує час на реагування на інциденти та допомагає аналітикам зосередитися на більш важливих аспектах безпеки. Різноманітні алгоритми машинного навчання, зокрема методи класифікації, кластеризації та виявлення аномалій, дають змогу ефективно працювати з великими обсягами логових даних, виявляючи корисні патерни для подальшого аналізу.

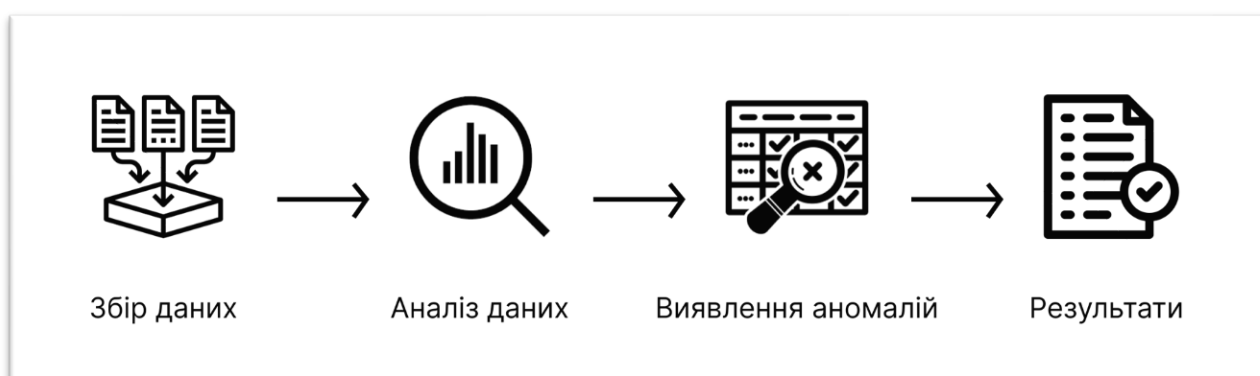


Рисунок 2.5 – Етапи обробки логів

Основні етапи обробки логів на основі ІІІ представлені на рисунку 2.4.

1) Збір та передобробка даних.

Етап, на якому здійснюється збір логів з різних джерел, їх фільтрація, нормалізація і перетворення у формат, зручний для аналізу. Важливими аспектами є визначення типу даних, їх кодування і усунення неточностей.

2) Аналіз і класифікація даних.

Крок автоматичного визначення подій є важливим для подальшого аналізу. Машинне навчання та штучний інтелект дозволяють створювати моделі для класифікації подій, таких як розпізнавання шкідливих дій, виявлення аномальних спроб доступу, виявлення атак і т. д.

3) Виявлення аномалій

Використовуючи алгоритми без нагляду (наприклад, кластеризацію або методи глибокого навчання) або великі мовні моделі, можна виявити нові або невідомі загрози, які не були б помічені традиційними методами. Виявлення таких аномалій дозволяє оперативно реагувати на зловмисні дії, навіть якщо вони не підпадають під відомі патерни атак.

4) Інтерпретація та візуалізація результатів.

Представлення результатів у вигляді, зручному для аналітиків. Візуалізація даних (наприклад, за допомогою графіків або таблиць) допомагає швидко оцінити ситуацію і приймати обґрунтовані рішення.

Використання штучного інтелекту для обробки логів не обмежується лише загальними методами виявлення аномалій або класифікації подій. Для ефективного аналізу великих обсягів інформації, що зберігається у вигляді логів, необхідно застосовувати специфічні методи, що дозволяють витягувати корисні патерни та розпізнавати схеми взаємодії між компонентами системи. Такі методи включають в себе не тільки традиційні підходи до аналізу, але й нові, спеціалізовані моделі, що дозволяють вирішувати специфічні завдання, характерні для обробки системних логів.

Одним із основних напрямків є застосування глибоких нейронних мереж для моделювання послідовностей подій, записаних у логах. Враховуючи, що логи часто мають часову залежність, рекурентні нейронні мережі (RNN) та їхні

варіанти, зокрема LSTM (Long Short-Term Memory), дозволяють моделювати ці часові залежності. За допомогою цих моделей можна будувати прогнози для майбутніх подій або виявляти відхилення в поведінці системи на основі попередніх записів логів. Наприклад, можна навчити модель на попередніх записах логів для виявлення нетипових дій, таких як спроби несанкціонованого доступу або виконання підозрілих команд [14].

Особливістю використання глибоких нейронних мереж є їхня здатність аналізувати складні та нелінійні взаємозв'язки між різними подіями у великих обсягах даних, а моделі, зокрема рекурентні та згорткові нейронні мережі, демонструють високу ефективність у виявленні як типових, так і прихованих загроз, що не піддаються простому аналізу через стандартні правила або сигнатури. Наприклад, нейронні мережі здатні виявляти поведінкові аномалії в трафіку, які можуть сигналізувати про приховані атаки, як-от складні фішингові кампанії або DDoS-атаки зі змінною інтенсивністю.

Автокодери – окремий тип глибоких нейронних мереж – є одним із інструментів для виявлення аномалій у логах, вони навчаються відтворювати нормальний стан системи, створюючи латентне представлення даних, яке фокусується лише на найбільш суттєвих характеристиках. У випадку, коли вхідні дані суттєво відрізняються від нормального стану, автокодери не здатні точно їх реконструювати і така невідповідність сигналізує про аномалію, яка може свідчити про спробу компрометації системи або інші нетипові події.

Наприклад, в умовах корпоративної мережі автокодери можуть навчитися розпізнавати звичайні патерни доступу користувачів до баз даних. Якщо ж користувач із невласивими для нього повноваженнями намагається отримати доступ до конфіденційних даних у незвичний час доби або з нетипового пристрою, модель ідентифікує відхилення від норми і може негайно надіслати попередження адміністратору безпеки.

Іншим важливим аспектом обробки системних логів є семантичний аналіз текстових записів. Логи, що генеруються різними пристроями та додатками, часто містять значну кількість текстових повідомлень, які можуть бути важкими

для інтерпретації без належної обробки. Тому для ефективного розуміння значення і контексту подій необхідно застосовувати методи обробки природної мови (NLP). Зокрема, сучасні методи NLP, такі як моделі трансформерів, зокрема BERT (Bidirectional Encoder Representations from Transformers) та GPT (Generative Pretrained Transformer), дозволяють не тільки автоматично класифікувати текстові повідомлення, але й розуміти контекст, в якому ці повідомлення були згенеровані.

Моделі на основі трансформерів, зокрема GPT, можуть аналізувати текстові дані з різних джерел і виявляти приховані значення, що допомагає у більш точному визначенні інцидентів або загроз. Наприклад, коли система реєструє повідомлення, яке містить інтерпретацію несанкціонованого доступу, вона не просто фіксує саму подію, але й оцінює її в контексті того, що сталося перед цим, чи є подібні попередні події і який саме тип атак може бути виявлений. У випадку з GPT модель може детектувати текстові патерни, що вказують на підозрілу активність, навіть якщо конкретні сигнатури атак ще не були зафіксовані [15].

Одним з прикладів застосування трансформерних моделей є класифікація повідомлень за категоріями загроз. Наприклад, логи можуть містити повідомлення про невдалі спроби входу в систему, доступ до конфіденційних даних або спроби виконати несанкціоновані операції. Традиційні методи аналізу можуть просто зафіксувати ці події як серії помилок або аномалій. Однак з використанням NLP та моделей на основі трансформерів можна не лише виявити ці події, а й класифікувати їх за рівнем загрози, а також визначити ймовірні причини цих інцидентів, що дозволяє значно підвищити ефективність реагування на загрози.

Моделі трансформерів можуть також допомогти у виявленні складних взаємозв'язків між подіями. Наприклад, якщо декілька незначних подій відбуваються в короткий проміжок часу або з певною географічною закономірністю, модель може ідентифікувати їх як потенційну атаку, навіть якщо кожна окрема подія здається незначною. Також моделі трансформерів,

здатні обробляти великі обсяги текстових даних, можуть допомогти у виявленні бажливих патернів в тексті логів. Наприклад, можна виявляти повторювані шаблони або дивні комбінації слів, що вказують на зловмисну активність. Семантичний аналіз дозволяє також автоматично класифікувати типи подій (помилки, попередження, успішні спроби входу тощо), що полегшує подальший аналіз.

Використання такого підходу дозволяє здійснювати аналіз в реальному часі та може бути корисним для виявлення складних загроз, таких як фішинг, соціальна інженерія або внутрішні загрози.

Застосування методів штучного інтелекту для обробки логів є важливим етапом розвитку сучасних систем моніторингу та захисту. Використання глибоких нейронних мереж, семантичного аналізу текстів, когнітивних моделей та графових підходів дозволяє значно підвищити ефективність виявлення аномалій та зниження числа хибних спрацьовувань. Зважаючи на постійну еволюцію кіберзагроз, ці методи дозволяють адаптувати системи до нових умов і забезпечувати належний рівень безпеки для інформаційних систем.

Інтеграція методів ШІ в традиційні системи обробки логів, такі як Graylog, ELK Stack або Splunk, дозволяє значно покращити їх функціональність, наприклад:

- автоматизація збору та попереднього аналізу логів (завдяки машинному навчанню можна автоматично виявляти важливі події та фільтрувати їх за допомогою правил та моделей);
- реальний час обробки (ШІ дозволяє аналізувати логи в реальному часі, швидко реагуючи на аномалії і загрози);
- інтелектуальні сповіщення (на основі результатів аналізу можна автоматично генерувати сповіщення про підозрілі події, що дозволяє зменшити навантаження на операторів та системи безпеки).

Попри очевидні переваги, інтеграція ШІ в обробку системних логів стикається з низкою викликів:

- якість і повнота даних (неякісні або неповні логи можуть призвести до помилкових результатів аналізу).
- обсяг даних (великі обсяги логів можуть бути складними для обробки навіть за допомогою сучасних методів ШІ, що вимагає високих обчислювальних потужностей).

Моделі обробки системних логів на основі штучного інтелекту стають важливим елементом сучасних систем моніторингу безпеки. Вони дозволяють значно підвищити точність виявлення загроз, знизити навантаження на операторів та підвищити ефективність реагування на інциденти. Проте важливою складністю є забезпечення якості даних та оптимізація алгоритмів для роботи з великими обсягами логів. Враховуючи це, подальший розвиток ШІ у цій галузі дозволить значно покращити безпеку інформаційних систем.

2.4 Інтеграція ШІ у традиційну архітектуру SIEM

Інтеграція штучного інтелекту в архітектуру традиційних SIEM-систем є новим етапом розвитку інструментів для захисту від кіберзагроз. Традиційні SIEM-системи, що зазвичай покладаються на фіксовані правила та шаблони для виявлення інцидентів, мають певні обмеження, для прикладу, такі системи потребують постійного ручного оновлення правил, через що часом не здатні швидко реагувати на нові типи загроз. У впровадженні ШІ в архітектуру таких систем криється можливість значно підвищити ефективність роботи завдяки автоматичному аналізу даних та адаптації до нових моделей поведінки. Такий підхід дозволяє системі самостійно розпізнавати нові типи загроз, аналізуючи в реальному часі великий обсяг даних та виявляючи аномалії, що можуть бути ознакою потенційного вторгнення або кіберзлому [16].

Ще однією суттєвою перевагою впровадження штучного інтелекту у систему SIEM є здатність значно зменшувати кількість хибних спрацьовувань, які є серйозною проблемою для традиційних систем моніторингу безпеки. Традиційні SIEM-системи часто генерують велику кількість тривог, більшість із яких виявляються помилковими, що додаткове навантаження на команди

кібербезпеки, які змушені витрачати ресурси на перевірку цих тривог. З часом така ситуація може призводити до зниження продуктивності команд і навіть до «втоми від тривог», коли реальні загрози залишаються поза увагою через надмірну кількість хибних сигналів.

Інтеграція штучного інтелекту, зокрема алгоритмів машинного навчання, дозволяє значно покращити точність аналізу. Алгоритми, здатні навчатися на історичних даних, аналізують поведінкові шаблони користувачів і систем, а також враховують контекст подій у реальному часі. Наприклад, поведінковий аналіз може допомогти визначити, чи є підозріла активність реальною загрозою або звичайною дією, яка раніше вже спостерігалася в системі. Після цього з'являється можливість автоматично фільтрувати велику кількість тривог і залишати на розгляд операторів лише ті події, які мають високу ймовірність загрози [16].

Окрім того, алгоритми ШІ можуть використовувати контекстуальний аналіз для виявлення аномалій, які не відповідають стандартним шаблонам, і зменшувати ймовірність помилкових висновків. Конкретніше про використання контекстів, система може враховувати географічне розташування, час доби, історію доступів користувача до певних ресурсів і поєднувати ці фактори для формування точного висновку про наявність загрози.

ШІ також значно розширює можливості SIEM-систем щодо інтеграції зі сторонніми джерелами інформації. Замість обробки лише внутрішніх логів, система з ШІ може отримувати та обробляти дані з зовнішніх ресурсів, таких як соціальні мережі, публічні бази даних загроз та інші інформаційні системи. Таке розширення дає змогу отримувати актуальні дані про нові загрози та швидше адаптувати систему безпеки до змінного середовища.

Завдяки інтеграції штучного інтелекту, SIEM-системи мають можливість не лише аналізувати минулі інциденти, але й отримувати здатність до прогнозування потенційних загроз. Використання передбачувальних моделей на основі штучного інтелекту дозволяє виявляти приховані закономірності та

тенденції в масивах даних, що часто залишаються непомітними для традиційних підходів.

Передбачувальні моделі, що побудовані на аналізі історичних інцидентів і трендів, мають здатність ідентифікувати поведінку, що передує загрозам. Наприклад, ШІ може виявити аномальну активність у мережі, яка відповідає початковим етапам кібератаки, таким як сканування портів чи тестування наявних вразливостей. Завдяки такому підходу система може завчасно попередити операційний центр безпеки про можливий ризик ще до того, як загроза реалізується, а також це значно скорочує час на реагування і дозволяє вжити заходів для мінімізації потенційної шкоди.

Швидкість і ефективність реагування на інциденти значно підвищуються завдяки впровадженню штучного інтелекту в системи SIEM. Замість трудомісткого аналізу інцидентів вручну, який вимагає багато часу і людських ресурсів, автоматизовані системи з інтегрованим ШІ здатні виявляти загрози практично миттєво, аналізуючи логи в режимі реального часу. Це особливо важливо в умовах сучасного кіберпростору, де швидкість реагування може бути критичною для мінімізації шкоди, завданої атакою.

Штучний інтелект дозволяє не тільки ідентифікувати інциденти, але й автоматично визначати їх пріоритетність, спираючись на оцінку потенційних ризиків. Для прикладу, інциденти, які можуть спричинити значні фінансові втрати чи порушення роботи критично важливих систем, аналітично розраховуються і отримують найвищий пріоритет.

Водночас впровадження штучного інтелекту в архітектуру SIEM-систем супроводжується низкою викликів, які потребують комплексного вирішення. Одним із ключових аспектів є необхідність точного налаштування алгоритмів машинного навчання. Для цього потрібно не лише визначити оптимальні параметри роботи моделей, а й забезпечити їх навчання на великих обсягах якісних даних. Наявність таких даних є критичною, адже точність і ефективність роботи алгоритмів безпосередньо залежить від репрезентативності та повноти навчальних вибірок. Проте, збір і підготовка даних можуть вимагати значних

зусиль та ресурсів, зокрема для забезпечення їх структурованості, очищення від шумів і аномалій, а також відповідності до специфіки задачі.

Ще одним суттєвим викликом є необхідність забезпечення безпеки самої системи штучного інтелекту. Оскільки ШІ виступає ключовим елементом у процесі аналізу та реагування на загрози, він може стати потенційною мішенню для зловмисників. Успішна компрометація системи ШІ може не лише вивести її з ладу, але й спровокувати помилкові дії, наприклад, викликати фальшиві спрацьовування або приховати реальні загрози. Для запобігання цьому необхідно інтегрувати додаткові механізми захисту, такі як шифрування даних, багатфакторна автентифікація для доступу до системи та моніторинг її стану.

Окремою проблемою є забезпечення прозорості роботи штучного інтелекту. У багатьох випадках алгоритми, особливо ті, що базуються на глибокому навчанні, працюють як "чорні ящики", ускладнюючи розуміння їхніх рішень, що може викликати складнощі у випадках, коли потрібно обґрунтувати конкретне рішення системи або провести аудит її роботи. Забезпечення прозорості та пояснюваності ШІ вимагає використання підходів, які дозволяють аналізувати та пояснювати процес прийняття рішень, наприклад, методів Explainable AI (XAI). Прозорість роботи ШІ є важливим фактором для підтримки довіри з боку фахівців і забезпечення можливості взаємодії з системою у критичних ситуаціях.

У підсумку проведеного аналізу та опису системи впровадження штучного інтелекту в SIEM-системи було окреслено ключові аспекти, які відображають як переваги, так і виклики, пов'язані з інтеграцією таких рішень у сфері кібербезпеки.

Використовуючи всю надану інформацію було сформовано та зображено алгоритм впровадження ШІ в наявну систему SIEM, який продемонстровано на рисунку 2.6.

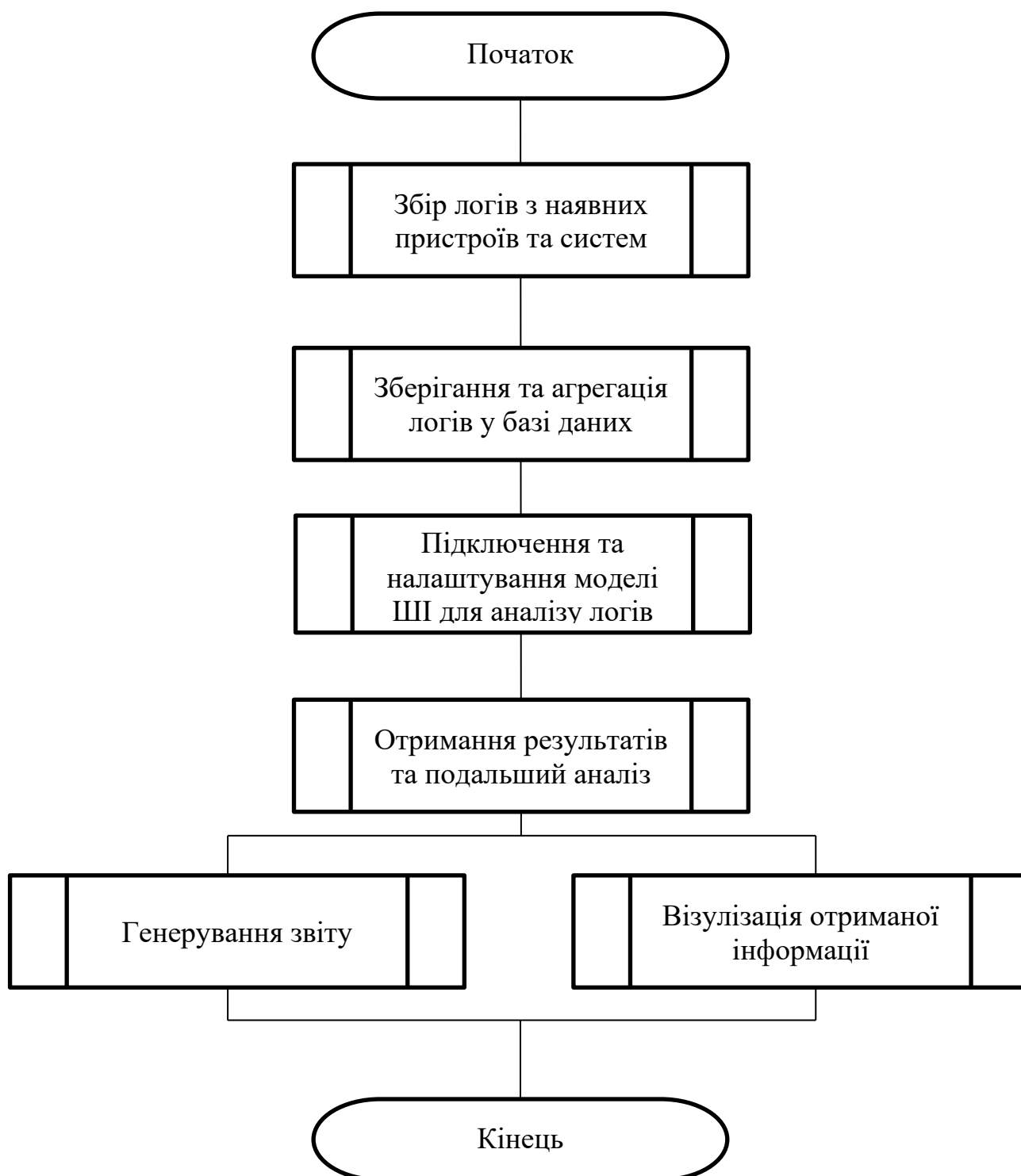


Рисунок 2.6 – Схема алгоритму впровадження штучного інтелекту в систему SIEM

Інтеграція штучного інтелекту в існуючу архітектуру SIEM-системи планується здійснитись через кілька етапів. Оскільки в межах виконання роботи вже був розроблений кіберполігон із налаштованою SIEM-системою, наступним

кроком стало впровадження ШІ для покращення аналізу логів та забезпечення більш швидкого реагування на кіберзагрози.

Першим етапом потрібно підготувати інфраструктуру для інтеграції ШІ, що включає налаштування серверів і забезпечення відповідного програмного середовища для обробки великих обсягів даних, зокрема логів, що генеруються різними джерелами в межах SIEM, що дозволило налаштувати систему для роботи в реальному часі.

