

ДЕРЖАВНИЙ УНІВЕРСИТЕТ ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ
НАВЧАЛЬНО-НАУКОВИЙ ІНСТИТУТ КІБЕРБЕЗПЕКИ ТА ЗАХИСТУ
ІНФОРМАЦІЇ

КАФЕДРА СИСТЕМ ТА ТЕХНОЛОГІЙ КІБЕРБЕЗПЕКИ

КВАЛІФІКАЦІЙНА РОБОТА

на тему:

«Технологія розслідування кіберінцидентів з використанням методів штучного
інтелекту»

зі спеціальності

125 Кібербезпека та захист інформації

(код, найменування спеціальності)

освітньо-професійної програми

Інформаційна та кібернетична безпека

(назва програми)

*Кваліфікаційна робота містить результати власних досліджень. Використання ідей,
результатів і текстів інших авторів мають посилання на відповідне джерело*

_____ Владислав ШЕВЧУК

(підпис)

Виконав: здобувач(ка) вищої освіти групи БСДМ-63

ШЕВЧУК Владислав

(прізвище, ім'я)

Керівник

к.т.н., ВЛАСЕНКО Вадим

(науковий ступінь, вчене звання, прізвище, ім'я)

Рецензент

_____ (науковий ступінь, вчене звання, прізвище, ім'я)

ДЕРЖАВНИЙ УНІВЕРСИТЕТ ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ
НАВЧАЛЬНО-НАУКОВИЙ ІНСТИТУТ КІБЕРБЕЗПЕКИ ТА ЗАХИСТУ ІНФОРМАЦІЇ

Кафедра Систем та технологій кібербезпеки

Ступінь вищої освіти Магістр

Спеціальність 125 Кібербезпека

Освітньо-професійна програма Інформаційна та кібернетична безпека

ЗАТВЕРДЖУЮ

Завідувач кафедри

Систем та технологій
кібербезпеки

Галина ГАЙДУР
“___” грудня 2024 року

З А В Д А Н Н Я
НА КВАЛІФІКАЦІЙНУ РОБОТУ

ШЕВЧУКУ Владиславу Ігоровичу

(прізвище, ім'я)

1. Тема кваліфікаційної роботи: «Технологія розслідування кіберінцидентів
з використанням методів штучного інтелекту»

керівник кваліфікаційної роботи ВЛАСЕНКО Вадим Олександрович, к.т.н.
(прізвище, ім'я, науковий ступінь, вчене звання)

затверджені наказом Державного університету інформаційно-комунікаційних
технологій від «15» жовтня 2024 року № 320.

2. Строк подання здобувачем вищої освіти кваліфікаційної роботи 15.12.2024 р.

3. Вихідні дані до кваліфікаційної роботи
актуальні загрози у кіберпросторі;
автоматизація розслідування кіберінцидентів штучним інтелектом;
наукова та технічна література, експлуатаційна документація, нормативні
документи, міжнародні стандарти.

4. Зміст розрахунково-пояснювальної записки (перелік питань, які потрібно
розробити)

1. Дослідження проблеми виявлення та аналізу кіберінцидентів.

2. Аналіз методів та інструментів для розслідування кіберінцидентів.

3. Автоматизований аналіз кіберінцидентів із застосуванням штучного

інтелекту.

5. Перелік графічного матеріалу
Презентація PowerPoint.

6. Дата видачі завдання 01.10.2024 р.

КАЛЕНДАРНИЙ ПЛАН

№ зп	Назва етапів кваліфікаційної роботи	Строк виконання етапів роботи	Примітка
1.	Визначення актуальності проблеми виявлення та аналізу загроз.	01.10.2024 р.	
2.	Аналіз наукової та технічної літератури з питань теми кваліфікаційної роботи.	22.10.2024 р.	
3.	Аналіз методів та засобів виявлення та аналізу кіберінцидентів штучним інтелектом.	27.10.2024 р.	
4.	Технологія виявлення та аналізу кіберінцидентів штучним інтелектом.	03.11.2024 р.	
5.	Розроблення рекомендацій щодо виявлення та аналізу кіберінцидентів штучним інтелектом.	15.11.2024 р.	
6.	Оформлення результатів дослідження.	26.11.2024 р.	
7.	Підготовка доповіді до захисту.	15.12.2024 р.	