Наступним етапом йде вибір моделі ШІ для аналізу логів. Враховуючи важливість роботи з великими масивами текстових даних, було прийнято рішення використовувати метод GPT (Generative Pre-trained Transformer), так як цей метод, завдяки своїм можливостям працювати з текстовими даними, є особливо ефективним для обробки неструктурованих логів і здатний розпізнавати патерни в даних та аналізувати аномалії в поведінці системи, що дозволяє виявляти загрози ще до їх фактичної реалізації.

Для інтеграції GPT у SIEM-систему буде створено додатковий скрипт, який передає логи на обробку в модель ШІ та отримує результати у вигляді тексту та графіків, що дозволило здійснювати обробку логів у реальному часі без необхідності попереднього їх структурного аналізу. Після обробки даних GPT генерує попередження або рекомендації щодо потенційних загроз, які потім можуть передаватись відділу інформаційної безпеки для подальшого реагування.

Завдяки такій інтеграції, SIEM-система отримає значне підвищення точності та швидкості виявлення аномалій, зменшивши кількість хибних спрацьовувань і автоматизувавши процеси обробки даних. Це дозволяє фахівцям з кібербезпеки зосередитися на найважливіших інцидентах та ефективно реагувати на загрози в режимі реального часу.

2.5 Вибір GPT для вирішення задач обробки логів

На ринку штучного інтелекту існує кілька розробок моделей на основі технології GPT, кожна з яких має свої особливості, переваги й обмеження. Найбільш відомими розробниками є OpenAI, Google та інші компанії, що

пропонують трансформерні архітектури для роботи з текстовими даними. Розглянемо основні переваги та недоліки наявних моделей:

1. OpenAI GPT (GPT-4, GPT-4o Mini, GPT-3.5) [17]

Переваги:

- гнучкість та універсальність;
- інтуїтивний API;
- широка підтримка та документація;
- оптимальні витрати.

Недоліки:

- залежність від інтернету;
- обмеження продуктивності в локальних середовищах.

Моделі GPT від OpenAI здатні працювати з широким спектром задач, включаючи генерацію тексту, аналіз логів, обробку природної мови та контекстну класифікацію. Також, OpenAI пропонує легкий у використанні API, який дозволяє швидко інтегрувати GPT у будь-яке середовище. Розробники отримують доступ до розгорнутої документації, прикладів використання та активної спільноти. Різна кількість наявних версій пропонує конкурентоспроможну тарифну модель із низькою вартістю обробки вхідних токенів, що робить її придатною для середніх і великих організацій. GPT добре справляється з аналізом складних текстових даних, таких як журнали подій, звіти або записи взаємодій.

Щодо недоліків, для використання моделей потрібне стабільне підключення до сервісів OpenAI. GPT працює лише через API, що обмежує можливість роботи без підключення до мережі.

2. Google Gemini AI (Pathways Language Model) [18]

Переваги:

- підтримка великих обсягів даних;
- мультизадачність;
- інтеграція з екосистемою Google.

Недоліки:

- висока вартість;
- складність інтеграції.

PaLM розроблено для аналізу масштабних наборів даних, що дозволяє використовувати її у великих корпоративних середовищах. Щодо моделі, то вона здатна виконувати різноманітні завдання одночасно завдяки багатопоточності. PaLM легко інтегрується в інфраструктуру Google Cloud, що спрощує налаштування для клієнтів, які вже використовують сервіси Google. Попри такі переваги, така модель має ряд недоліків, використання PaLM значно дорожче порівняно з GPT від OpenAI та для налаштування PaLM потрібен високий рівень технічної експертизи, а документація менш розгорнута.

3. Claude від Anthropic [19]

Переваги:

- етичний підхід;
- гнучкість у діалогових завданнях.

Недоліки:

- обмеженість функціоналу;
- нижча продуктивність.

Claude приділяє особливу увагу безпеці й етиці використання, що робить її привабливою для організацій із високими вимогами до захисту даних. Модель добре справляється із завданнями, що потребують підтримки тривалих діалогів або аналізу текстів. Claude більше орієнтована на спілкування, ніж на технічний аналіз даних, що обмежує її застосування для задач SIEM. У порівнянні з OpenAI GPT, модель має менший контекст і нижчу точність у складних завданнях.

4. Mistral та Falcon [20]

Переваги:

- відкритий код;
- оптимізація продуктивності.

Недоліки:

- потреба в значній технічній підготовці;
- вузька спеціалізація.

Моделі доступні для локального розгортання, що забезпечує максимальну незалежність і контроль над даними. Falcon і Mistral побудовані з акцентом на ефективність використання ресурсів. Для налаштування й інтеграції моделей потрібен високий рівень знань у галузі обчислювальної техніки та ІІІ. Моделі менш адаптовані для широкого спектру задач, зокрема текстового аналізу в SIEM.

Після проведеного аналізу на сучасному ринку, перевагу та подальше використання отримують моделі від компанії OpenAI.

Серед доступних моделей OpenAI є кілька варіантів GPT, кожна з яких має свої характеристики, переваги та обмеження. Вибір моделі залежить від вимог до обробки логів, обсягу даних і специфіки завдань системи SIEM. Таблиця з порівнянням моделей представлена нижче під номером 2.1 [17].

Таблиця 2.1 – Порівняння наявних моделей від OpenAI

Критерій	GPT-4o	GPT-4o Mini	GPT-3.5 Turbo-0125
Контекст (токенів)	128K	128K	~16K
Точність аналізу	Дуже висока	Висока	Середня
Обробка складних задач	Відмінна	Добра	Обмежена
Вартість / 1M input токенів	\$2.50	\$0.150	\$0.50
Вартість / 1M output токенів	\$1.25	\$0.075	\$1.50
Застосування для SIEM	Великі організації	Невеликі та середні	Малий бюджет

Для інтеграції в систему SIEM оптимальним вибором є модель GPT-4o Mini, яка поєднує високу продуктивність із доступною вартістю, забезпечуючи

ефективний аналіз великих обсягів логів за мінімальних витрат. Завдяки підтримці широкого контексту до 128 тисяч токенів, модель здатна обробляти великі журнали подій, виконуючи складні запити та аналізуючи багаторівневі дані. Економічна ефективність GPT-4o Mini особливо помітна, оскільки вона є значно дешевшою за повну версію GPT-4o, зберігаючи при цьому функціональність, достатню для завдань у сфері обробки логів. Додатковою перевагою є мультимодальність, що дозволяє працювати з текстовими й візуальними даними, що підвищує універсальність моделі при роботі з різноформатними джерелами інформації.

В результаті для інтеграції в SIEM систему обрано модель GPT-4o Mini як найкращий баланс між вартістю, точністю та функціональністю, що забезпечить ефективний аналіз логів, може підтримувати складні запити та знижує витрати на обробку великих обсягів даних.

2.6 Висновок до розділу

Отже, в ході опрацювання даного розділу було з'ясовано, що інтеграція штучного інтелекту у традиційну архітектуру SIEM-систем надає значні переваги для підвищення рівня кібербезпеки сучасних організацій. Завдяки використанню алгоритмів машинного навчання та прогнозувальних моделей, SIEM-системи здатні автоматично виявляти аномальні дії, попереджати про потенційні загрози ще до їхньої реалізації та зменшувати кількість хибних спрацьовувань, що знижує навантаження на команди безпеки. Інтеграція ШІ також забезпечує можливість отримання й обробки логів із додаткових зовнішніх джерел, що підвищує обізнаність щодо поточних кіберзагроз та швидкість реагування на інциденти.

Водночас, було виявлено певні виклики, які можуть виникнути під час впровадження ШІ, зокрема складність налаштування та потребу у великих обсягах якісних даних для навчання моделей, а також забезпечення прозорості та захисту самих алгоритмів ШІ від атак. Попри ці виклики, використання штучного інтелекту в SIEM-системах показує високий потенціал у підвищенні

ефективності кіберзахисту, роблячи його більш адаптивним, масштабованим і здатним до проактивного реагування на загрози.

Впродовж опрацювання цього розділу було описано метод впровадження штучного інтелекту в систему SIEM для аналізу логів. Розглянуто етапи інтеграції, починаючи від підготовки інфраструктури та налаштування серверів, до вибору та інтеграції методу GPT для ефективного оброблення неструктурованих даних.

Також, було проведено аналіз наявних компаній на ринку, які надають доступ до моделей GPT, розставлено переваги та недоліки, після чого було обрано основну компанію OpenAI та проведено аналіз наявних моделей штучного інтелекту.

У процесі вибору моделі було зважено на ефективність роботи з великими обсягами логів, економічну вигоду та вимоги до продуктивності. Враховуючи ці фактори, було вирішено зупинитися на GPT-4o Mini. Ця модель поєднує високу ефективність обробки текстових даних, підтримує 128K токенів для контексту, що дозволяє працювати з великими наборами логів і складними запитамі, а також має значно нижчу вартість порівняно з GPT-4, що робить її економічно доцільною для широкого застосування.

3 ЕКСПЕРЕМЕНТАЛЬНІ ДОСЛІДЖЕННЯ

3.1 Опис програмного середовища та інструментів

У даному підрозділі розглянуто програмне середовище та інструменти, використані для розробки та інтеграції штучного інтелекту в SIEM-систему з метою автоматизованого аналізу логів і мережевого трафіку. Особливістю цієї роботи є використання кіберполігону, створеного в рамках бакалаврської роботи, що забезпечує тестування та оцінку ефективності рішень у контрольованому середовищі [21]. Основні інструменти й технології було обрано з урахуванням вимог до масштабованості, продуктивності та зручності інтеграції.

Кіберполігон, розроблений у рамках попередніх досліджень, забезпечує реалістичне середовище для симуляції мережевого трафіку та генерації логів з різних систем. Створене середовище дозволяє налаштувати та проводити атаки, створюючи аномалії для подальшого аналізу. Використання такого тестового середовища дозволяє перевірити роботу моделей на основі штучного інтелекту у близьких до реальних умов, що забезпечує об'єктивність результатів [21]. Топологія мережі кіберполігону зображена на рисунку 3.1.

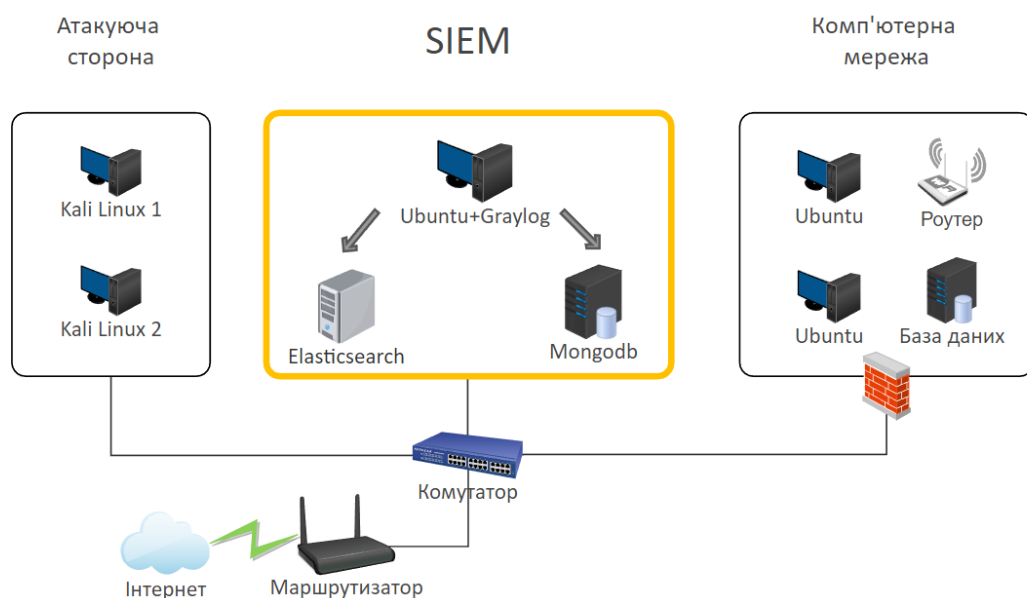


Рисунок 3.1 – Топологія мережі кіберполігону

Основна робота з інтеграції штучного інтелекту в SIEM-систему зосереджується на забезпеченні автоматизації збору, обробки та аналізу логів для підвищення ефективності виявлення загроз.

Для ефективної реалізації прототипу системи аналізу логів у рамках кіберполігону було обрано програмне середовище, побудоване на мові програмування Python. Вибір зумовлений низкою переваг, які забезпечують відповідність Python до вимог завдання. Мова Python вирізняється простотою у написанні коду, широкою екосистемою бібліотек для обробки даних і можливістю інтеграції з інструментами штучного інтелекту, такими як GPT-моделі. Крім того, Python підтримує багатопотоковість, роботу з великими обсягами даних і має високу популярність у сфері кібербезпеки, що забезпечує доступ до багатьох готових рішень і спільноти розробників.

Програмне середовище, яке буде створено в рамках цієї роботи, націлене на автоматизацію збору, попередньої обробки та аналізу логів, що є важливою частиною процесу кібербезпеки. Відповідно до цього, інтеграція Python із компонентами кіберполігону дозволить забезпечити зручне й ефективне збирання логів із різноманітних джерел у реальному часі, для прикладу: сервери, мережеве обладнання, додатки, а також інші компоненти інфраструктури, які генерують критичні дані для моніторингу та виявлення аномалій.

Після цього дані передаються до модулю штучного інтелекту, інтегрованого в систему SIEM для аналізу. У цьому контексті GPT-модель грає важливу роль: вона отримує структуровані запити, що містять ключову інформацію про події, зібрані з різних джерел. Python забезпечує легкість формування цих запитів та передачу їх через API, що дозволяє модулю GPT отримувати точні дані для подальшого аналізу. Завдяки цьому взаємодія між різними компонентами системи стає не лише ефективною, а й гнучкою, адже Python є універсальним інструментом для інтеграції з різними типами систем та протоколами.

Для збору та зберігання даних в рамках кіберполігону буде використано наступні інструменти:

- Graylog – платформа для управління логами, яка збирає та індексує логи з різних джерел у кіберполігоні. Завдяки Graylog лог-файли централізовано зберігаються для подальшого аналізу.
- MongoDB – NoSQL база даних, яка зберігає неструктуровані дані та дозволяє ефективно обробляти великі обсяги логів. MongoDB використовується для зберігання даних, які надходять від кіберполігону, а також для проміжного зберігання результатів аналізу.
- Elasticsearch – пошукова система, яка забезпечує швидкий доступ до даних для їх подальшого аналізу й пошуку. Elasticsearch інтегрується з Python для зручного та швидкого оброблення логів у рамках кіберполігону.

Проект використовує Python та інструменти обробки запитів до штучного інтелекту для обробки та аналізу логів, зібраних у кіберполігоні [22]:

- Plotly та Pandas – бібліотеки для аналізу даних, що допомагають у створенні таблиць та графіків аналізу логів і трафіку.
- OpenAI API – додаткова бібліотека для інтеграції алгоритмів штучного інтелекту, яка автоматизує аналіз логів і допомагає виявляти аномалії й потенційні загрози. API дозволяє використовувати сучасні моделі глибокого навчання для обробки логів та мережевого трафіку з кіберполігону.

Для візуалізації результатів буде використовуватися мова програмування Python, що дозволяє створювати кастомізовані дашборди. Інтегрована система Python збирає логи з кіберполігону, проводить попередню обробку, аналізує їх за допомогою моделей ШІ, а результати подає через дашборди для візуалізації, що дозволяє миттєво оцінити ситуацію та отримати інформацію про стан безпеки.

Розроблені на Python дашборди дозволяють налаштовувати та відображати результати аналізу у вигляді графіків, таблиць та індикаторів, які легко інтерпретувати. На відміну від стандартних рішень, кастомізована візуалізація забезпечує гнучкість, швидке оновлення даних та можливість інтеграції результатів ШІ-аналізу у реальному часі. Графічне представлення всього шляху від збору логів до його аналізу за допомогою ШІ представлено на рисунку 3.2.

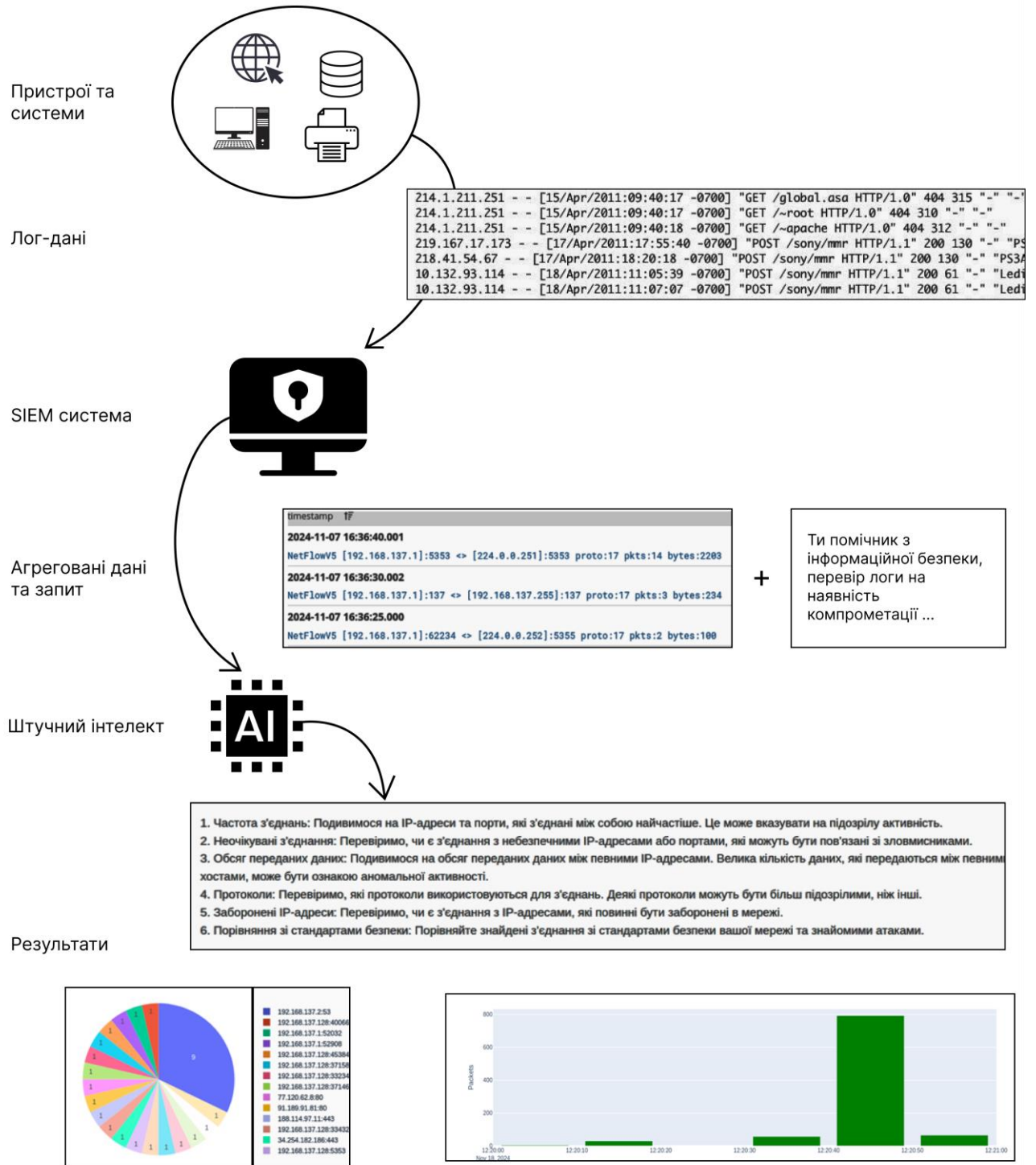


Рисунок 3.2 – Графічне представлення шляху, який проходять логи від збору до формування звіту

Отже, в результаті було проаналізовано та забезпечено реалістичні умови для тестування SIEM-системи з інтегрованим ШІ, завдяки використанню кіберполігону, що було створеного під час попередніх досліджень. Поєднання Python з MongoDB, Elasticsearch, Graylog і OpenAI API дозволяє побудувати

ефективну, гнучку систему аналізу логів і трафіку. Візуалізація результатів через дашборди на Python надає користувачам зручний та інформативний інтерфейс, що підвищує ефективність роботи з даними й пришвидшує процес реагування на загрози. Перейдемо до розробки прототипу системи.

3.2 Розробка прототипу системи

У цьому розділі описано процес створення прототипу SIEM-системи, доповненої функціями штучного інтелекту для аналізу логів та виявлення кіберзагроз. Розробка прототипу включала визначення архітектури системи, налаштування середовища для роботи з великими обсягами даних, реалізацію алгоритмів на основі моделі GPT, інтеграцію модулів для обробки різних типів логів та тестування системи.

Перший етап розробки прототипу включав розробку архітектури системи. Враховуючи вимоги до продуктивності, обсягу даних і швидкості обробки логів, було вирішено розмістити SIEM та модуль штучного інтелекту на серверній інфраструктурі з високою обчислювальною потужністю, що в подальшому дозволить забезпечити ефективну обробку даних в режимі реального часу та підтримувати стабільну роботу системи під навантаженням.

Архітектура прототипу включає наступні основні компоненти:

- 1) Модуль збору логів – для отримання даних з різних джерел (мережевих пристроїв, серверів, баз даних, додатків тощо).
- 2) Система зберігання даних – база даних для логів (MongoDB), що дозволяє зберігати дані у великому обсязі.
- 3) Модуль обробки та аналізу логів – на базі GPT, інтегрований через API, що аналізує зібрані дані та виявляє аномалії.
- 4) Інтерфейс користувача – для візуалізації результатів обробки, аналітики та звітів про загрози.

Налаштування прототипу включало створення середовища для роботи з даними у форматі логів, так як всі вони спочатку збираються, а потім індексуються у Graylog. Також є можливість побачити, в якому вигляді вони

зберігаються, що представлено на рисунку 3.3. Тут показано, як логи виглядають всередині бази даних Graylog та як виглядають в агрегованому вигляді на інтерактивній панелі.

```
{
  "_index": "graylog_0",
  "_type": "doc",
  "_id": "7a745440-077b-11ee-b064-000c2990507e",
  "_score": null,
  "_source": {
    "nf_dst": "239.255.255.250:1900",
    "nf_stop": "2023-06-10T10:41:11.891Z",
    "nf_version": 5,
    "gl2_remote_ip": "127.0.0.1",
    "gl2_remote_port": 49033,
    "nf_pkts": 4,
    "source": "127.0.0.1",
    "gl2_source_input": "646bae9bcf82341db70912e1",
    "nf_tcp_flags": 0,
    "nf_dst_port": 1900,
    "nf_bytes": 792,
    "nf_src_as": 0,
    "nf_proto": 17,
    "nf_dst_address": "239.255.255.250",
    "nf_src_port": 61292,
    "gl2_source_node": "72db8cfe-a25e-41d8-ae67-f703c2ccfe64",
    "timestamp": "2023-06-10 10:42:20.000",
    "nf_src": "192.168.137.1:61292",
    "nf_dst_mask": 0,
    "gl2_accounted_message_size": 543,
    "nf_flow_packet_id": 22859,
    "streams": [
      "00000000000000000000000000000001"
    ],
    "gl2_message_id": "01H2JECQW5M7BEG97401PCDPJG",
    "message": "NetFlowV5 [192.168.137.1]:61292 <> [239.255.255.250]:1900 proto:17 pkts:4 bytes:792",
    "nf_dst_as": 0,
    "nf_tos": 0,
    "nf_src_address": "192.168.137.1",
    "nf_proto_name": "UDP",
    "nf_src_mask": 0,
    "nf_snmp_input": 0,
    "nf_snmp_output": 0,
    "nf_start": "2023-06-10T10:41:08.862Z"
  },
  "sort": [
    1686393740000
  ]
}
```

a)

application_name	kernel
facility	kernel
facility_num	0
level	4
message	[3276.192018] workqueue: blk_mq_run_work_fn hogged CPU for >10000us 256 times, consider switching to WQ_UNBOUND
source	user-VirtualBox
timestamp	2024-11-21T15:01:54.073Z

б)

Рисунок 3.3 – Приклад вигляду логів (мережевий (а), системний (б))

Наступним кроком, потрібно правильно обробити логи. Для цього було використано промпти (запити), які виглядають як комбінація двох компонентів:

role system і role user, які задають контекст функціонування ШІ та характер його взаємодії із запитом. Параметр role system визначає основний контекст або "особистість" ШІ, наприклад, вказуючи, що ШІ виступає як експерт із кібербезпеки, здатний аналізувати логи для виявлення та запобігання кіберзагрозам. Таким чином є можливість поекспериментувати, паралельно задаючи різні запити, що можуть бути в нагоді для покращення роботи та отриманих результатів. На рисунку 3.4 представлено приклад запиту, який демонструє комбінацію обох компонентів.

```
messages=[
  {"role": "system", "content": "Ви помічник, який має досвід аналізу журналів мережевого трафіку, щоб виявити можливі кіберзагрози."},
  {"role": "user", "content": f"Проаналізуйте наведені нижче журнали мережевого трафіку на наявність будь-яких аномалій або підозрілих дій:\n{joined_logs}"}]
```

Рисунок 3.4 – Приклад вигляду сформованого запиту

Після завершення етапу збору, отримані журнали передаються на сервер OpenAI для подальшої обробки. Повний запит до системи ШІ містить низку параметрів, які дозволяють налаштувати процес обробки інформації відповідно до конкретних потреб. Наприклад, є можливість вказати модель, яка буде використовуватись для виконання обчислень і генерації відповідей, що визначає рівень деталізації, специфіку та швидкість роботи. Окрім цього, можна задати максимальну кількість токенів, що обмежує обсяг тексту, контролюючи тривалість і змістовність результату. Параметр температури, регулює ступінь випадковості у відповідях моделі: вищі значення сприяють більш креативним та варіативним відповідям, тоді як нижчі забезпечують більш передбачувані та строгі результати. Налаштування цих параметрів відбувається в коді програми, приклад повного запиту зображений на рисунку 3.5.

```

try:
    response = client.chat.completions.create(
        model="gpt-3.5-turbo-0125",
        max_tokens=1000,
        temperature=0.5,
        messages=[
            {"role": "system", "content": "Ви помічник, який має досвід аналізу журналів мережевого трафіку, щоб виявити можливі кіберзагрози."},
            {"role": "user", "content": f"Проаналізуйте наведені нижче журнали мережевого трафіку на наявність будь-яких аномалій або підозрілих дій:\n{joined_logs}"
        ]
    )
    analysis = response.choices[0].message.content.replace("\n", "<br>")
    return analysis
except Exception as e:
    return f"Помилка аналізу журналів мережі: {e}"

```

Рисунок 3.5 – Приклад повного запиту з параметрами

У результаті формується підсумковий звіт, який наведено на рисунку 3.6.

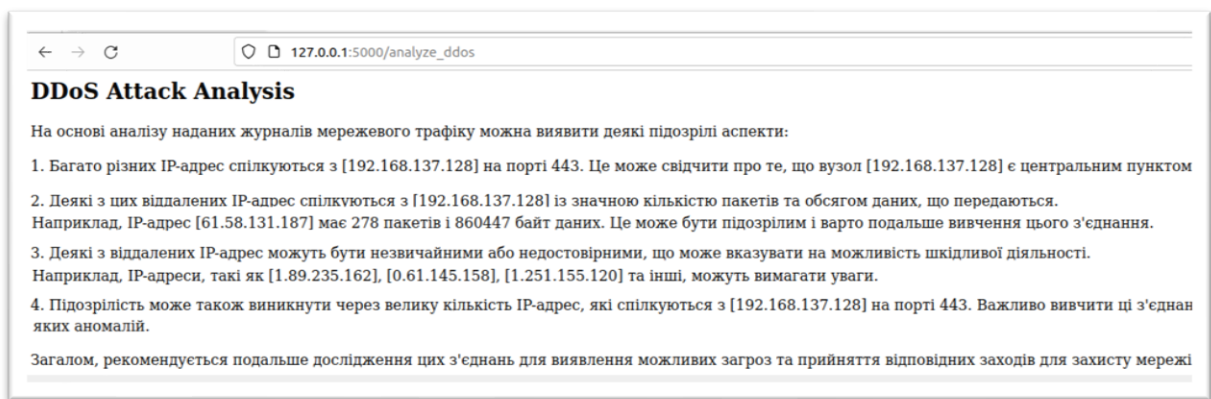


Рисунок 3.6 – Приклад отриманого результату

Основна роль GPT полягала у виявленні аномалій та патернів, що свідчать про потенційні кіберзагрози. GPT-алгоритми були налаштовані для роботи з текстовими логами, щоб зчитувати дані, ідентифікувати аномалії та формувати звіти про можливі інциденти.

Для візуалізації отриманих результатів їх було представлено у вигляді HTML-сторінки з графіками і таблицями, що дозволяють ефективно проаналізувати дані та зробити відповідні висновки, які зображено на рисунку 3.4. Для створення графіків і таблиць використовувався Python разом із бібліотекою Plotly та Pandas, що дозволило створити деталізовані, динамічні графіки з можливістю масштабування, вибору окремих елементів та інших

інтерактивних функцій, що значно покращило наочність і зручність аналізу результатів.

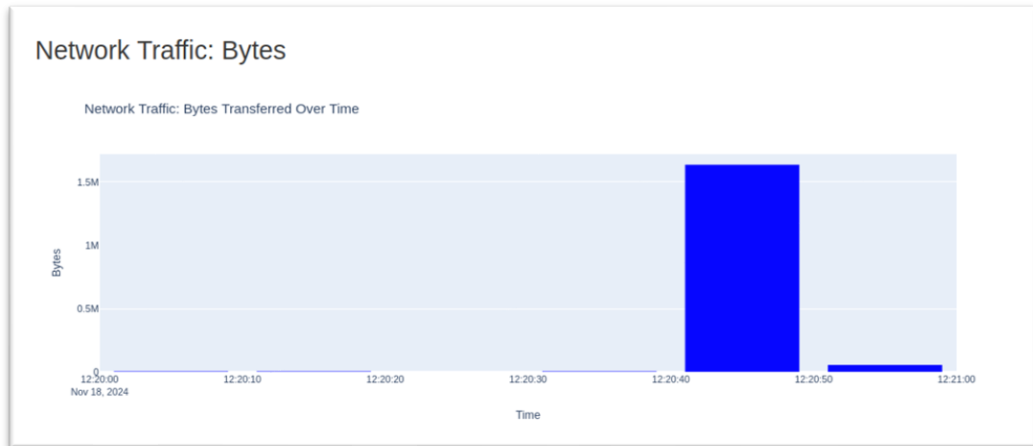


Рисунок 3.7 – Вигляд графіку кількості надісланих байтів інформації на веб-сторінці з результатами

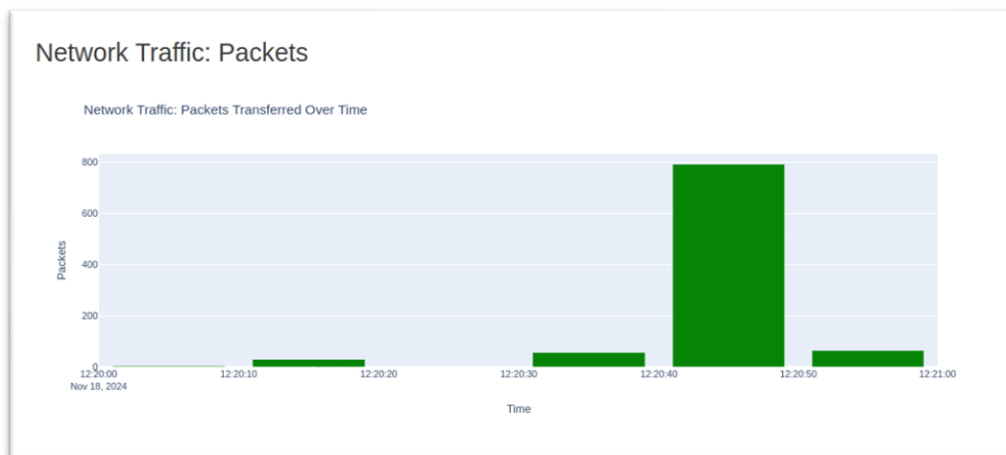


Рисунок 3.8 – Вигляд графіку кількості надісланих пакетів даних

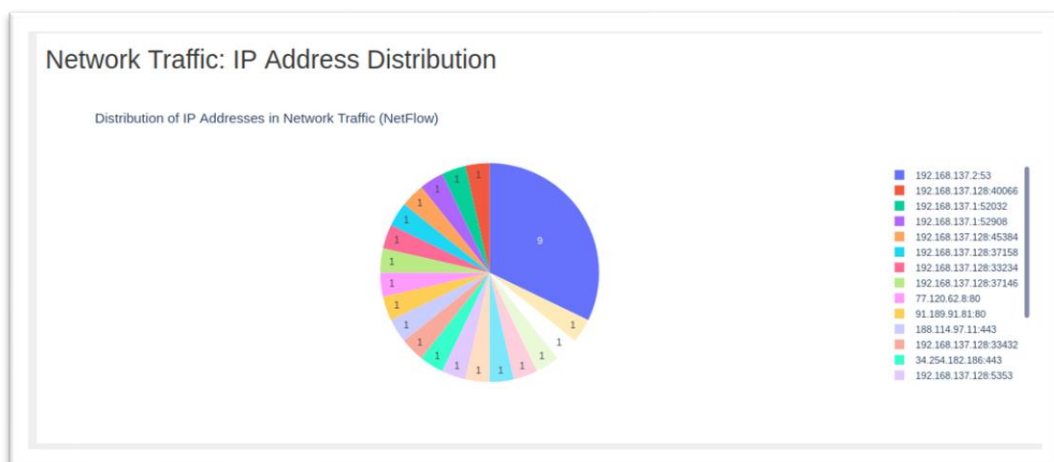


Рисунок 3.9 – Вигляд кругової діаграми з IP-адресами та їх кількостями, які фігурують в логах

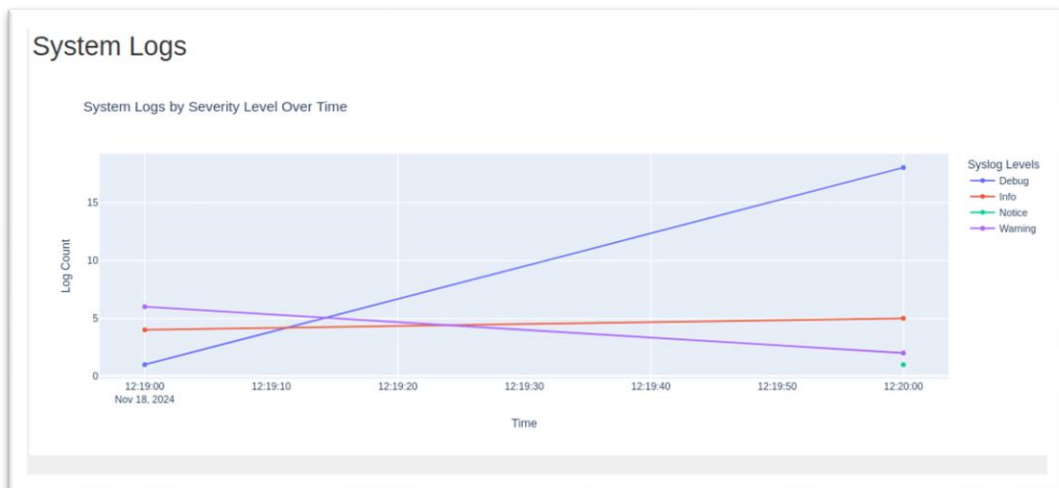


Рисунок 3.10 – Вигляд графіку кількості системних логів та їх рівні

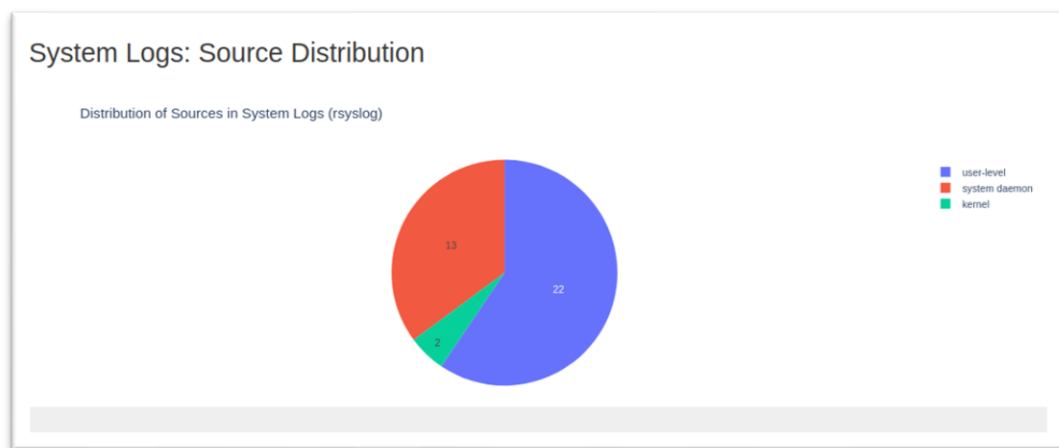


Рисунок 3.11 – Вигляд кругової діаграми з кількістю та джерелами системних логів

Як можна побачити з рисунків вище (див. рис. 3.7-3.11), таке представлення значано спрощує вигляд для подальшого аналізу даних. Такі графіки показують наочне відображення інформації, сприяють виявленню ключових тенденцій та взаємозв'язків, а також завдяки інтерактивності дозволяють глибше досліджувати дані.

Також, на сторінці наявні рекомендаційні заходи, які пропонує ШІ, також зображені на рисунку 3.12.

AI Analysis: Network Traffic

На основі аналізу наданих журналів мережевого трафіку можна виявити деякі підозрілі аспекти:

1. Багато різних IP-адрес спілкуються з [192.168.137.128] на порті 443. Це може свідчити про те, що вузол (192.168.137.128) є центральним пунктом для багатьох з'єднань.
2. Деякі з цих віддалених IP-адрес спілкуються з 192.168.137.128 із значною кількістю пакетів та обсягом даних, що передаються. Наприклад, IP-адрес (61.58.131.187) має 278 пакетів і 860,447 байт даних. Це може бути підозрілим, і варто подальше вивчення цього з'єднання.
3. Деякі з віддалених IP-адрес можуть бути незвичайними або недостовірними, що може вказувати на можливість шкідливої діяльності. Наприклад, IP-адреси, такі як (11.89.235.162), (0.61.145.158), [1.251.155.120] та інші, можуть вимагати уваги.
4. Підозрілість може також виникнути через велику кількість IP-адрес, які спілкуються з [192.168.137.128] на порті 443. Важливо вивчити ці з'єднання для виявлення будь-яких аномалій.

Загалом, рекомендується подальше дослідження цих з'єднань для виявлення можливих загроз та прийняття відповідних заходів для захисту мережі.

a)

AI Analysis: System Logs

За результатами аналізу системних журналів можна виявити наступні проблеми та незвичайну активність:

1. Помилка управління монтажем простору імен, що призвела до відмови в доступі через permission denied:
update.go:85: cannot change mount namespace according to change mount (/var/lib/snapd/hostfs/usr/local/share/doc /usr/local/share/doc none bind,ro 0 0): cannot open directory "/usr/local/share": permission denied
2. Помилки оновлення джерел пакетів, що можуть вказувати на проблеми з мережею або недоступністю деяких репозиторіїв:
Failed to fetch https://packages.graylog2.org/repo/debian/dists/stable/InRelease
Failed to fetch https://repo.mongodb.org/apt/ubuntu/dists/focal/mongodb-org/4.4/InRelease
Failed to fetch https://repo.mongodb.org/apt/ubuntu/dists/jammy/mongodb-org/6.0/InRelease
Some index files failed to download. They have been ignored, or old ones used instead.
3. Повідомлення про роботу з чергами, що може вказувати на можливі проблеми з продуктивністю:
[94.903626] workqueue: blk_mq_run_work_fn hogged CPU for >10000us 8 times, consider switching to WQ_UNBOUND
4. Повідомлення про невдачу спробу оновлення іконки для livepatch:
unable to update icon for livepatch
5. Повідомлення про відсутність оновлень для деяких пакетів:
storehelpers.go:923: cannot refresh: snap has no updates available: "bare", "core18", "gnome-3-38-2004", "gnome-42-2204", "gtk-common-themes", "shellcheck", "snapd", "snapd-desktop-integration"

Загалом, в системних журналах виявлено декілька проблем, які можуть вплинути на безпеку та продуктивність системи. Рекомендується подальше вивчення та вирішення цих проблем для забезпечення стабільної та безпечної роботи системи.

б)

Рисунок 3.12 – Рекомендації від ШІ на веб-сторінці з результатами

Отже, етап створення прототипу SIEM-системи, доповненої функціями штучного інтелекту для аналізу логів та виявлення кіберзагроз, пройшов успішним та в результаті отримано прототип системи.

3.3 Аналіз результатів та оцінка ефективності

Наступним етапом стала перевірка коректності роботи прототипу SIEM-системи з інтегрованим штучним інтелектом у середовищі кіберполігону.

Основну увагу було зосереджено на трьох аспектах: достовірність виявлення загроз, швидкості обробки логів та точності класифікації виявлених інцидентів.

Для перевірки функціональності та ефективності розробленої системи буде використовуватися попередньо згенеровані набори даних (datasets), які імітують кілька атак, зокрема перший датасет імітує атаку типу DDoS. Всередині dataset містить приклади мережевого трафіку з високим навантаженням на сервери, включаючи різні патерни поведінки, що характерні для розподілених атак на відмову в обслуговуванні. Наступним кроком будуть перевірені дані, які в результаті дозволять оцінити здатність системи виявляти аномалії всередині логів, ефективно аналізувати великий обсяг запитів і своєчасно ідентифікувати ознаки атак.