Здобувач(ка) вищої освіти _____
(підпис)

Владислав ШЕВЧУК
(ім'я, прізвище)

Керівник кваліфікаційної роботи _____
(підпис)

Вадим ВЛАСЕНКО
(ім'я, прізвище)

ВІДГУК РЕЦЕНЗЕНТА

на кваліфікаційну роботу

здобувача(ки) ШЕВЧУКА Владислава

на тему: «Технологія розслідування кіберінцидентів з використанням методів штучного інтелекту»

Актуальність: Сучасний розвиток цифрових технологій створює не лише нові можливості, але й нові загрози для інформаційної безпеки, що робить проблему розслідування кіберінцидентів надзвичайно актуальною. Зростання складності атак, таких як DDoS, фішинг, експлойти та шкідливе програмне забезпечення, ускладнює їх своєчасне виявлення та ефективну нейтралізацію. У цьому контексті традиційні методи аналізу часто виявляються недостатньо швидкими та точними.

Застосування штучного інтелекту дозволяє автоматизувати процеси виявлення, аналізу та класифікації кіберзагроз, що значно підвищує ефективність і знижує витрати часу на реагування. Алгоритми машинного навчання здатні виявляти приховані патерни в мережевих логах, які складно розпізнати вручну. Таким чином, розробка та впровадження технологій розслідування кіберінцидентів із використанням штучного інтелекту не лише сприяє забезпеченню інформаційної безпеки, але й відповідає викликам часу, які потребують сучасних інноваційних рішень.

Позитивні сторони:

1. На основі проведеного аналізу, в роботі встановлено зміст проблеми виявлення та розслідування кіберінцидентів з використанням штучного інтелекту.
2. Досліджено методи та засоби розслідування кіберінцидентів з використанням штучного інтелекту
3. Розглянуто зміст технології виявлення та розслідування кіберінцидентів з використанням штучного інтелекту та розроблено рекомендації щодо її застосування.
4. Текст викладено достатньо ґрантно, послідовно. Сформульовано чіткі та змістовні висновки. Графічний матеріал оформлено якісно. Список науково-технічної літератури свідчить про вміння користуватись матеріалами за темою кваліфікаційної роботи.

Недоліки:

1. У кваліфікаційній роботі бажано було б провести аналіз умов застосування кіберінцидентів з використанням штучного інтелекту у великих організаціях.
2. Запропоновану технологію виявлення та аналізу кіберінцидентів з використанням штучного інтелекту бажано було б показати на прикладі конкретної організації.

Відзначені зауваження не впливають на загальну позитивну оцінку кваліфікаційної роботи.

Висновок: Враховуючи недоліки, кваліфікаційна робота заслуговує оцінку “добре”, а здобувач(ка) **ШЕВЧУК Владислав** – присвоєння кваліфікації магістр з кібербезпеки за освітньо-професійною програмою Інформаційна та кібернетична безпека.

Рецензент:

(науковий ступінь,
вчене звання)

(підпис)

(ім'я, прізвище)

**ДЕРЖАВНИЙ УНІВЕРСИТЕТ ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ
НАВЧАЛЬНО-НАУКОВИЙ ІНСТИТУТ КІБЕРБЕЗПЕКИ ТА ЗАХИСТУ ІНФОРМАЦІЇ**

**ПОДАННЯ
ГОЛОВІ ЕКЗАМЕНАЦІЙНОЇ КОМІСІЇ
ЩОДО ЗАХИСТУ КВАЛІФІКАЦІЙНОЇ РОБОТИ**

Направляється здобувач ШЕВЧУК Владислав до захисту кваліфікаційної роботи
(прізвище, ім'я)
спеціальності 125 Кібербезпека та захист інформації
освітньо-професійної програми Інформаційна та кібернетична безпека
(шифр і назва спеціальності)
на тему: «Технологія розслідування кіберінцидентів з використанням методів штучного
інтелекту».
Кваліфікаційна робота і рецензія додаються.

Директор інституту

(підпис)

Євгенія ІВАНЧЕНКО

(ім'я, прізвище)

Висновок керівника кваліфікаційної роботи

Здобувач(ка) ШЕВЧУК Владислав обрав тему роботи, метою якої було дослідити технологію розслідування кіберінцидентів з використанням методів штучного інтелекту та розробка рекомендацій щодо її реалізації. Перелік використаних джерел свідчить про вміння здобувачем розбиратись в наукових питаннях та застосовувати їх при дослідженнях. Під час виконання кваліфікаційної роботи ШЕВЧУК Владислав показав добру теоретичну та практичну підготовку, вміння самостійно вирішувати питання і робити висновки. Роботу виконував сумлінно, акуратно та вчасно за планом.