Приклад кількох логів із датасету, що використовувався для проведення досліджень, представлено нижче на рисунку 3.13, а весь датасет можна знайти за посиланням у Github [23].

```
{
  "nf_dst": "192.168.137.128:443",
  "nf_stop": "2024-10-08T14:52:07.206Z",
  "nf_version": 5,
  "g12_remote_ip": "194.139.139.226",
  "g12_remote_port": 22074,
  "nf_pkts": 199,
  "source": "114.86.27.213",
  "g12_source_input": "646bae9bcf82341db70912e1",
  "nf_tcp_flags": 24,
  "nf_dst_port": 443,
  "nf_bytes": 630419,
  "nf_src_as": 0,
  "nf_proto": 6,
  "nf_dst_address": "192.168.137.128",
  "nf_src_port": 50488,
  "g12_source_node": "72db8cfe-a25e-41d8-ae67-f703c2ccfe64",
  "timestamp": "2024-10-08T14:52:07.206Z",
  "nf_src": "74.32.200.76:19844",
  "nf_dst_mask": 0,
  "g12_accounted_message_size": 527,
  "nf_flow_packet_id": 9660,
  "streams": [
    {
      "g12_message_id": "01J8S00AYYE5W2XRK1T5NNE0001",
      "message": "NetFlowV5 [17.144.164.92]:22228 <> [192.168.137.128]:443 proto:6 pkts:676 bytes:68306",
      "nf_dst_as": 0,
      "nf_tos": 0,
      "nf_src_address": "211.207.95.140",
      "nf_proto_name": "TCP",
      "nf_src_mask": 0,
      "nf_snmp_input": 0,
      "nf_snmp_output": 0,
      "nf_start": "2024-10-08T14:52:07.206Z"
    }
  ],
  "nf_dst": "192.168.137.128:443",
  "nf_stop": "2024-10-08T14:52:09.206Z",
  "nf_version": 5,
  "g12_remote_ip": "172.18.170.173",
  "g12_remote_port": 38147,
  "nf_pkts": 820,
  "source": "227.37.10.183",
  "g12_source_input": "646bae9bcf82341db70912e1",
  "nf_tcp_flags": 27,
  "nf_dst_port": 443,
  "nf_bytes": 318948,
  "nf_src_as": 0,
  "nf_proto": 6,
  "nf_dst_address": "192.168.137.128",
  "nf_src_port": 47081,
  "g12_source_node": "72db8cfe-a25e-41d8-ae67-f703c2ccfe64",
  "timestamp": "2024-10-08T14:52:09.206Z",
  "nf_src": "211.75.53.23:64344",
  "nf_dst_mask": 0,
  "g12_accounted_message_size": 921,
  "nf_flow_packet_id": 8111,
  "streams": [
    {
      "g12_message_id": "01J8S00AYYE5W2XRK1T5NNE0002",
      "message": "NetFlowV5 [209.239.242.40]:13526 <> [192.168.137.128]:443 proto:6 pkts:915 bytes:842081",
      "nf_dst_as": 0,
      "nf_tos": 0,
      "nf_src_address": "149.128.78.172",
      "nf_proto_name": "TCP",
      "nf_src_mask": 0,
      "nf_snmp_input": 0,
      "nf_snmp_output": 0,
      "nf_start": "2024-10-08T14:52:09.206Z"
    }
  ]
}
```

Рисунок 3.13 – Вигляд частини датасету

Під час тестування було з'ясовано, що результати аналізу значною мірою залежать від якості формулювання запитів до штучного інтелекту. Прості, узагальнені запити часто давали недостатньо конкретні результати. Відповіді GPT у таких випадках обмежувалися базовими рекомендаціями, які не дозволяли глибоко проаналізувати лог-файли чи виявити конкретні аномалії. Наприклад,

під час тестування прототипу були зафіксовані випадки, коли GPT, реагуючи на узагальнені запити, пропонував типові рекомендації, такі як перевірка конфігурацій мережі або проведення додаткового аудиту, без конкретного вказання на можливі джерела проблеми. Узагальнені запити могли описати можливі загальні підходи до аналізу даних, залишаючи невирішеними питання специфічних проблем, виявлених у логах, що свідчило про необхідність внесення змін у підхід до формування запитів (рис. 3.14).

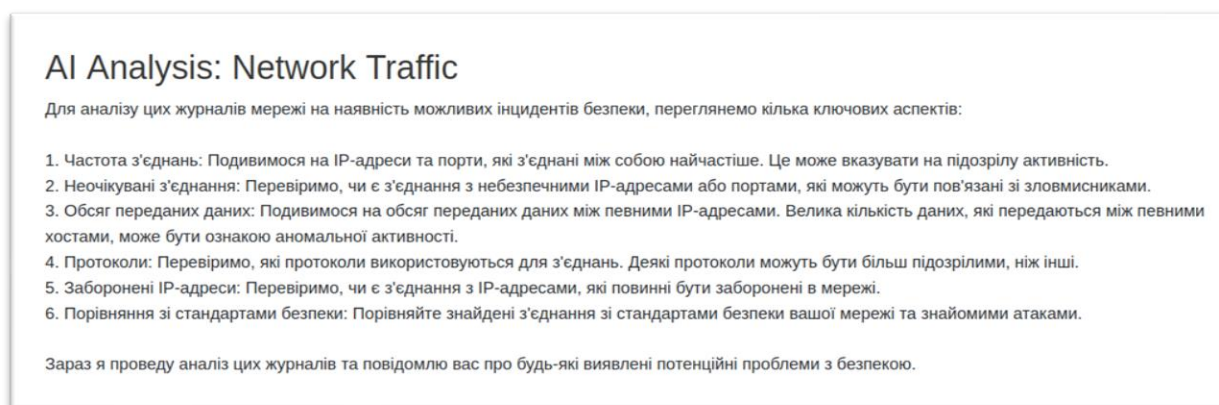


Рисунок 3.14 – Відповідь штучного інтелекту на прості запити

Використання такого підходу виявилось особливо проблематичним у сценаріях із великою кількістю даних, коли модель не могла коректно виокремити найважливіші аспекти з отриманих логів, що підкреслило необхідність більш ретельного формулювання запитів для адаптації роботи моделі до специфіки поставлених завдань.

Важливим аспектом оцінки ефективності стало порівняння роботи GPT із традиційними методами аналізу. У разі нестандартних форматів логів або роботи з великими обсягами даних GPT виявився більш гнучким і адаптивним. Традиційні підходи часто стикалися з труднощами під час обробки неструктурованих даних, тоді як GPT зміг надати корисні результати завдяки своїй здатності до самонавчання і розпізнавання патернів у даних.

Для прикладу, до вже наявних логів з DDoS-атакою, було додано список з різними форматами, приклад яких наведено нижче:

```

2024-11-21 10:15:45 - SecuritySystem - Warning: Multiple failed login attempts
detected - IP: 192.168.1.7 - Port 22 - Reason: Possible brute force attack.
2024-11-01 10:26:00, 192.168.1.7, "Login Attempt", "Rejected", "Failed"
[Router-4321] Binary Event Data: 1010111100001111 | Packet dropped due to high
network latency | 2024-11-01 10:25:00
Security database updated successfully - Latest threat definitions applied. - 2024-
11-21 15:35:05 - SecuritySystem - Information.
High CPU usage detected - 95% utilization - Consider reviewing running processes. -
2024-11-21 13:50:12 - Warning.

```

Після чого, було зроблено запит до ШІ та отримано результат, представлений на рисунку 3.15.

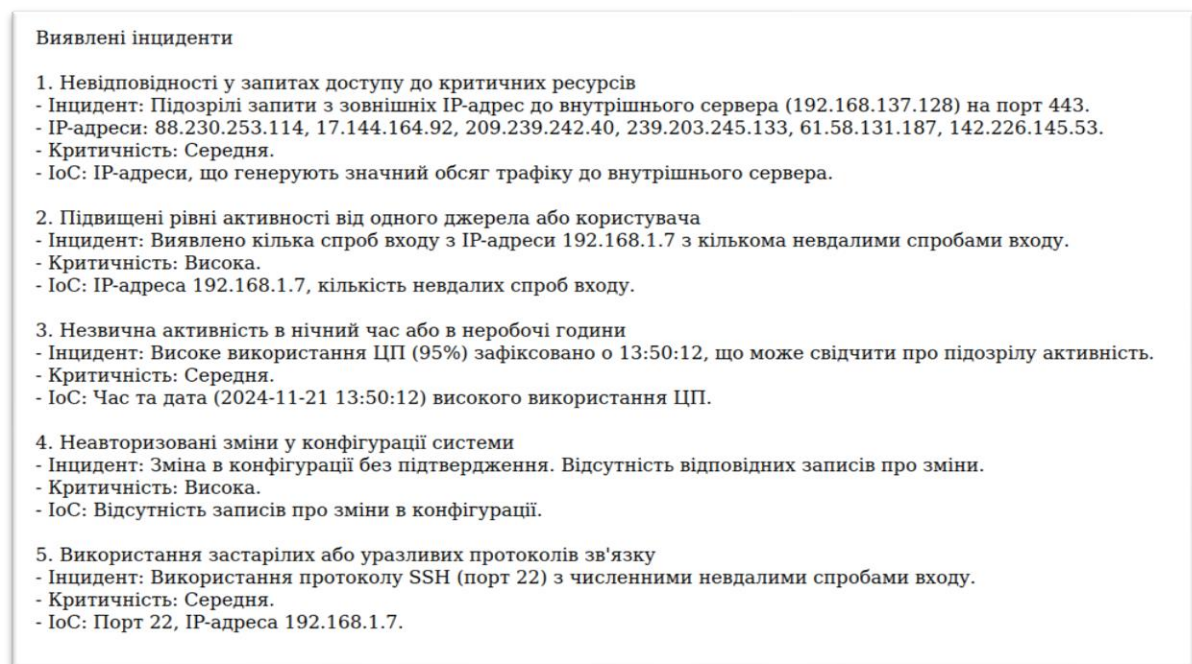


Рисунок 3.15 – Отриманий результат з додатковими логами різних форматів

Проаналізувавши результати, можна дійти висновку, що робота ШІ пройшла успішно навіть попри різні формати даних та їх структуру, що є додатковою перевагою при інтеграції ШІ в систему SIEM.

Під час дослідження було проведено експеримент, метою якого було з'ясувати, як система реагує на невідповідні або нерелевантні приклади в запитах. Для цього було використано набір даних, що містить різні типи атак, включаючи DDoS, фішингові спроби, експлойти та інші. Кожен приклад мав свою специфіку та вимагав відповідного аналізу.

1. Частота з'єднань: З журналів видно, що IP-адреса 192.168.137.128 (зазвичай це внутрішня адреса в мережі) активно взаємодіє з багатьма зовнішніми IP-адресами на порту 443 (HTTPS). Це може свідчити про те, що комп'ютер здійснює велику кількість запитів, що може бути ознакою компрометації.
2. Обсяг трафіку: Деякі з'єднання мають значний обсяг переданих байтів (наприклад, [239.203.245.133]:36991 з 270266 байт). Це може свідчити про передачу великої кількості даних, що може бути характерним для передачі конфіденційної інформації або завантаження шкідливого ПЗ.
3. Дивні або підозрілі IP-адреси: Багато з IP-адрес, з якими взаємодіє 192.168.137.128, можуть бути підозрілими. Наприклад:
 - IP-адреси, що починаються з 239, 222, 223, 240, 250, 251, 254 та інші можуть належати до класу адрес, які часто використовуються для шкідливих дій.
 - Рекомендується перевірити ці IP-адреси на предмет відомих шкідливих дій або занесення до чорних списків.
4. Патерни трафіку: Якщо трафік відбувається в незвичний час або має аномальні патерни (наприклад, раптове збільшення запитів), це може бути ознакою компрометації.
5. Аналіз доменів: Якщо у вас є доступ до DNS-записів або журналів, варто перевірити, з якими доменами намагається з'єднатися 192.168.137.128. Якщо домени виглядають підозріло або схожі на відомі фішингові сайти, це є серйозним сигналом.
6. Використання шифрування: З'єднання на порту 443 є зашифрованими, що ускладнює аналіз вмісту переданих даних. Проте, якщо на комп'ютері встановлено шкідливе ПЗ, воно може використовувати шифрування для приховування своїх дій.
7. Аномалії у поведінці користувача: Якщо ви знаєте, що користувач, пов'язаний з IP-адресою 192.168.137.128, не виконує дій, які вимагають такої кількості з'єднань або обсягу трафіку, це може бути ще одним знаком компрометації.

Рисунок 3.16 – Перевірка результатів експерименту з нерелевантністю атак

Результати показали, що передача нерелевантного прикладу (наприклад, даних про фішингову атаку в контексті тестування системи, налаштованої на виявлення DDoS) може призвести до некоректної обробки або навіть до пропуску важливих сигналів. У таких випадках алгоритми, налаштовані на певний шаблон загроз, ігнорують невідповідні дані або генерують хибнопозитивні результати.

Також було проведено детальний аналіз наявного датасету із DDoS-атаками із використанням фреймворку MITRE ATT&CK [24] та бази відомих вразливостей CVE (Common Vulnerabilities and Exposures) [25].

Аналіз датасету із DDoS-атаками за допомогою фреймворку MITRE ATT&CK та бази вразливостей CVE не дав очікуваних результатів. Зокрема, жоден із застосованих методів не надав достатньо релевантної інформації щодо виявлення або класифікації ключових аспектів атак.

MITRE ATT&CK, хоч і дозволяє структурувати дані про атаки за тактиками й техніками, не зміг запропонувати чіткої кореляції між представленими діями атакуючих і даними в наявному наборі. Інформація,

отримана через цей метод, залишилася загальною і не відповідала специфіці аналізованих DDoS-атак.

У свою чергу, використання бази CVE для ідентифікації вразливостей не дозволило розкрити деталі, які могли б прояснити механізми виконання атак. Виявлені записи у базі CVE стосувалися загальних проблем у програмному забезпеченні або обладнанні, але їхній зв'язок із подіями в датасеті залишився непідтвердженим.

Таким чином, підхід із застосуванням цих інструментів не зміг розкрити необхідної інформації про природу атак, що свідчить про потребу вдосконалення методів аналізу або підготовки більш деталізованого датасету.

Аналіз наданих журналів мережі за допомогою фреймворку MITRE ATT&CK може виявити можливі кіберзагрози та атаки. Нижче наведено основні аспекти, на які слід звернути увагу:

1. Аномальна активність

- Велика кількість з'єднань: Відзначається велика кількість з'єднань з однією і тією ж IP-адресою (192.168.137.128) на порт 443 (HTTPS). Це може вказувати на DDoS-атаку або на те, що система є мішенню для сканування вразливостей.
- Різноманітність IP-адрес: З'єднання з багатьох різних IP-адрес може свідчити про ботнет або про те, що зловмисник використовує проксі-сервери для маскування своєї активності.

2. Аналіз трафіку

- Обсяг трафіку: Деякі з'єднання мають значний обсяг переданих даних (до 1 Мб). Це може вказувати на передачу даних або на завантаження шкідливого програмного забезпечення.
- Частота пакетів: Велика кількість пакетів (до 1000) може вказувати на активну взаємодію або на спроби зловмисників отримати доступ до системи.

3. Можливі техніки з MITRE ATT&CK

- T1071.001 - Application Layer Protocol: Web Protocols: Використання HTTPS для комунікації може вказувати на спроби приховати зловмисну активність.
- T1496 - Resource Hijacking: Якщо з'єднання відбуваються з великої кількості IP-адрес, це може свідчити про використання ресурсів жертви.
- T1583 - Acquire Infrastructure: Використання різних IP-адрес може вказувати на спроби зловмисників отримати контроль над інфраструктурою.

4. Рекомендації

- Моніторинг трафіку: Рекомендується впровадити системи для моніторингу аномальної активності в мережі, щоб виявляти потенційні загрози.
- Фільтрація трафіку: Впровадження фільтрів для блокування підозрілих IP-адрес або обмеження з'єднань з невідомими джерелами.
- Аналіз з'єднань: Провести детальний аналіз з'єднань з підозрілими IP-адресами, щоб виявити можливі загрози або компрометації.

Висновок

Загалом, наявність великої кількості з'єднань з однією і тією ж IP-адресою на порт 443 може свідчити про потенційні загрози. Рекомендується вжити заходів для моніторингу та аналізу трафіку, щоб виявити та запобігти можливим атакам.

Рисунок 3.17 – Перевірка результатів з запитом на основі фреймворку MITRE ATT&CK

На основі проведеного аналізу журналів мережі, який представлено на рисунку 3.17, було виявлено кілька ключових показників аномальної активності, які можуть свідчити про потенційні загрози, такі як DDoS-атаки, використання ботнетів чи спроби маскування через проксі-сервери. Основна увага приділялася

виявленню великих обсягів трафіку на одну і ту ж IP-адресу, частоті пакетів та кількості підключень із різних IP-адрес.

Проте результати свідчать, що запропоновані методи ідентифікують лише загальні патерни підозрілої активності без чіткого визначення джерела чи механізмів атак.

Аналіз журналів мережі, поданих у вашому запиті, вказує на велику кількість з'єднань з IP-адресою 192.168.137.128, яка, ймовірно є внутрішньою адресою у вашій мережі. Основна мета цих з'єднань – порт 443, який зазвичай використовується для захищених HTTPS-з'єднань.

1. Частота з'єднань:

- У журналі зафіксовано велику кількість з'єднань з різних зовнішніх IP-адрес (більше 200), що може свідчити про активну взаємодію або потенційно шкідливу активність.

2. Обсяг переданих даних:

- Деякі з'єднання передають значні обсяги даних (наприклад, 994583 байт, 929105 байт, 918805 байт тощо). Це може вказувати на можливе використання каналу для завантаження або вивантаження великої кількості даних, що може бути тривожним сигналом.

3. Аналіз IP-адрес:

- Багато IP-адрес, які з'єднуються з 192.168.137.128, можуть бути підозрілими. Рекомендується перевірити ці IP-адреси на предмет відомих загроз або вразливостей, використовуючи базу даних CVE (Common Vulnerabilities and Exposures).

Рисунок 3.18 – Перевірка результатів з запитом на основі CVE

Те саме відбулось при запиті на основі CVE, результати якого представлені на рисунку 3.18, штучний інтелект не зміг видати релевантної інформації, а лише залишив загальні рекомендації.

Після проведеного аналізу результатів було виявлено необхідність внесення змін до структури запиту для покращення його роботи. Вони були розширені за рахунок включення контекстуальних параметрів, таких як конкретні рекомендації, рівня критичності та індикаторів компрометації [26].

В результаті, цей підхід дозволив суттєво покращити якість отримуваних відповідей. GPT-модель почала генерувати більш конкретні та деталізовані рекомендації, що включали, наприклад, конкретні аномалії в логах, алгоритми їхнього виявлення та попереджувальні заходи. На рисунку 3.6 наведено приклад оновленого та розширеного запиту, що ілюструє внесені модифікації та їх вплив на формування запиту.

```

messages=[
  {"role": "system", "content": "Ви помічник, який знає інформацію про кібербезпеку та як аналізувати журнали для запобігання кіберзагрозам атакам."},
  {
    "role": "user",
    "content": f"""
    Проаналізуйте події безпеки в наведених журналах системи. Визначте всі аномалії, підозрілі активності та потенційні загрози, які можуть свідчити про компрометацію системи. Використайте поведінковий аналіз для виявлення відхилень від нормальних дій користувачів та системних процесів.

    Конкретно зверніть увагу на:
    - Невідповідності ☐ запитах доступу до критичних ресурсів.
    - Підвищені рівні активності від одного джерела ☐ користувача.
    - Незвичну активність в нічний час ☐ в неробочі години.
    - Використання застарілих ☐ уразливих протоколів зв'язку.
    - Неавторизовані зміни ☐ конфігурації системи.

    Класифікуйте знайдені інциденти за рівнем критичності: висока, середня, низька. Вкажіть індикатори компрометації (IoC), якщо вони присутні, та надайте рекомендації щодо подальших дій для усунення потенційних загроз. Також згенеруйте звіт про ключові метрики безпеки для загального огляду системи на основі зібраних даних. Для створення звіту ☐ вас є 2000 токенів, використайте їх як варто.

    Ось журнали для аналізу:
    {joined_logs}
    """
  }
]

```

Рисунок 3.19 – Вигляд зміненого та розширеного запиту

За допомогою аналізу журналів NetFlowV5 було виявлено наступні аномалії та підозрілі активності:

- **Велика кількість підключень до одного джерела**:** Існує підозра на DDoS-атаку або спробу перехоплення даних через зловмисні дії.
- **Значна кількість переданих даних**:** Існує підозра на великі обсяги витоку даних або незаконне використання ресурсів.
- **Нехарактерна активність в нічний час**:** Зазвичай нічний час використовується для хакерських атак, оскільки може бути менше ко.
- **Підозріле використання застарілих протоколів зв'язку**:** Це може створити вразливості, які можуть бути використані для атак.
- **Велика кількість пакетів від одного джерела**:** Це може свідчити про спробу переповнення буфера або іншу форму атаки.

Класифікація знайдених інцидентів за рівнем критичності:

- **Висока критичність**:** Велика кількість підключень та переданих даних, підозріла активність в нічний час.
- **Середня критичність**:** Нехарактерна активність в нічний час, використання застарілих протоколів зв'язку.
- **Низька критичність**:** Підозріле використання застарілих протоколів зв'язку, велика кількість пакетів від одного джерела.

Індикатори компрометації (IoC):

- IP-адреса: 88.230.253.114
- IP-адреса: 17.144.164.92
- IP-адреса: 209.239.242.40
- IP-адреса: 239.203.245.133
- IP-адреса: 61.58.131.187

Рекомендації щодо подальших дій:

- **Блокування IP-адрес**:** Заблокуйте зазначені IP-адреси для запобігання подальшим атакам.
- **Моніторинг активності**:** Постійно моніторте активність мережі для виявлення подібних підозрілих активностей.
- **Оновлення протоколів зв'язку**:** Переконайтеся, що використовуються оновлені та безпечні протоколи зв'язку.
- **Зміна паролів та перевірка безпеки**:** Перегляньте паролі та забезпечте високий рівень безпеки для критичних ресурсів.