Все це дозволяє оцінити виконану кваліфікаційну роботу здобувача(ки) ШЕВЧУКА Владислава на оцінку «добре» та присвоїти йому кваліфікацію магістр з кібербезпеки за освітньо-професійною програмою Інформаційна та кібернетична безпека.

Керівник кваліфікаційної роботи

(підпис)

Вадим ВЛАСЕНКО

(ім'я, прізвище)

“ ____ ” грудня 2024 року

Висновок кафедри про кваліфікаційну роботу

Кваліфікаційна робота розглянута. Здобувач(ка) Шевчук Владислав допускається до захисту даної кваліфікаційної роботи в Державній екзаменаційній комісії.

Завідувач кафедри Систем та технологій кібербезпеки

(назва)

(підпис)

Галина ГАЙДУР

(ім'я, прізвище)

РЕФЕРАТ

Текстова частина кваліфікаційної роботи: 67 сторінок, 10 рисунків, 16 джерел.

Об'єкт дослідження – розслідування кіберінцидентів з використанням методів штучного інтелекту.

Предмет дослідження – технологія розслідування кіберінцидентів з використанням методів штучного інтелекту.

Мета роботи – розробити порядок застосування технології виявлення та розслідування кіберінцидентів з використанням методів штучного інтелекту.

Методи дослідження – опрацювання літератури за даною темою, аналіз експлуатаційної документації, міжнародних стандартів та їх порівняння.

Швидкість розвитку технологій та зростання кількості кіберзагроз роблять розслідування кіберінцидентів одним із ключових напрямів забезпечення інформаційної безпеки. Традиційні підходи до аналізу мережевої активності вимагають значних ресурсів і часу, що може уповільнити реагування на загрози. У відповідь на це, методи штучного інтелекту відкривають нові горизонти в автоматизації процесів виявлення, класифікації та розслідування кіберінцидентів. Завдяки машинному навчанню, алгоритми штучного навчання здатні виявляти складні патерни в даних, аналізувати мережеві логи у реальному часі та формувати прогнози щодо можливих атак. Це дозволяє підвищити точність аналізу та скоротити час реагування на загрози, що є критично важливим у сучасному світі.

В роботі досліджено проблему розслідування кіберінцидентів з використанням методів штучного інтелекту, визначено його мету та завдання. Визначено існуючі підходи до розслідування кіберінцидентів з використанням методів штучного інтелекту. Проаналізовано методи розслідування кіберінцидентів з використанням штучного інтелекту. Визначено призначення, основні функції та склад рішення розслідування кіберінцидентів з використанням методів штучного інтелекту.

На основі досліджень проведених в роботі запропоновано порядок застосування технології розслідування кіберінцидентів з використанням методів штучного інтелекту.

Галузь використання – кібербезпека інформаційних систем організації.

ІНФОРМАЦІЙНА СИСТЕМА ОРГАНІЗАЦІЇ, ВИЯВЛЕННЯ ТА АНАЛІЗ ЗАГРОЗ, ТЕХНОЛОГІЇ РОЗСЛІДУВАННЯ КІБЕРІНЦИДЕНТІВ З ВИКОРИСТАННЯМ МЕТОДІВ ШТУЧНОГО ІНТЕЛЕКТУ

ABSTRACT

Master's thesis: 67 pages, 10 figures, 16 sources.

The object of the study - the investigation of cyber incidents using artificial intelligence methods.

The subject of the study - the technology of investigating cyber incidents using artificial intelligence methods.

The purpose of the work - to develop a procedure for applying the technology of detecting and investigating cyber incidents using artificial intelligence methods.

Research methods - the study of literature on this topic, analysis of operational documentation, international standards and their comparison.

The speed of technological development and the growth of the number of cyber threats make the investigation of cyber incidents one of the key areas of ensuring information security. Traditional approaches to analyzing network activity require significant resources and time, which can slow down the response to threats. In response to this, artificial intelligence (AI) methods open up new horizons in automating the processes of detecting, classifying and investigating cyber incidents. Thanks to machine learning, AI algorithms are able to detect complex patterns in data, analyze network logs in real time, and make predictions about possible attacks. This allows you to increase the accuracy of analysis and reduce the time to respond to threats, which is critically important in the modern world.

The paper examines the problem of investigating cyber incidents using artificial intelligence methods, defines its purpose and objectives. Defines existing approaches to investigating cyber incidents using artificial intelligence methods. Analyzes methods and investigating cyber incidents using artificial intelligence methods. Defines the purpose, main functions, and composition of the solution for investigating cyber incidents using artificial intelligence methods.