Звіт про ключові метрики безпеки:

- **Кількість підключень**:** Велика кількість підключень від одного джерела.
- **Обсяг переданих даних**:** Великий обсяг переданих даних від певних IP-адрес.
- **Час активності**:** Нехарактерна активність в нічний час.
- **Використання протоколів зв'язку**:** Використання застарілих протоколів зв'язку.

Додатковий моніторинг та аналіз можуть допомогти виявити та запобігти можливим кіберзагрозам та атакам.

Рисунок 3.20 – Отриманий результат з оновленим запитом

Результати, що зображені на рисунку 3.20, отримані за допомогою оновлених запитів і показали значне покращення, конкретно, система почала повертати не лише деталізовані рекомендації щодо виявлення проблем, але й конкретні алгоритми фільтрації даних, вказівки на можливі аномалії та сценарії реагування. Наприклад, у випадках виявлення підозрілих IP-адрес система могла пропонувати сценарії блокування трафіку або детальніше описувати зв'язки між різними подіями в логах.

Отож, фінальною задачею стала перевірка результатів на хибнопозитивність, або ж помилки першого та другого роду. Для оцінки точності роботи інтегрованої моделі GPT у системі SIEM було проведено дослідження з використанням двох різних датасетів (датасет з DDoS-атакою та датасет з різними атаками). Хибнопозитивні результати першого роду свідчать про помилкове ідентифікування звичайного трафіку як загрози, тоді як другого роду — про пропуск реальних атак. Метою дослідження було визначити, наскільки адаптована модель здатна працювати з різними типами даних, і які обмеження можуть виникати під час обробки складних логів. Розрахуємо відсотки для хибнопозитивних результатів.

Хибнопозитивний результат першого роду – випадок, в якому модель помилково ідентифікує нормальну подію як аномалію, другого роду - класифікація аномальної події як нормальної. Відсоток розраховується за формулами:

$$FP_1 = \frac{\text{Кількість хибнопозитивних 1 роду}}{\text{Загальна кількість нормальних подій}} \times 100\%$$

$$FP_2 = \frac{\text{Кількість хибнопозитивних 2 роду}}{\text{Загальна кількість нормальних подій}} \times 100\%$$

Після проведення дослідження з наведеними датасетами, було отримано такі дані:

$$1) \text{ Хибнопозитивні першого роду } FP_1 = \frac{4}{50} \times 100\% = 8\%.$$

$$\text{Хибнопозитивні другого роду } FP_2 = \frac{6}{50} \times 100\% = 12\%.$$

$$2) \text{ Хибнопозитивні першого роду } FP_1 = \frac{5}{50} \times 100\% = 10\%.$$

$$\text{Хибнопозитивні другого роду } FP_2 = \frac{7}{50} \times 100\% = 14\%.$$

Зведені результати наведено в таблиці 3.1 з представленими кількістю досліджень, відсоткової ваги результатів та додаткові примітки.

Таблиця 3.1 – Аналіз хибнопозитивних результатів для різних датасетів

Датасет	Кількість досліджень	Хибнопозитивні першого роду, %	Хибнопозитивні другого роду, %	Примітки
DDoS-атаки	50	8	12	Типові аномалії з високим обсягом трафіку.
Різні атаки	50	10	14	Комплексні атаки, включаючи фішингові спроби та SQL-ін'єкції.

Виходячи з результатів, для датасету DDoS-атак рівень хибнопозитивних результатів першого роду був помірним, що вказує на високу специфічність моделі у виявленні цього типу загроз, однак хибнопозитивні результати другого роду свідчать про можливі складнощі у виявленні маскованих аномалій або трафіку, схожого на нормальний. У датасеті, що включав різні типи атак, спостерігався вищий рівень хибнопозитивних результатів як першого, так і другого роду, що обумовлено більшою різноманітністю атак і складністю їхньої класифікації, особливо в разі фішингових атак або SQL-ін'єкцій, які мають складніші патерни, ніж DDoS.

Попри всі досягнення, тестування виявило й деякі обмеження інтегрованої системи. GPT-модель, хоча й демонструє високу ефективність у багатьох аспектах, залежить від якості навчання та оновлення бази знань. У разі виникнення нових типів загроз або змін у форматі логів система може потребувати додаткового налаштування, щоб забезпечити адекватну реакцію.

Прототип успішно продемонстрував здатність виявляти загрози в режимі реального часу та знижувати кількість хибних спрацьовувань, що свідчить про ефективність інтеграції GPT для обробки логів у системі SIEM.

3.4 Виявлені проблеми та пропозиції щодо вдосконалення

У процесі впровадження прототипу SIEM-системи з інтеграцією GPT були ідентифіковані певні проблеми, які потребують вирішення для підвищення ефективності системи. Основні проблеми стосуються продуктивності, адаптивності до нових кіберзагроз і масштабованості системи.

Однією з ключових проблем у роботі SIEM-систем є затримка обробки великих обсягів даних у режимі реального часу. Сучасні мережі створюють гігантські масиви логів із серверів, додатків, мережевого обладнання та інших джерел. Зі збільшенням кількості підключених пристроїв і зростанням складності кіберзагроз обчислювальні потужності SIEM починають зазнавати перевантажень, що може призводити до значних затримок у виявленні загроз.

Особливо це критично під час швидких атак, таких як DDoS або експлойти нульового дня, де навіть кілька секунд зволікання можуть дозволити зловмисникам завдати непоправної шкоди. Іноді затримки могли сягати критичних 60 або більше секунд, що дуже знижувало ефективність аналізу логів в реальному часі та більше походило для пасивного аналізу.

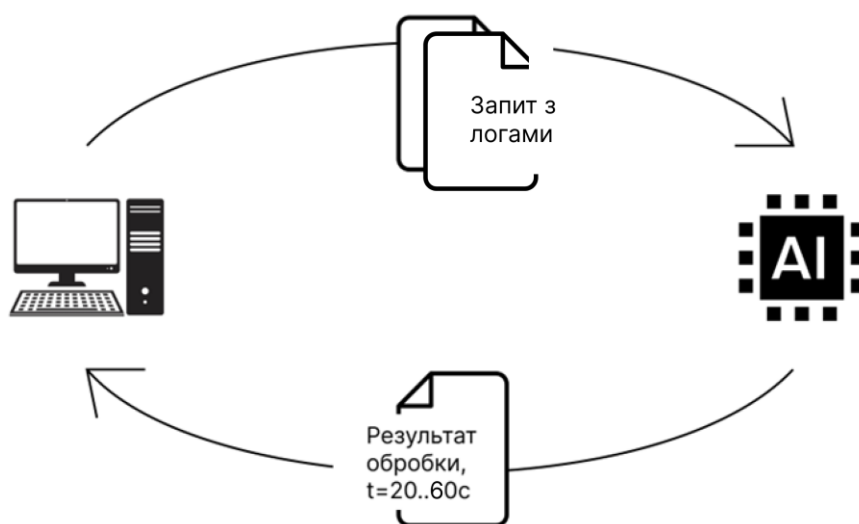


Рисунок 3.21 – Графічне представлення проблеми затримки

Наприклад, під час аналізу атак на кіберполігоні було виявлено, що затримки в обробці логів призводять до втрати можливості оперативного виявлення проникнень. У випадках, коли SIEM потребувала кілька хвилин на

обробку інформації, атаки встигали досягти своїх цілей, включно з викраденням даних або порушенням функціонування критичних систем, демонструючи нагальну потребу в модернізації інфраструктури обробки даних.

Для вирішення цієї проблеми пропонується кілька підходів. По-перше, впровадження оптимізованих алгоритмів обробки даних може суттєво зменшити час обробки логів. Алгоритми, що базуються на технологіях потокової обробки (stream processing), дозволяють аналізувати дані в реальному часі без потреби в їхньому збереженні у великих обсягах [27]. По-друге, розподілені обчислювальні системи є ефективним рішенням для рівномірного розподілу навантаження між кількома вузлами, забезпечуючи паралельну обробку логів. Наприклад, використання кластерних рішень на базі Apache Kafka може значно прискорити аналіз [28].

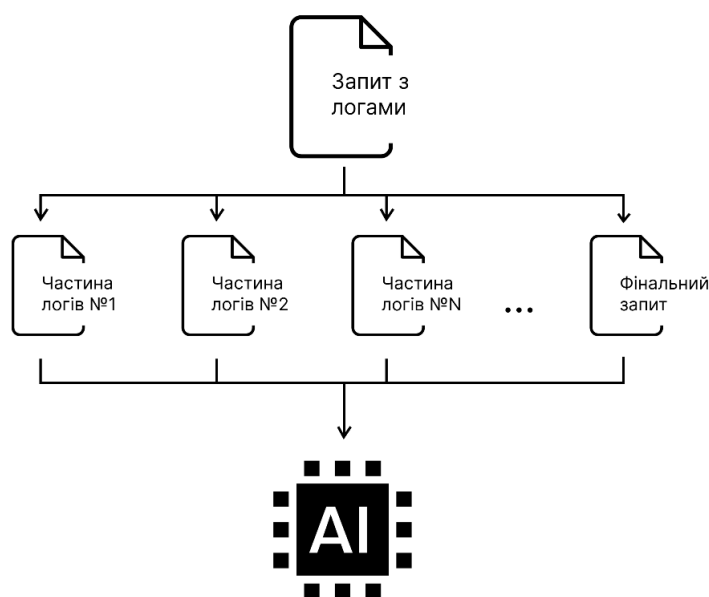


Рисунок 3.21 – Графічне представлення вирішення проблем та вдосконалення системи шляхом потокової обробки

Ще одним підходом є інтеграція хмарних платформ для масштабування обчислювальних потужностей у пікові моменти. Платформи, такі як Amazon Web Services або Microsoft Azure, дозволяють автоматично збільшувати ресурси залежно від навантаження, що допомагає уникнути перевантаження локальної інфраструктури. Зокрема, під час атак великого масштабу хмара може

забезпечити додаткові потужності, необхідні для обробки різкого збільшення обсягів логів [29].

Сучасні загрози, зокрема комплексні атаки, які охоплювали широкий спектр вразливостей, продемонстрували обмеження традиційного підходу на основі сигнатур і правил, який використовується в багатьох SIEM-системах. Традиційні методи не здатні ідентифікувати нові типи атак, для яких ще не створено сигнатур, що створює серйозні ризики для безпеки, особливо у випадках складних атак, спрямованих на певну інфраструктуру.

Для підвищення адаптивності системи запропоновано розширити використання методів машинного навчання, які дозволяють аналізувати поведінкові шаблони та виявляти аномалії. GPT, інтегрований у SIEM, потребує вдосконалення алгоритмів, що враховують контекст та специфіку роботи кожної окремої системи, з якою він взаємодіє.

Водночас важливим аспектом впровадження є формулювання коректних і конкретних запитів до моделі GPT. Неправильно сформований промпт може призвести до нерелевантних результатів, наприклад, перевірки логів на DDoS-атаки у випадку фішингової активності, і в результаті така помилка може спотворити аналітичні висновки та ускладнити виявлення реальних загроз. Щоб уникнути цього, рекомендовано використовувати загальні, але гнучкі запити, які враховують широкий спектр можливих сценаріїв. Правильна побудова промптів гарантує адекватність результатів аналізу та максимальну ефективність системи, дозволяючи їй адаптуватися до нових і складних викликів [26].

Традиційні SIEM-системи стикаються зі складнощами під час роботи в умовах збільшення обсягу даних, які надходять на аналіз. Централізовані сховища мають свої фізичні та обчислювальні межі, що призводить до затримок або збоїв у роботі системи.

Запропоновано впровадження більш гнучких архітектур на основі мікросервісів, які дозволяють масштабувати систему залежно від навантаження. Наприклад, використання розподілених баз даних і систем обробки потокових даних може допомогти ефективно зберігати та обробляти великі обсяги логів.

На основі проведеного аналізу розроблено кілька пропозицій для покращення роботи системи. Насамперед, це створення адаптивної структури запитів до GPT, яка включає контекстуальні параметри та метадані, що дозволить забезпечити більш точні відповіді та релевантні рекомендації від штучного інтелекту. Додатково, розробка інструментів для динамічного моніторингу продуктивності SIEM дозволить виявляти вузькі місця системи в режимі реального часу та оперативно їх усувати.

Підсумовуючи, ідентифіковані проблеми стали основою для формування стратегії вдосконалення SIEM-системи. Запропоновані рішення спрямовані на підвищення ефективності обробки логів, точності виявлення загроз і здатності системи адаптуватися до нових викликів у сфері кібербезпеки.

3.5 Висновок до розділу

Впродовж опрацювання третього розділу було виконано розробку, тестування та аналіз прототипу SIEM-системи з інтегрованим штучним інтелектом. Описано програмне середовище та інструменти, використані для реалізації проєкту, визначено основні етапи збирання, обробки логів та їх подальшого аналізу GPT-моделлю. Було акцентовано увагу на важливості деталізації запитів до штучного інтелекту, адже саме чіткість формулювання значно підвищила якість отриманих результатів.

Результати тестування прототипу в умовах кіберполігону засвідчили його здатність до аналізу логів у реальному часі, що забезпечило оперативне виявлення загроз. Оновлена структура запитів дозволила знизити кількість хибних тривог та отримати більш точні рекомендації щодо реагування на інциденти безпеки, що підтверджує доцільність інтеграції GPT-моделі в традиційну архітектуру SIEM і демонструє високий потенціал використання штучного інтелекту для вирішення завдань у сфері кібербезпеки.

У результаті, отримано робочий прототип системи з інтегрованим штучним інтелектом, перевірено його працездатність в польових умовах та отримано конкретні результати, які були зображені графічно у розділі.

4 ЕКОНОМІЧНА ЧАСТИНА

4.1 Оцінювання наукового ефекту

Основними та важливими ознаками наукового ефекту науково-дослідної роботи є новизна роботи, рівень її теоретичного опрацювання, перспективність, рівень розповсюдження результатів, можливість реалізації. Науковий ефект НДР на тему «Засіб моніторингу подій в інформаційній системі з використанням засобів штучного інтелекту» визначається за такими критеріями, як ступінь наукової новизни та рівень глибини теоретичного опрацювання.

Впродовж написання роботи було обрано експертів для оцінки та визначення рівня ступеня новизни. Лукічов Віталій Володимирович, Кафедра Захисту інформації, доцент; Войтович Олеся Петрівна, Кафедра Захисту інформації, к.т.н., доцент.; Каплун Валентина Аполінаріївна, Кафедра Захисту інформації, ст. викладач.

Значення зазначених критеріїв у балах, що дозволяє кількісно оцінити значущість та унікальність проведеної роботи, представлені у таблицях 4.1 та 4.2.

Таблиця 4.1 – Показники ступеня новизни науково-дослідної роботи виставлені експертами

Ступінь новизни	Характеристика ступеня новизни	Значення ступеня новизни, бали		
		Експерти (ПІБ)		
		Лукічов.В. В.	Войтович О.П.	Каплун. В.А.
Принципово нова	Робота якісно нова за постановкою задачі і ґрунтується на застосуванні оригінальних методів дослідження. Результати дослідження відкривають новий напрям в цій галузі науки і техніки. Отримано принципово нові факти, закономірності; розроблено нову теорію. Створено принципово новий пристрій, спосіб, метод	0	0	0
Нова	Отримано нову інформацію, яка суттєво зменшує невизначеність наявних значень (по-новому або вперше пояснено відомі факти, закономірності, впроваджено нові поняття, розкрито структуру змісту). Проведено суттєве вдосконалення, доповнення і уточнення раніше досягнутих результатів	59	58	63

Відносно нова	Робота має елементи новизни в постановці задачі і методах дослідження. Результати дослідження систематизують і узагальнюють наявну інформацію, визначають шляхи подальших досліджень; вперше знайдено зв'язок (або знайдено новий зв'язок) між явищами. В принципі, відомі положення поширено на велику кількість об'єктів, в результаті чого знайдено ефективне рішення. Розроблено більш прості способи для досягнення відомих результатів. Проведено часткову раціональну модифікацію (з ознаками новизни)	0	0	0
Традиційна	Робота виконана за традиційною методикою. Результати дослідження мають інформаційний характер. Підтверджені або поставлені під сумнів відомі факти та твердження, які потребують перевірки. Знайдено новий варіант рішення, який не дає суттєвих переваг в порівнянні з існуючим	0	0	0
Не нова	Отримано результат, який раніше зафіксований в інформаційному полі, та не був відомий авторам	0	0	0
Середнє значення балів експертів		60		

Залежно від середнього значення визначених експертних балів ступінь новизни характеризується як відносно нова, тобто отримана нова інформація, яка систематизує та узагальнює наявну інформацію.

Таблиця 4.2 – Показники рівня теоретичного опрацювання науково-дослідної роботи виставлені експертами

Характеристика рівня теоретичного опрацювання	Значення показника рівня теоретичного опрацювання, бали		
	Експерти (ПІБ)		
	Лукічов. В.В.	Войтович О.П.	Каплун. В.А.
Відкриття закону, розробка теорії	0	0	0
Глибоке опрацювання проблеми: багатоаспектний аналіз зв'язків, взаємозалежності між фактами з наявністю пояснень, наукової систематизації з побудовою евристичної моделі або комплексного прогнозу	0	0	0
Розробка способу (алгоритму, програми), пристрою, отримання нової речовини	57	60	63
Елементарний аналіз зв'язків між фактами та наявною гіпотезою, класифікація, практичні рекомендації для окремого випадку тощо	0	0	0
Опис окремих елементарних фактів, викладення досвіду, результатів спостережень, вимірювань тощо	0	0	0
Середнє значення балів експертів	60		

Згідно отриманого середнього значення балів експертів рівень теоретичного опрацювання науково-дослідної роботи характеризується як глибоке опрацювання проблеми: багатоаспектний аналіз зв'язків,

взаємозалежності між фактами з наявністю пояснень, наукової систематизації з побудовою евристичної моделі або комплексного прогнозу.

Показник, який характеризує рівень наукового ефекту, визначаємо за формулою [30]:

$$E_{\text{нау}} = 0,6 \cdot k_{\text{нов}} + 0,4 \cdot k_{\text{теор}}, \quad (4.1)$$

де $k_{\text{нов}}$, $k_{\text{теор}}$ – показники ступеня новизни та рівня теоретичного опрацювання науково-дослідної роботи, де $k_{\text{нов}} = 60$, $k_{\text{теор}} = 60$ балів;

0,6 та 0,4 – питома вага (значимість) показників ступеня новизни та рівня теоретичного опрацювання науково-дослідної роботи.

$$E_{\text{нау}} = 0,6 \cdot k_{\text{нов}} + 0,4 \cdot k_{\text{теор}} = 60 \text{ балів}$$

Визначення характеристики показника $E_{\text{нау}}$ і проводиться на основі висновків експертів виходячи з граничних значень, які наведені в таблиці 4.3.

Таблиця 4.3 – Граничні значення показника наукового ефекту

Досягнутий рівень показника	Кількість балів
Високий	70...100
Середній	50...69
Достатній	15...49
Низький (помилкові дослідження)	1...14

Відповідно до визначеного рівня наукового ефекту проведеної науково-дослідної роботи на тему «Модель атак на інформаційні ресурси кіберполігону», даний рівень становить 60 балів і відповідає статусу - середній рівень. Тобто у даному випадку можна вести мову про потенційну фактичну ефективність науково-дослідної роботи.

4.2 Розрахунок витрат на здійснення науково-дослідної роботи

4.2.1. Витрати на оплату праці

Основна заробітна плата дослідників розраховується відповідно до посадових окладів працівників, за формулою:

$$Z_o = \sum ((M_{pi} \times t_i) / T_p), \quad (4.2)$$

де k – кількість посад дослідників, залучених до процесу досліджень;

M_{pi} – місячний посадовий оклад конкретного дослідника, грн;

t_i – кількість днів роботи конкретного дослідника, дн.;

T_p – середня кількість робочих днів в місяці, $T_p=20$.

$$Z_o = 35000 \cdot 20 / 20 = 35000 \text{ грн.}$$

Таблиця 4.4 – Витрати на заробітну плату дослідників:

Найменування посади	Місячний посадовий оклад, грн	Оплата за робочий день, грн	Кількість днів роботи	Витрати на заробітну плату, грн.
Програміст	25000	1250	20	25000
Аналітик з кібербезпеки	35000	1750	20	35000
Керівник проекту	21000	700	5	3500
Всього				63500

Основна заробітна плата робітників

Витрати на основну заробітну плату робітників (Z_p) за відповідними найменуваннями робіт розраховуємо за формулою [30]:

$$Z_p = \sum_{i=1}^n C_i \cdot t_i, \quad (4.3)$$

де C_i – погодинна тарифна ставка робітника відповідного розряду, за виконану відповідну роботу, грн/год;

t_i – час роботи робітника при виконанні визначеної роботи, год.

Погодинну тарифну ставку робітника відповідного розряду C_i можна визначити за формулою:

$$C_i = \frac{M_M \cdot K_i \cdot K_c}{T_p \cdot t_{зм}}, \quad (4.4)$$

де M_M – розмір прожиткового мінімуму працездатної особи, або мінімальної місячної заробітної плати (в залежності від діючого законодавства), прийmemo $M_M = 8000,00$ грн;

K_i – коефіцієнт міжкваліфікаційного співвідношення для встановлення тарифної ставки робітнику відповідного розряду ;

K_c – мінімальний коефіцієнт співвідношень місячних тарифних ставок робітників першого розряду з нормальними умовами праці виробничих об'єднань і підприємств до законодавчо встановленого розміру мінімальної заробітної плати.

T_p – середнє число робочих днів в місяці, приблизно $T_p = 20$ дн;

$t_{зм}$ – тривалість зміни, год.

$$C_1 = 8000,00 * 1,5 * 1,65 / (20 * 8) = 123,75 \text{ грн.}$$

$$З_{p1} = 123,75 * 40 = 4950 \text{ грн.}$$

Таблиця 4.5 – Величина витрат на основну заробітну плату робітників

Найменування робіт	Тривалість роботи, год	Розряд роботи	Тарифний коефіцієнт	Погодинна тарифна ставка, грн	Величина оплати на робітника, грн
Розробка API інтеграції GPT	40	4	1,5	123,75	4950
Тестування API та оптимізація	32	4	1,5	123,75	3960
Налаштування середовища	16	3	1,35	111,37	1782
Аналіз логів та кіберзагроз	40	5	1,7	140,25	5610
Генерація тестового датасету	24	4	1,5	111,37	2673
Документування процесу	16	3	1,35	111,37	1782
Разом:					20757

Додаткову заробітну плату розраховуємо як 10 ... 12% від суми основної заробітної плати дослідників та робітників за формулою:

$$З_{\text{дод}} = (З_{\text{о}} + З_{\text{п}}) \cdot \frac{H_{\text{дод}}}{100\%}, \quad (4.5)$$

де $H_{\text{дод}}$ – норма нарахування додаткової заробітної плати. Прийmemo 10%.

$$З_{\text{дод}} = (63500 + 20757) \cdot 10 / 100\% = 8425,7 \text{ грн.}$$

4.2.2. Витрати на оплату праці

Нарахування на заробітну плату дослідників та робітників розраховуємо як 22% від суми основної та додаткової заробітної плати дослідників і робітників за формулою [30]:

$$Z_n = (Z_o + Z_p + Z_{\text{дод}}) \cdot \frac{H_{zn}}{100\%}, \quad (4.6)$$

де H_{zn} – норма нарахування на заробітну плату, де $H_{zn} = 22\%$.

$$Z_n = (63500 + 20757 + 8425,7) \cdot 22 / 100\% = 20390,2 \text{ грн.}$$

4.2.3. Розрахунок витрат на матеріали

Майже всі процеси відбуваються електронно, але є деяка необхідність у папері для записів, моделювання та друку. Витрати на матеріали (М), у вартісному вираженні розраховуються окремо по кожному виду матеріалів за формулою [30]:

$$M = \sum_{j=1}^n H_j \cdot C_j \cdot K_j - \sum_{j=1}^n B_j \cdot C_{vj}, \quad (4.7)$$

де H_j – норма витрат матеріалу j -го найменування, кг;

n – кількість видів матеріалів;

C_j – вартість матеріалу j -го найменування, грн/кг;

K_j – коефіцієнт транспортних витрат, ($K_j = 1,1 \dots 1,15$);

B_j – маса відходів j -го найменування, кг;

C_{vj} – вартість відходів j -го найменування, грн/кг.

$$M_1 = 2,0 \cdot 172,00 \cdot 1,1 - 0 \cdot 0 = 378,4 \text{ грн.}$$

Таблиця 4.6 – Витрати на матеріали

Найменування матеріалу, марка, тип, сорт	Ціна за 1 кг,шт/грн	Норма витрат, кг,шт	Вартість витраченого матеріалу, грн
Папір канцелярський офісний (500шт)(A4)	172,00	2	378,4
Стартовий набір в картонній коробці AS111/M	1389,00	1	1527,9
Настільний набір	378,00	2	831,6
Всього			3617,9

4.2.4. Розрахунок витрат на комплектуючі

Залежно від потреби, кількість комплектуючих кіберполігону може суттєво відрізнятись, для навчального кіберполігону на рівні університету або невеликого комерційного кіберполігону не потрібно досить серйозних затрат, проте деякі комплектуючі можуть обійтись доволі дорого, розрахуємо їх за формулою [30]:

$$K_{\text{с}} = \sum_{j=1}^n H_j \cdot C_j \cdot K_j \quad (4.8)$$

де H_j – кількість комплектуючих j -го виду, шт.;

C_j – покупна ціна комплектуючих j -го виду, грн;

K_j – коефіцієнт транспортних витрат, ($K_j = 1,1 \dots 1,15$).

$K_{\text{в}} = 1 * 150\,000 * 1,1 = 62243.5$ грн

Таблиця 4.7 – Витрати на комплектуючі

Найменування комплектуючих	Кількість, шт.	Ціна за штуку, грн	Вартість, грн
Сервер (16 ядер, 64 ГБ RAM, модель Dell PowerEdge R740)	1	150 000	165 000
Серверний жорсткий диск (2 ТБ, модель Seagate Exos X18)	2	8 000	17 600
Мережевий комутатор (Cisco Catalyst 2960)	1	12 000	13 200
Блок безперебійного живлення (APC Smart-UPS C 1000VA)	1	5 000	5 500
Кабель Ethernet (Cat 6, 10 м, модель UGREEN Ethernet Cable)	4	500	2 200
Система охолодження сервера (Cooler Master MasterAir MA620M)	1	7 000	7 700
Комплект ПК (процесор Intel Core i5-12400F, 16 ГБ RAM, SSD 512 ГБ, Lenovo V530 Tower)	2	25 000	55 000
Монітор (24-дюймовий Full HD, модель Dell P2419H)	2	6 000	13 200
Разом:			279 400

4.2.5. Розрахунок амортизації обладнання, програмних засобів та приміщень

В спрощеному вигляді амортизаційні відрахування по кожному виду обладнання, приміщень та програмному забезпеченню тощо, розраховуємо з використанням прямолінійного методу амортизації за формулою [30].

$$A_{обл} = \frac{Ц_б}{T_б} \cdot \frac{t_{вик}}{12}, \quad (4.8)$$

де $Ц_б$ – балансова вартість обладнання, програмних засобів, приміщень тощо, які використовувались для проведення досліджень, грн;

$t_{вик}$ – термін використання обладнання, програмних засобів, приміщень під час досліджень, місяців;

$T_б$ – строк корисного використання обладнання, програмних засобів, приміщень тощо, років.

$$A_{обл} = (165\,000 \cdot 1) / (2 \cdot 12) = 6875 \text{ грн.}$$

Таблиця 4.8 – Амортизаційні відрахування по кожному виду обладнання

Найменування комплектуючих	Балансова вартість, грн	Строк корисного використання, років	Термін використання обладнання, місяців	Амортизаційні відрахування, грн
Сервер (16 ядер, 64 ГБ RAM, модель Dell PowerEdge R740)	165 000	2	1	6875
Серверний жорсткий диск (2 ТБ, модель Seagate Exos X18)	17 600	2	1	733,3
Мережевий комутатор (Cisco Catalyst 2960)	13 200	2	1	550
Блок безперебійного живлення (APC Smart-UPS C 1000VA)	5 500	2	1	229,2
Система охолодження сервера (Cooler Master MasterAir MA620M)	7 700	2	1	320,8
Комплект ПК (процесор Intel Core i5-12400F, 16 ГБ RAM, SSD 512 ГБ, Lenovo V530 Tower)	55 000	2	1	2291,7
Монітор (24-дюймовий Full HD, модель Dell P2419H)	13 200	2	1	550
Разом:				11550

4.2.6. Паливо та енергія для науково-виробничих цілей

Витрати на силову електроенергію (B_e) розраховуємо за формулою [30]:

$$B_e = \sum_{i=1}^n \frac{W_{yi} \cdot t_i \cdot Ц_e \cdot K_{eni}}{\eta_i}, \quad (4.9)$$

де W_{yi} – встановлена потужність обладнання на визначеному етапі розробки, кВт;

t_i – тривалість роботи обладнання на етапі дослідження, год;

C_e – вартість 1 кВт-години електроенергії, грн; (вартість електроенергії визначається за даними енергопостачальної компанії), $C_e = 11$ грн;

K_{eni} – коефіцієнт, що враховує використання потужності, $K_{eni} < 1$;

η_i – коефіцієнт корисної дії обладнання, $\eta_i < 1$

$B_e = 0,5 * 720 * 11 * 0,95 / 0,97 = 3878,3$ грн.

Проведені розрахунки зведемо до таблиці.

Таблиця 4.9 – Витрати на електроенергію

Найменування комплектуючих	Встановлена потужність, кВт	Тривалість роботи, год	Сума, грн.
Сервер (16 ядер, 64 ГБ RAM, модель Dell PowerEdge R740)	0,5	720	3878,3
Серверний жорсткий диск (2 ТБ, модель Seagate Exos X18)	0,02	720	155,1
Мережевий комутатор (Cisco Catalyst 2960)	0,1	720	775,7
Блок безперебійного живлення (APC Smart-UPS C 1000VA)	0,06	720	465,4
Система охолодження сервера (Cooler Master MasterAir MA620M)	0,15	720	1163,5
Комплект ПК (процесор Intel Core i5-12400F, 16 ГБ RAM, SSD 512 ГБ, Lenovo V530 Tower)	0,35	160	603,3
Монітор (24-дюймовий Full HD, модель Dell P2419H)	0,025	160	43,1
Разом:			7084,4

4.2.7. Накладні (загальновиробничі) витрати

Витрати за статтею «Накладні (загальновиробничі) витрати» розраховуємо як 100...150% від суми основної заробітної плати дослідників та робітників за формулою:

$$B_{нзв} = (3_o + 3_p) \cdot \frac{H_{нзв}}{100\%}, \quad (4.10)$$

де $H_{нзв}$ – норма нарахування за статтею «Накладні (загальновиробничі) витрати», $H_{нзв} = 100\%$.

$B_{нзв} = (63500 + 20757) * 100 / 100 = 84257$ грн.

Для інтеграції моделі GPT-4o mini від OpenAI у середовище кіберполігону основні витрати пов'язані з оплатою API-доступу до моделі відповідно до тарифікації [31]:

\$0.150 за 1М вхідних токенів (input tokens);

\$0.075 за 1М вхідних токенів при зберіганні контексту (context retention);

\$0.600 за 1М вихідних токенів (output tokens).

Далі розрахуємо витрати залежно від типового сценарію використання системи. Аналіз витрат на використання GPT-4o mini. Середній сценарій роботи системи:

Кількість запитів на добу: 5 000.

Середній обсяг вхідних даних на запит: 3 000 токенів.

Середній обсяг вихідних даних на запит: 750 токенів.

Робочі дні на місяць: 30.

Розрахуємо місячний обсяг токенів:

Вхідні токени (input):

$$T_{in} = \text{Кількість запитів} * \text{Вхідні токени на запит} * \text{Кількість днів}$$

$$T_{in} = 5000 * 3000 * 30 = 450000000 \text{ токенів.}$$

Вихідні токени (output):

$$T_{out} = \text{Кількість запитів} * \text{Вихідні токени на запит} * \text{Кількість днів}$$

$$T_{in} = 5000 * 750 * 30 = 112500000 \text{ токенів.}$$

Розрахуємо вартість витрат на місяць:

$$\begin{aligned} \text{Вартість вхідних токенів} &= \frac{T_{in}}{1\text{М токенів}} * \text{Тариф} = \frac{450000000}{1000000} * 0.150 = 67.5 \text{ USD} \\ &= 67.5 * 42 = 2835 \text{ грн.} \end{aligned}$$

$$\text{Вартість вихідних токенів} = \frac{112500000}{1000000} * 0.600 = 67.5 \text{ USD} = 67.5 * 42 = 2835 \text{ грн.}$$

Загальна середня вартість використання запитів до штучного інтелекту рівна 5670 грн.

Загалом розрахуємо всі компоненти для проведення науково-дослідної роботи шляхом суми всіх попередніх значень:

$$63500 + 20757 + 8425,7 + 20390,2 + 3617,9 + 279400 + 11500 + 7084,4 + 84257 + 5670 = 504602,2 \text{ грн.}$$

Загальні витрати $ЗВ$ на завершення науково-дослідної (науково-технічної) роботи та оформлення її результатів розраховується за формулою [30]:

$$ЗВ = \frac{B_{\text{заг}}}{\eta}, \quad (4.11)$$

де η - коефіцієнт, який характеризує етап (стадію) виконання науково-дослідної роботи, прийmemo $\eta = 0,5$

$$ЗВ = 504602,2 / 0,5 = 1009204,4$$

Отже, проведено розрахунки та отримано загальну вартість витрат для проведення науково-дослідної роботи.

4.3 Оцінювання важливості та наукової значимості науково-дослідної роботи

Оцінювання та доведення ефективності виконання науково-дослідної роботи фундаментального чи пошукового характеру є достатньо складним процесом і часто базується на експертних оцінках, тому має вірогідний характер.

Для обґрунтування доцільності виконання науково-дослідної роботи на тему «Засіб моніторингу подій в інформаційній системі з використанням засобів штучного інтелекту» використовується спеціальний комплексний показник, що враховує важливість, результативність роботи, можливість впровадження її результатів у виробництво, величину витрат на роботу.

Комплексний показник K_P рівня науково-дослідної роботи може бути розрахований за формулою [30]:

$$K_P = \frac{I^n * T_C * R}{B * t} \quad (4.12)$$

де I – коефіцієнт важливості роботи. $I = 3$;

n - коефіцієнт використання результатів роботи; $n = 0$, коли результати роботи не будуть використовуватись; $n = 1$, коли результати роботи будуть використовуватись частково; $n = 2$, коли результати роботи будуть використовуватись в дослідно-конструкторських розробках; $n = 3$, коли

результати можуть використовуватись навіть без проведення дослідно-конструкторських розробок. $n=2$;

T_C – коефіцієнт складності роботи, $T_C = 3$;

R – коефіцієнт результативності роботи; якщо результати роботи плануються вище відомих, то $R = 4$; якщо результати роботи відповідають відомому рівню, то $R = 3$; якщо нижче відомих результатів, то $R = 2$;

B – вартість науково-дослідної роботи, тис. грн. $B = 504602,2$ грн;

t – час проведення дослідження. $t = 0,08$ років

$$K_P = \frac{I^n * T_C * R}{B * t} = \frac{3^2 * 3 * 4}{504,6 * 0,08} = 2,67$$

Якщо $K_P > 1$, то науково-дослідну роботу на тему «Засіб моніторингу подій в інформаційній системі з використанням засобів штучного інтелекту» можна вважати ефективною з високим науковим, технічним і економічним рівнем.

4.4 Висновки до розділу

У результаті було розраховано витрати на проведення науково-дослідної роботи, які складають 504602,2 грн. Відповідно до проведеного аналізу та розрахунків рівень науково-економічного ефекту проведеної науково-дослідної роботи на тему «Засіб моніторингу подій в інформаційній системі з використанням засобів штучного інтелекту» є середній, а дослідження актуальними, рівень доцільності виконання науково-дослідної роботи $K_P > 1$, що свідчить про потенційну ефективність з високим науковим, технічним і економічним рівнем.

ВИСНОВКИ

При виконанні магістерської кваліфікаційної роботи було проведено комплексний аналіз проблем забезпечення інформаційної безпеки в сучасних системах SIEM, зокрема в контексті збільшення обсягів даних і складності атак. Основна увага була приділена інтеграції штучного інтелекту для підвищення ефективності аналізу логів, зниження кількості хибних спрацювань та ідентифікації складних багатовекторних загроз.

У межах роботи обґрунтовано використання GPT як ефективної моделі штучного інтелекту, що здатна працювати з великими обсягами даних і різними форматами логів. Проведено аналіз можливостей моделі у виявленні аномалій, що важко розпізнати за допомогою традиційних алгоритмів.

Емпіричні дослідження включали тестування інтегрованої системи в умовах кіберполігону з використанням наборів даних різних типів атак, а отримані результати підтвердили, що використання GPT дозволяє підвищити точність виявлення загроз у реальному часі. Було виявлено залежність результатів аналізу від релевантності побудови запитів, що підкреслює важливість точного формулювання задач для ШІ.

Було проведено детальний розрахунок економічної доцільності проєкту, що охоплював усі основні аспекти витрат і ресурсів, необхідних для його реалізації. Зокрема, враховано витрати на впровадження системи, включаючи інтеграцію GPT-4o Mini у кіберполігон. У рамках розрахунків оцінено потребу в апаратному забезпеченні, що включає сервери, системи зберігання даних, мережеве обладнання та периферійні пристрої.

Отже, в результаті виконаної роботи було виконано поставлені завдання в повному обсязі та отримано робочу частину кіберполігону з додатковим функціоналом у вигляді інтегрованого штучного інтелекту, який доповнить наявний функціонал та завершить повну систему.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. González-Granadillo, Gustavo, Susana González-Zarzosa, and Rodrigo Diaz. "Security information and event management (SIEM): analysis, trends, and usage in critical infrastructures." *Sensors* 21.14 (2021): 4759.
2. Якімов О. П. Вплив використання штучного інтелекту та оптимізації промптів на ефективність виявлення загроз у системах SIEM [Електронний ресурс] / О. П. Якімов, О. П. Войтович // Матеріали Всеукраїнської науково-практичної інтернет-конференції "Молодь в науці: дослідження, проблеми, перспективи (МН-2025)", Вінниця, 15-16 червня 2025 р. – Електрон. текст. дані. – 2024. – Режим доступу: <https://conferences.vntu.edu.ua/index.php/mn/mn2025/paper/view/22360>
3. H. M. J. Almohri and S. A. Almohri, "Security Evaluation by Arrogance: Saving Time and Money," *2017 IEEE/ACM 1st International Workshop on Software Engineering for Startups (SoftStart)*, Buenos Aires, Argentina, 2017, pp. 12-16. URL: <https://doi.org/10.1109/SoftStart.2017.1> (date of access: 06.11.2024).
4. López Velásquez, J.M., Martínez Monterrubio, S.M., Sánchez Crespo, L.E. *et al.* Systematic review of SIEM technology: SIEM-SC birth. *Int. J. Inf. Secur.* 22, 691–711 (2023). URL: <https://doi.org/10.1007/s10207-022-00657-9> (date of access: 06.11.2024).
5. SIM, SEM, and SIEM: Definitions and Choosing the Right Enterprise Solution. *OpenText Community for Micro Focus products*. URL: <https://community.microfocus.com/cyberres/b/cybersecurity-blog/posts/sim-sem-and-siem-definitions-and-choosing-the-right-enterprise-solution> (date of access: 06.11.2024).
6. Di Mauro M., Di Sarno C. Improving SIEM capabilities through an enhanced probe for encrypted Skype traffic detection. *Journal of Information Security and Applications*. 2018. Vol. 38. P. 85–95. URL: <https://doi.org/10.1016/j.jisa.2017.12.001> (date of access: 06.11.2024).

7. Anilkumar, Akhila, et al. "Detecting and Analysing Network Logs Using Machine Learning Techniques." *REVISTA GEINTEC-GESTAO INOVACAO E TECNOLOGIAS* 11.3 (2021): 271-286.
8. Intrusion Detection System Using machine learning Algorithms / R. Tahri et al. *ITM Web of Conferences*. 2022. Vol. 46. P. 02003. URL: <https://doi.org/10.1051/itmconf/20224602003> (date of access: 06.11.2024).
9. Ayodele, Taiwo Oladipupo. "Types of machine learning algorithms." *New advances in machine learning* 3.19-48 (2010): 5-1.
10. Parimala, Venkata Krishna. "Introductory Chapter: Anomaly Detection—Recent Advances, AI and ML Perspectives and Applications." *Anomaly Detection—Recent Advances, AI and ML Perspectives and Applications* (2024).
11. What is Log Management? 4 Best Practices & More | CrowdStrike. *CrowdStrike: We Stop Breaches with AI-native Cybersecurity*. URL: <https://www.crowdstrike.com/en-us/cybersecurity-101/next-gen-siem/log-management/> (date of access: 06.11.2024).
12. Brad Hale, "Estimating Log Generation for Security Information Event and Log Management", URL: http://content.solarwinds.com/creative/pdf/Whitepapers/estimating_log_generation_white_paper.pdf (date of access: 06.11.2024).
13. Cao, Qimin, Yinrong Qiao, and Zhong Lyu. "Machine learning to detect anomalies in web log analysis." *2017 3rd IEEE international conference on computer and communications (ICCC)*. IEEE, 2017.
14. Zhao, Zhijun, Chen Xu, and Bo Li. "A LSTM-based anomaly detection model for log analysis." *Journal of Signal Processing Systems* 93.7 (2021): 745-751.
15. Bertero, Christophe, et al. "Experience report: Log mining using natural language processing and application to anomaly detection." *2017 IEEE 28th International Symposium on Software Reliability Engineering (ISSRE)*. 2017.
16. Jangampet, Vinay Dutt. "The Rise of The Machines: AI-Driven SIEM User Experience for Enhanced Decision-Making." *International Journal of Engineering & Technology*.12. 74-83.