Based on the research conducted in the paper, the procedure for applying the technology for investigating cyber incidents using artificial intelligence methods is proposed.

Field of application - cybersecurity of information systems of the organization.

INFORMATION SYSTEM OF THE ORGANIZATION, TECHNOLOGY OF
INVESTIGATION OF CYBER INCIDENTS USING ARTIFICIAL INTELLIGENCE
METHODS

ЗМІСТ

ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ	9
ВСТУП.....	11
1. ДОСЛІДЖЕННЯ ПРОБЛЕМИ ВИЯВЛЕННЯ ТА АНАЛІЗУ	
КІБЕРІНЦИДЕНТІВ	13
1.1 Дослідження актуальних загроз у кіберпросторі.....	13
1.2 Аналіз підходів до виявлення кіберінцидентів	23
1.3 Огляд сучасних рішень для ідентифікації кіберінцидентів.....	31
2. АНАЛІЗ МЕТОДІВ ТА ІНСТРУМЕНТІВ ДЛЯ РОЗСЛІДУВАННЯ	
КІБЕРІНЦИДЕНТІВ	37
2.1 Характеристика основних методів розслідування кіберінцидентів	37
2.2 Порівняння інструментів для аналізу логів та даних кібербезпеки для розслідування кіберінцидентів за допомогою штучного інтелекту.....	40
2.2 Вибір оптимального підходу до автоматизації розслідування за допомогою штучного інтелекту	52
3. ТЕХНОЛОГІЯ РОЗСЛІДУВАННЯ КІБЕРІНЦИДЕНТІВ З	
ВИКОРИСТАННЯМ МЕТОДІВ ШТУЧНОГО ІНТЕЛЕКТУ	57
3.1 Порядок виявлення та розслідування кіберінцидентів за допомогою застосування штучного інтелекту.....	57
3.2 Рекомендації щодо виявлення та аналізу загроз для подальшого розслідування кіберінцидентів	65
3.3 Оптимізація та перспективи розвитку автоматизованого аналізу з подальшим розслідуванням кіберінцидентів	67
ВИСНОВКИ	70
ПЕРЕЛІК ПОСИЛАНЬ.....	71
ДЕМОНСТРАЦІЙНІ МАТЕРІАЛИ (Презентація).....	73

ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ

ОС – операційна система

ШІ – штучний інтелект

ІТ - Інформаційні технології

ШПЗ - Шкідливе програмне забезпечення

IR – Incident response

DDoS - Denial-of-service attack

WMI - Windows Management Instrumentation

MitM - Man-in-the-Middle

DNS - Domain Name System

IP - Internet Protocol

HTTPS - Hypertext transfer protocol secure

DoS - Denial-of-service attack

UDP - User Datagram Protocol

ICMP - Internet Control Message Protocol

NTP - Network Time Protocol

SQL - Structured query language

LDAP - Lightweight Directory Access Protocol

XXE - eXternal Entities Injection

XSS - Cross-site scripting

IDS - Intrusion detection system

IPS - Intrusion detection system

GDPR – General Data Protection Regulation

SVM – Support vector machine

LOF - Local Outlier Factor

DL - Deep learning architectures

CNN - Convolutional Neural Networks

RNN - Recurrent Neural Networks

GAN - Generative Adversarial Networks

RL - Reinforcement Learning

NSM - Network Security Monitoring

TLS – Transport Layer Security

FTP – File Transfer Protocol

SSH - Secure Shell

ML - Machine Learning

SOAR - Security orchestration, automation and response

IOC - Indicator of compromise

SIEM - Security Information and Event Management)

EDR - Endpoint Detection and Response

RCA - Root cause analysis

AWS - Amazon Web Services

ВСТУП

Актуальність дослідження. Методи штучного інтелекту відкривають нові можливості для автоматизації та підвищення ефективності процесів розслідування кіберінцидентів. Завдяки використанню алгоритмів машинного навчання, систем обробки великих даних і нейронних мереж, штучний інтелект дозволяє аналізувати величезні обсяги логів, виявляти аномалії, ідентифікувати складні загрози та прогнозувати потенційні атаки. Це робить технології штучного інтелекту незамінними в умовах зростаючих кіберзагроз.

Вищесказане визначає актуальність теми даної кваліфікаційної роботи, основний зміст якої становлять дослідження технології розслідування кіберінцидентів з використанням методів штучного інтелекту.

Об'єкт дослідження – виявлення та розслідування загроз технологією штучного інтелекту.

Предмет дослідження – технологія виявлення та розслідування загроз штучним інтелектом.