17. OpenAI API | OpenAI. URL: <https://openai.com/index/openai-api/> (date of access: 15.11.2024).
18. Gemini in Google SecOps | Google Security Operations | Google Cloud. *Google Cloud*. URL: <https://cloud.google.com/chronicle/docs/secops/gemini-chronicle> (date of access: 15.11.2024).
19. Build with Claude. *Home* | *Anthropic*. URL: <https://www.anthropic.com/api> (date of access: 15.11.2024).
20. Falcon vs. Mistral: Which LLM is Better? | Sapling. *Language Model API Toolkit and Copilot* | *Sapling*. URL: <https://sapling.ai/llm/falcon-vs-mistral> (date of access: 15.11.2024).
21. Якімов О. П. Кіберполігон для дослідження подій інформаційної безпеки. Частина 3. Обробка на основі Elasticsearch: Бакалаврська дипломна робота: спец. 125 Кібербезпека / науковий керівник Войтович О. П., Вінниця – Електрон. текст. дані. – 2024. – Режим доступу: <https://iq.vntu.edu.ua/repository/card.php?lang=uk&id=6396>
22. Welcome to Python.org. *Python.org*. URL: <https://www.python.org> (date of access: 25.11.2024).
23. GitHub - wallxandr/mkr-testing-datasets: All the datasets used to test on SIEM with integrated AI. *GitHub*. URL: <https://github.com/wallxandr/mkr-testing-datasets> (date of access: 26.11.2024).
24. MITRE ATT&CK®. *MITRE ATT&CK®*. URL: <https://attack.mitre.org> (date of access: 26.11.2024).
25. CVE. *CVE*. URL: <https://cve.mitre.org> (date of access: 26.11.2024).
26. Korzynski, Pawel, et al. "Artificial intelligence prompt engineering as a new digital competence: Analysis of generative AI technologies such as ChatGPT." *Entrepreneurial Business and Economics Review* 11.3 (2023): 25-37.
27. How To Use Event Streaming with GPT-4 for Real-Time Generative AI. *Confluent*. URL: <https://www.confluent.io/blog/chatgpt-and-streaming-data-for-real-time-generative-ai/> (date of access: 26.11.2024).

28. Apache Kafka + Vector Database + LLM = Real-Time GenAI. *Kai Waehner*. URL: <https://www.kai-waehner.de/blog/2023/11/08/apache-kafka-flink-vector-database-llm-real-time-genai/> (date of access: 26.11.2024).
29. Microsoft Azure vs Amazon Web Services (AWS): порівняння хмарних сервісів. *SoftwareOne*. URL: <https://www.softwareone.com/uk-ua/blog/articles/2024/05/30/microsoft-azure-vs-amazon-web-services> (date of access: 26.11.2024).
30. Методичні вказівки до виконання економічної частини магістерських кваліфікаційних робіт / Уклад. : В. О. Козловський, О. Й. Лесько, В. В. Кавецький. – Вінниця : ВНТУ, 2021. – 42 с.
31. Pricing OpenAI API | OpenAI. URL: <https://openai.com/api/pricing/> (date of access: 6.12.2024).

ДОДАТКИ

Додаток А

ПРОТОКОЛ ПЕРЕВІРКИ
НАВЧАЛЬНОЇ (КВАЛІФІКАЦІЙНОЇ) РОБОТИ

Назва роботи: Засіб моніторингу подій в інформаційній системі з використанням засобів штучного інтелекту
Автор роботи: Якімов Олександр Павлович
Тип роботи: магістерська кваліфікаційна робота
Підрозділ: кафедра захисту інформації ФІТКІ
Керівник: Войтович Олеся Петрівна, к.т.н., доцент

Показники звіту подібності

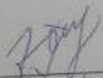
Turnitin	
Оригінальність	95 %
Загальна схожість	5 %

Аналіз звіту подібності (відмітити потрібне):

- ☒ 1. Запозичення, виявлені у роботі, оформлені коректно і не містять ознак плагіату.
- ☐ 2. Виявлені у роботі запозичення не мають ознак плагіату, але їх надмірна кількість викликає сумніви щодо цінності роботи і відсутності самостійності її автора. Роботу направити на доопрацювання.
- ☐ 3. Виявлені у роботі запозичення є недобросовісними і мають ознаки плагіату та/або в ній містяться навмисні спотворення тексту, що вказують на спроби приховування недобросовісних запозичень.

Заявляю, що ознайомлений з повним звітом подібності, який був згенерований Системою щодо роботи.

Автор роботи


(підпис)

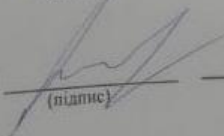
Олександр ЯКИМОВ
(прізвище, ініціали)

Особа, відповідальна за перевірку


(підпис)

Валентина КАПЛУН
(прізвище, ініціали)

Керівник роботи


(підпис)

Олеся ВОЙТОВИЧ
(прізвище, ініціали)

ІЛЮСТРАТИВНА ЧАСТИНА

Засіб моніторингу подій в інформаційній системі з використанням засобів
штучного інтелекту

(Назва магістерської кваліфікаційної роботи)

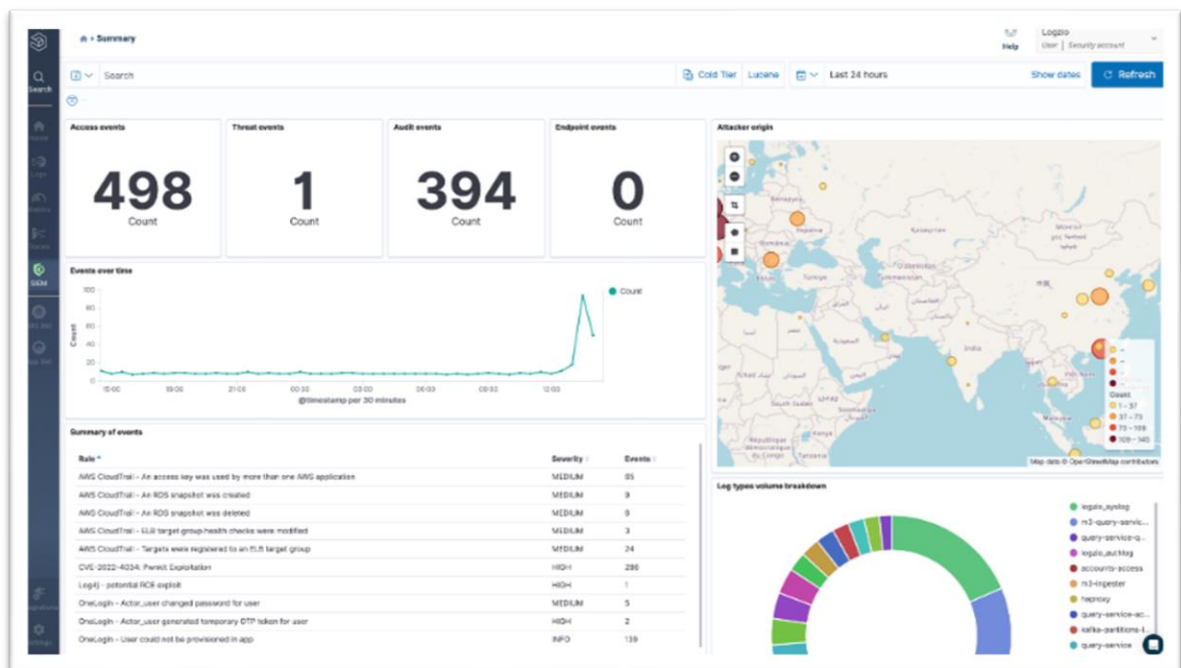
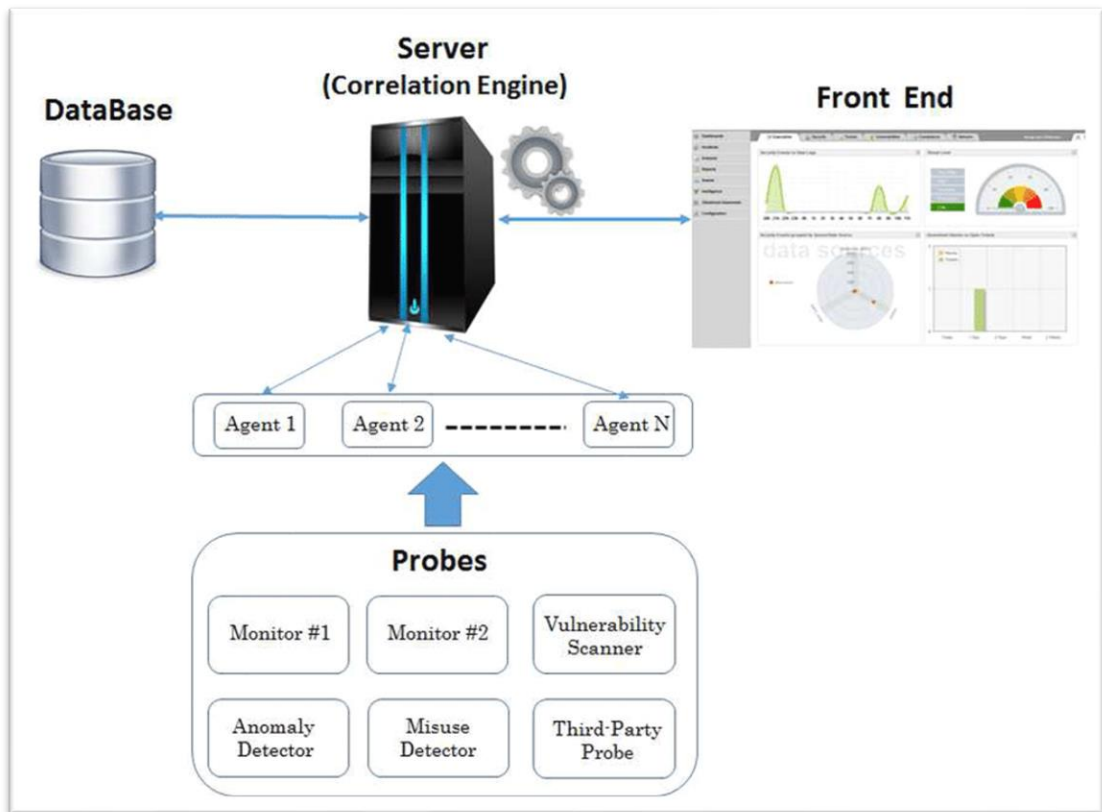
Виконав: студент 2 курсу групи ІБС-23м
спеціальності 125 Кібербезпека та захист
інформації

Якимов Олександр ЯКІМОВ
13.12 2024 р.

Керівник: к. т. н., доцент каф. ЗІ

Войтович Олеся ВОЙТОВИЧ
13.12 2024 р.

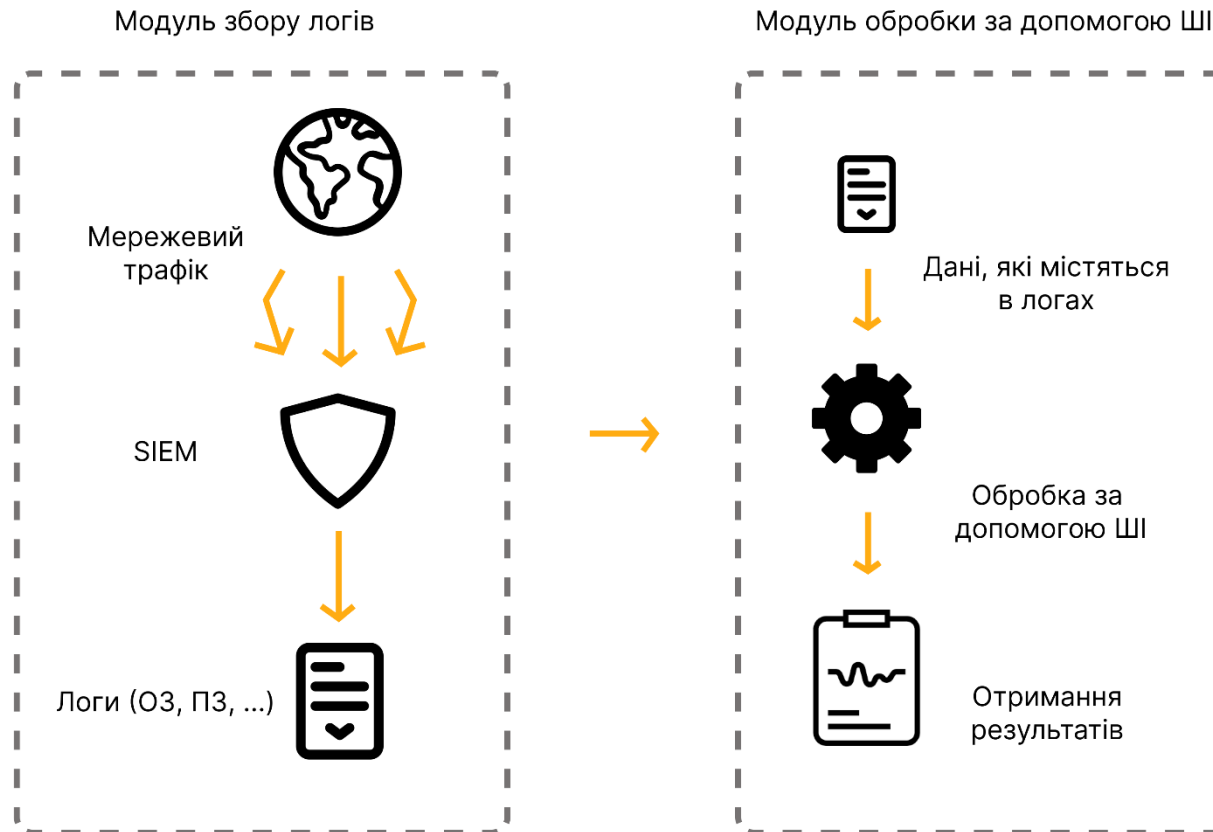
ВИГЛЯД ТИПОВОЇ SIEM-СИСТЕМИ



ЗАСТОСУВАННЯ МЕТОДІВ МАШИННОГО НАВЧАННЯ В КІБЕРБЕЗПЕЦІ

Концепція/ Метод	Опис	Застосування в кібербезпеці
Машинне навчання (ML)	Підрозділ штучного інтелекту, що фокусується на створенні алгоритмів, які можуть навчатися на даних і робити прогнози або рішення без прямого програмування.	Використовується для виявлення аномалій у мережевому трафіку, ідентифікації шкідливих програм.
Глибинне навчання (DL)	Підмножина машинного навчання, що використовує багаторівневі нейронні мережі для обробки великих обсягів даних і виявлення складних патернів.	Використовується для класифікації зображень, аудіо, а також у виявленні шкідливого ПЗ.
Аналіз поведінки користувачів (UBA)	Метод виявлення аномалій на основі вивчення звичок і поведінки користувачів у системах.	Використовується для виявлення несанкціонованого доступу або порушень політики безпеки.
Класифікація	Процес, у якому алгоритм машинного навчання призначає категорію або клас об'єкту на основі його характеристик.	Застосовується для класифікації вхідних логів за типами загроз.
Регресія	Метод, який використовує статистичні техніки для прогнозування числових значень на основі вхідних даних.	Може бути використана для прогнозування ймовірності атаки або фінансових втрат від інцидентів.
Кластери	Техніка, яка групує схожі дані в єдині категорії, що дозволяє виявляти аномалії шляхом вивчення об'єктів, що не підходять до жодної категорії.	Застосовується для виявлення нових типів загроз або атак, які раніше не були зафіксовані.
Адаптивні алгоритми	Алгоритми, які можуть змінювати свої параметри в залежності від зміни умов або вхідних даних.	Використовуються для динамічного налаштування порогів виявлення загроз у реальному часі.

ГРАФІЧНА СХЕМА ЕТАПІВ ОБРОБКИ ЛОГІВ З ШТУЧНИМ ІНТЕЛЕКТОМ



ТОПОЛОГІЯ МЕРЕЖІ

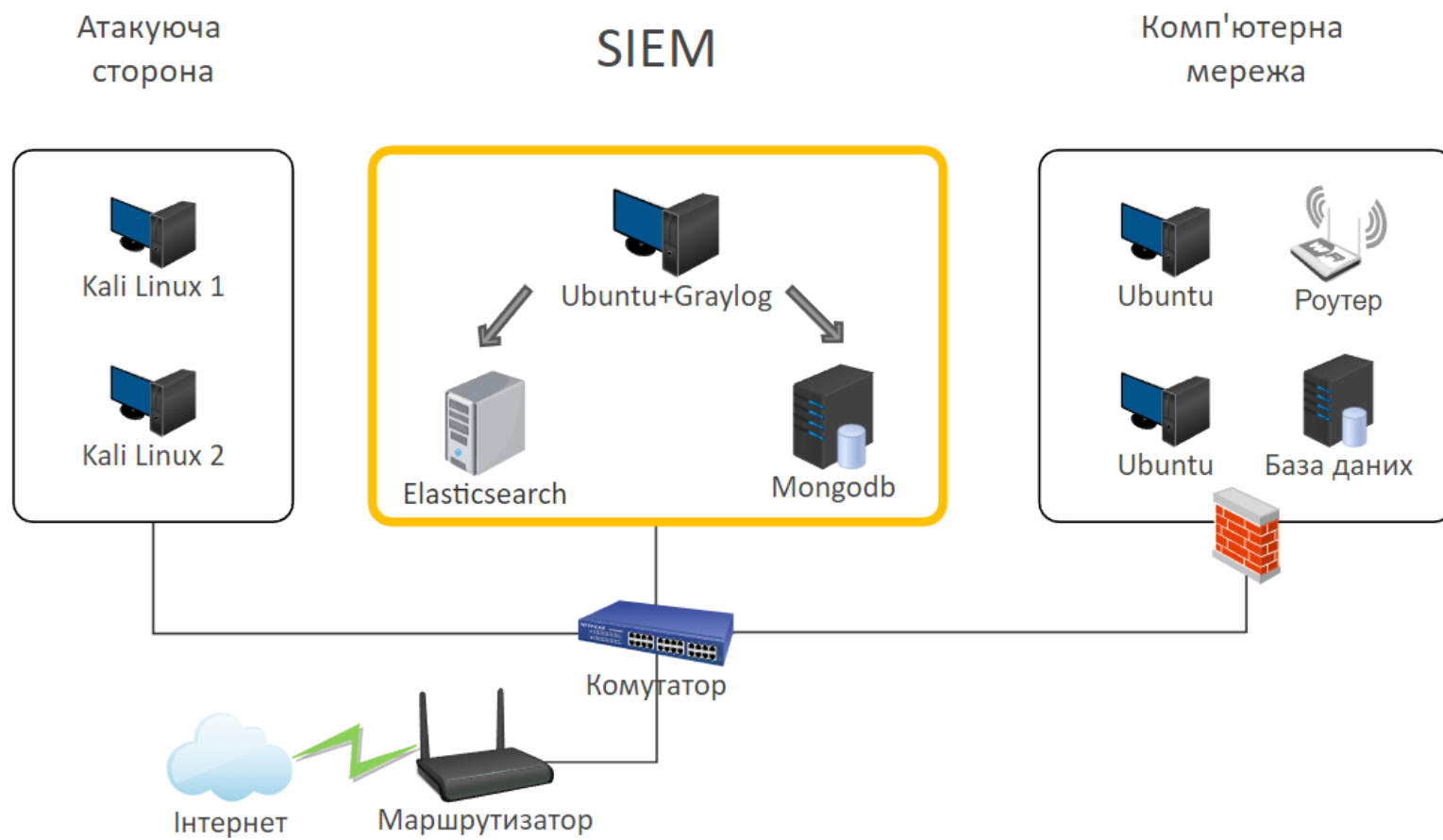
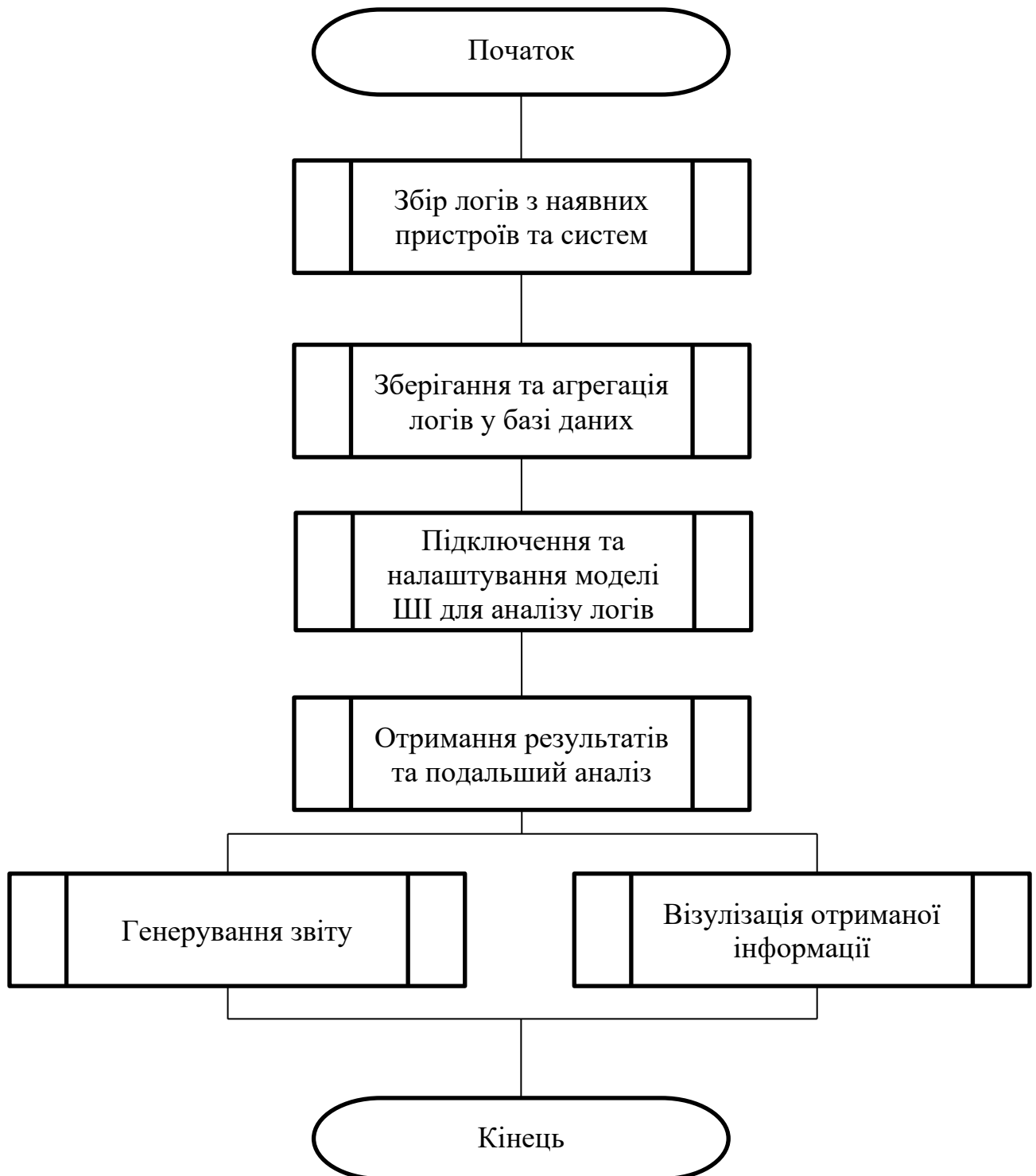
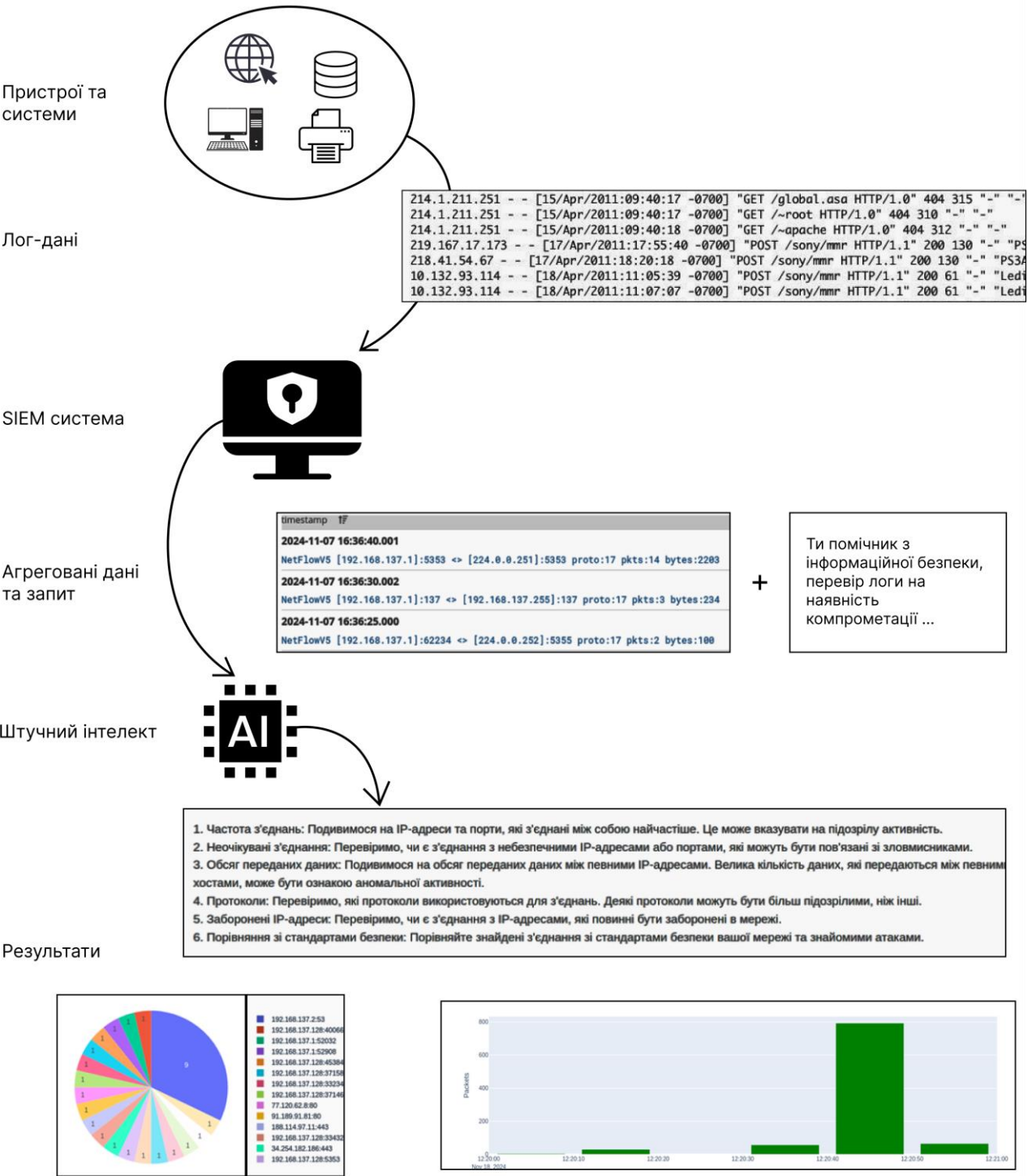


СХЕМА АЛОГРИТМУ ВПРОВАДЖЕННЯ ШТУЧНОГО ІНТЕЛЕКТУ В СИСТЕМУ SIEM



ГРАФІЧНЕ ПРЕДСТАВЛЕННЯ ШЛЯХУ ОБРОБКИ ЛОГІВ



ПРИКЛАДИ ВИГЛЯДУ ЛОГІВ

```
{
  "_index" : "graylog_0",
  "_type" : "_doc",
  "_id" : "7a745440-077b-11ee-b064-000c2990507e",
  "_score" : null,
  "_source" : {
    "nf_dst" : "239.255.255.250:1900",
    "nf_stop" : "2023-06-10T10:41:11.891Z",
    "nf_version" : 5,
    "gl2_remote_ip" : "127.0.0.1",
    "gl2_remote_port" : 49033,
    "nf_pkts" : 4,
    "source" : "127.0.0.1",
    "gl2_source_input" : "646bae9bcf82341db70912e1",
    "nf_tcp_flags" : 0,
    "nf_dst_port" : 1900,
    "nf_bytes" : 792,
    "nf_src_as" : 0,
    "nf_proto" : 17,
    "nf_dst_address" : "239.255.255.250",
    "nf_src_port" : 61292,
    "gl2_source_node" : "72db8cfe-a25e-41d8-ae67-f703c2ccfe64",
    "timestamp" : "2023-06-10 10:42:20.000",
    "nf_src" : "192.168.137.1:61292",
    "nf_dst_mask" : 0,
    "gl2_accounted_message_size" : 543,
    "nf_flow_packet_id" : 22859,
    "streams" : [
      "00000000000000000000000000000001"
    ],
    "gl2_message_id" : "01H2JECQW5M7BEG97401PCDPJG",
    "message" : "NetFlowV5 [192.168.137.1]:61292 <-> [239.255.255.250]:1900 proto:17 pkts:4 bytes:792",
    "nf_dst_as" : 0,
    "nf_tos" : 0,
    "nf_src_address" : "192.168.137.1",
    "nf_proto_name" : "UDP",
    "nf_src_mask" : 0,
    "nf_snmp_input" : 0,
    "nf_snmp_output" : 0,
    "nf_start" : "2023-06-10T10:41:08.862Z"
  },
  "sort" : [
    1686393740000
  ]
}
```

application_name
kernel
facility
kernel
facility_num
0
level
4
message
[3276.192018] workqueue: blk_mq_run_work_fn hogged CPU for >10000us 256 times, consider switching to WQ_UNBOUND
source
user-VirtualBox
timestamp
2024-11-21T15:01:54.073Z

ПРИКЛАДИ ВИГЛЯДУ ЗАПИТІВ

```
try:
    response = client.chat.completions.create(
        model="gpt-3.5-turbo-0125",
        max_tokens=1000,
        temperature=0.5,
        messages=[
            {"role": "system", "content": "Ви помічник, який має досвід аналізу журналів мережевого трафіку, щоб виявити можливі кіберзагрози."},
            {"role": "user", "content": f"Проаналізуйте наведені нижче журнали мережевого трафіку на наявність будь-яких аномалій або підозрілих дій:\n{joined_logs}"
        ]
    )
    analysis = response.choices[0].message.content.replace("\n", "<br>")
    return analysis
except Exception as e:
    return f"Помилка аналізу журналів мережі: {e}"
```

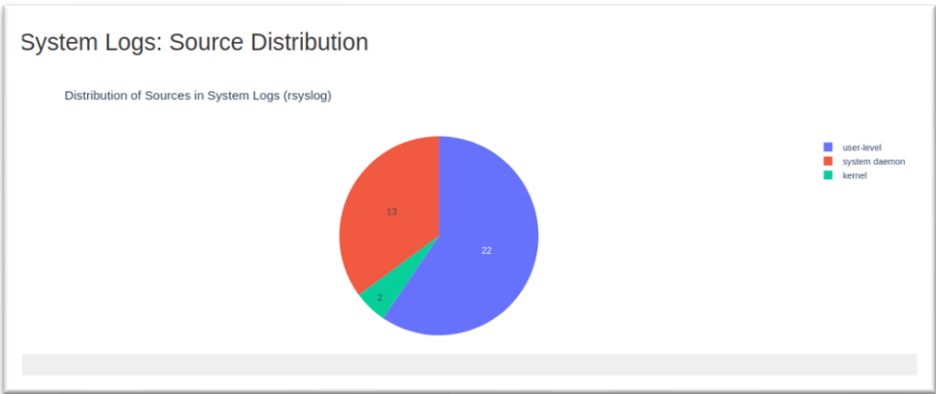
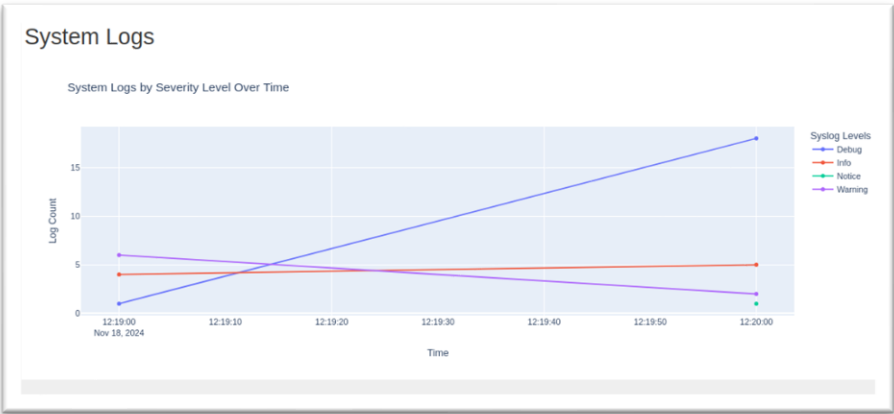
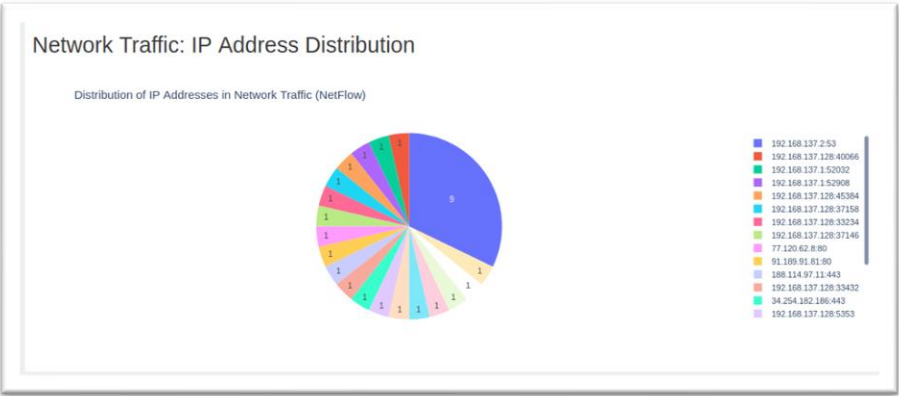
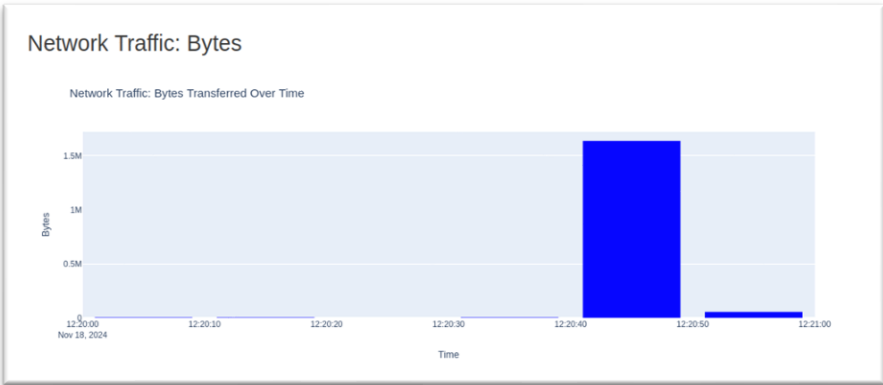
```
messages=[
    {"role": "system", "content": "Ви помічник, який знає інформацію про кібербезпеку та як аналізувати журнали для запобігання кіберзагрозам і атакам."},
    {
        "role": "user",
        "content": f"""
        Проаналізуйте події безпеки в наведених журналах системи. Визначте всі аномалії, підозрілі активності та потенційні загрози, які можуть свідчити про компрометацію системи. Використайте поведінковий аналіз для виявлення відхилень від нормальних дій користувачів та системних процесів.

        Конкретно зверніть увагу на:
        - Невідповідності у запитах доступу до критичних ресурсів.
        - Підвищені рівні активності від одного джерела або користувача.
        - Незвичну активність в нічний час або в неробочі години.
        - Використання застарілих або уразливих протоколів зв'язку.
        - Неавторизовані зміни у конфігурації системи.

        Класифікуйте знайдені інциденти за рівнем критичності: висока, середня, низька. Вкажіть індикатори компрометації (IoC), якщо вони присутні, та надайте рекомендації щодо подальших дій для усунення потенційних загроз. Також згенеруйте звіт про ключові метрики безпеки для загального огляду системи на основі зібраних даних. Для створення звіту у вас є 2000 токенів, використайте їх як варто.

        Ось журнали для аналізу:
        {joined_logs}
        """
    ]
```

ВІЗУАЛІЗАЦІЯ ДАНИХ ОТРИМАНИХ ПІСЛЯ ОБРОБКИ ШТУЧНИМ ІНТЕЛЕКТОМ



РЕКОМЕНДАЦІЇ ШТУЧНОГО ІНТЕЛЕКТУ ПІСЛЯ ОБРОБКИ ЛОГІВ

AI Analysis: Network Traffic

На основі аналізу наданих журналів мережевого трафіку можна виявити деякі підозрілі аспекти:

1. Багато різних IP-адрес спілкуються з [192.168.137.128] на порті 443. Це може свідчити про те, що вузол (192.168.137.128) є центральним пунктом для багатьох з'єднань.
2. Деякі з цих віддалених IP-адрес спілкуються з 192.168.137.128 із значною кількістю пакетів та обсягом даних, що передаються. Наприклад, IP-адрес (61.58.131.187) має 278 пакетів і 860,447 байт даних. Це може бути підозрілим, і варто подальше вивчення цього з'єднання.
3. Деякі з віддалених IP-адрес можуть бути незвичайними або недостовірними, що може вказувати на можливість шкідливої діяльності. Наприклад, IP-адреси, такі як (11.89.235.162), (0.61.145.158), [1.251.155.120] та інші, можуть вимагати уваги.
4. Підозрілість може також виникнути через велику кількість IP-адрес, які спілкуються з [192.168.137.128] на порті 443. Важливо вивчити ці з'єднання для виявлення будь-яких аномалій.

Загалом, рекомендується подальше дослідження цих з'єднань для виявлення можливих загроз та прийняття відповідних заходів для захисту мережі.

AI Analysis: System Logs

За результатами аналізу системних журналів можна виявити наступні проблеми та незвичайну активність:

1. Помилка управління монтажем простору імен, що призвела до відмови в доступі через permission denied:
update.go:85: cannot change mount namespace according to change mount (/var/lib/snapd/hostfs/usr/local/share/doc /usr/local/share/doc none bind,ro 0 0): cannot open directory "/usr/local/share": permission denied
2. Помилки оновлення джерел пакетів, що можуть вказувати на проблеми з мережею або недоступністю деяких репозиторіїв:
Failed to fetch https://packages.graylog2.org/repo/debian/dists/stable/InRelease
Failed to fetch https://repo.mongodb.org/apt/ubuntu/dists/focal/mongodb-org/4.4/InRelease
Failed to fetch https://repo.mongodb.org/apt/ubuntu/dists/jammy/mongodb-org/6.0/InRelease
Some index files failed to download. They have been ignored, or old ones used instead.
3. Повідомлення про роботу з чергами, що може вказувати на можливі проблеми з продуктивністю:
[94.903626] workqueue: blk_mq_run_work_fn hogged CPU for >10000us 8 times, consider switching to WQ_UNBOUND
4. Повідомлення про невдачу спробу оновлення іконки для livepatch:
unable to update icon for livepatch
5. Повідомлення про відсутність оновлень для деяких пакетів:
storehelpers.go:923: cannot refresh: snap has no updates available: "bare", "core18", "gnome-3-38-2004", "gnome-42-2204", "gtk-common-themes", "shellcheck", "snapd", "snapd-desktop-integration"

Загалом, в системних журналах виявлено декілька проблем, які можуть вплинути на безпеку та продуктивність системи. Рекомендується подальше вивчення та вирішення цих проблем для забезпечення стабільної та безпечної роботи системи.

РЕКОМЕНДАЦІЇ ШТУЧНОГО ІНТЕЛЕКТУ ДО ТА ПІСЛЯ ВНЕСЕНИХ ЗМІН У ВИГЛЯДІ КОНКРЕТИЗАЦІЇ ЗАПИТУ

AI Analysis: Network Traffic

Для аналізу цих журналів мережі на наявність можливих інцидентів безпеки, переглянемо кілька ключових аспектів:

1. Частота з'єднань: Подивимося на IP-адреси та порти, які з'єднані між собою найчастіше. Це може вказувати на підозрілу активність.
2. Неочікувані з'єднання: Перевіримо, чи є з'єднання з небезпечними IP-адресами або портами, які можуть бути пов'язані зі зловмисниками.
3. Обсяг переданих даних: Подивимося на обсяг переданих даних між певними IP-адресами. Велика кількість даних, які передаються між певними хостами, може бути ознакою аномальної активності.
4. Протоколи: Перевіримо, які протоколи використовуються для з'єднань. Деякі протоколи можуть бути більш підозрілими, ніж інші.
5. Заборонені IP-адреси: Перевіримо, чи є з'єднання з IP-адресами, які повинні бути заборонені в мережі.
6. Порівняння зі стандартами безпеки: Порівняйте знайдені з'єднання зі стандартами безпеки вашої мережі та знайомими атаками.

Зараз я проведу аналіз цих журналів та повідомлю вас про будь-які виявлені потенційні проблеми з безпекою.

За допомогою аналізу журналів NetFlowV5 було виявлено наступні аномалії та підозрілі активності:

1. ****Велика кількість підключень до одного джерела****: Існує підозра на DDoS-атаку або спробу перехоплення даних через зловмисні дії.
2. ****Значна кількість переданих даних****: Існує підозра на великі обсяги витоку даних або незаконне використання ресурсів.
3. ****Нехарактерна активність в нічний час****: Зазвичай нічний час використовується для хакерських атак, оскільки може бути менше ко...
4. ****Підозріле використання застарілих протоколів зв'язку****: Це може створити вразливості, які можуть бути використані для атак.
5. ****Велика кількість пакетів від одного джерела****: Це може свідчити про спробу переповнення буфера або іншу форму атаки.

Класифікація знайдених інцидентів за рівнем критичності:

1. ****Висока критичність****: Велика кількість підключень та переданих даних, підозріла активність в нічний час.
2. ****Середня критичність****: Нехарактерна активність в нічний час, використання застарілих протоколів зв'язку.
3. ****Низька критичність****: Підозріле використання застарілих протоколів зв'язку, велика кількість пакетів від одного джерела.

Індикатори компрометації (IoC):

1. IP-адреса: 88.230.253.114
2. IP-адреса: 17.144.164.92
3. IP-адреса: 209.239.242.40
4. IP-адреса: 239.203.245.133
5. IP-адреса: 61.58.131.187

Рекомендації щодо подальших дій:

1. ****Блокування IP-адрес****: Заблокуйте зазначені IP-адреси для запобігання подальшим атакам.
2. ****Моніторинг активності****: Постійно моніторте активність мережі для виявлення подібних підозрілих активностей.
3. ****Оновлення протоколів зв'язку****: Переконайтеся, що використовуються оновлені та безпечні протоколи зв'язку.
4. ****Зміна паролів та перевірка безпеки****: Перегляньте паролі та забезпечте високий рівень безпеки для критичних ресурсів.

Звіт про ключові метрики безпеки:

1. ****Кількість підключень****: Велика кількість підключень від одного джерела.
2. ****Обсяг переданих даних****: Великий обсяг переданих даних від певних IP-адрес.
3. ****Час активності****: Нехарактерна активність в нічний час.
4. ****Використання протоколів зв'язку****: Використання застарілих протоколів зв'язку.

Додатковий моніторинг та аналіз можуть допомогти виявити та запобігти можливим кіберзагрозам та атакам.