Мета роботи – розробити порядок застосування технології розслідування кіберінцидентів з використанням методів штучного інтелекту та рекомендації щодо її реалізації.

Наукові завдання:

дослідити сутність проблеми виявлення та розслідування загроз у кіберпросторі;

проаналізувати підходи до виявлення та розслідування загроз у кіберпросторі;

проаналізувати існуючі рішення для виявлення та розслідування загроз штучним інтелектом;

проаналізувати методи та засоби виявлення та розслідування загроз штучним інтелектом;

розкрити порядок реалізації технології виявлення та розслідування загроз штучним інтелектом;

Методи дослідження – опрацювання літератури за даною темою, аналіз експлуатаційної документації, міжнародних стандартів та їх порівняння.

Практичне значення одержаних результатів: запропоновано порядок застосування технології розслідування кіберінцидентів з використанням методів штучного інтелекту, а також розроблено рекомендації фахівцям з кібербезпеки щодо її реалізації.

Результати кваліфікаційної роботи апробовані на Всеукраїнській науковій конференції «Актуальні проблеми кібербезпеки», яка відбулася 25 жовтня 2024 року в Державному університеті інформаційно-комунікаційних технологій, м. Київ.

1. ДОСЛІДЖЕННЯ ПРОБЛЕМИ ВИЯВЛЕННЯ ТА АНАЛІЗУ КІБЕРІНЦИДЕНТІВ

1.1 Дослідження актуальних загроз у кіберпросторі

Сучасне життя стало набагато комфортнішим завдяки різноманітним цифровим пристроям та Інтернету для їх підтримки. У всього хорошого є зворотний бік, і це також стосується сучасного цифрового світу. Інтернет приніс позитивні зміни в сьогодинішнє життя, але разом із цим постає величезна проблема щодо захисту даних. Це призводить до кібератак.

Кібератаки мають безліч негативних наслідків. Здійснення атаки може призвести до витоку даних, що призведе до втрати даних або їх маніпулювання. Організації зазнають фінансових втрат, довіра клієнтів підривається, а репутація завдає шкоди. Сьогодні у світі існує багато різновидів кібератак. Якщо відомо про різні типи кібератак, фахівцям стає легше захистити мережі та системи.

Перелік поширених джерел кіберзагроз для організацій:

Національні держави — ворогуючі країни можуть здійснювати кібератаки проти місцевих компаній та установ з метою втручання в комунікації, спричинення безладу та завдання шкоди.

Терористичні організації — терористи здійснюють кібератаки, спрямовані на знищення або зловживання критичною інфраструктурою, загрозу національній безпеці, підлив економіки та нанесення тілесних ушкоджень громадянам.

Злочинні групи — організовані групи хакерів, спрямовані на злам комп'ютерних систем для отримання економічної вигоди. Ці групи використовують фішинг, спам, шпигунське та шкідливе програмне забезпечення для вимагання, крадіжки особистої інформації та онлайн-шахрайства.

Хакери — окремі хакери атакують організації, використовуючи різноманітні методи атак. Зазвичай вони керуються особистою вигодою, помстою, фінансовою вигодою або політичною діяльністю. Хакери часто створюють нові загрози, щоб

розвинути свої злочинні здібності та покращити свій особистий статус у хакерській спільноті.

Зловмисні інсайдери — працівники, які мають законний доступ до активів компанії та зловживають своїми привілеями, щоб викрасти інформацію чи пошкодити комп'ютерні системи заради економічної чи особистої вигоди. Інсайдери можуть бути співробітниками, підрядниками, постачальниками або партнерами цільової організації. Вони також можуть бути сторонніми особами, які зламали привілейований обліковий запис і видають себе за його власника.

Типи загроз кібербезпеці

Атаки шкідливих програм

Шкідливе програмне забезпечення — це абревіатура від «шкідливого програмного забезпечення», яке включає віруси, хробаки, трояни, шпигунське програмне забезпечення та програми-вимагачі, і є найпоширенішим типом кібератак. ШПЗ проникає в систему, як правило, через посилання на ненадійному веб-сайті, електронну пошту чи завантаження небажаного програмного забезпечення. ШПЗ розгортається в цільовій системі, збирає конфіденційні дані, маніпулює та блокує доступ до мережевих компонентів і може знищити дані або повністю вимкнути систему.

Ось деякі з основних типів атак ШПЗ:

Віруси — частина коду впроваджується в програму. Під час запуску програми виконується шкідливий код.

Хробаки — ШПЗ, яке використовує вразливості програмного забезпечення та бекдори для отримання доступу до операційної системи. Після встановлення в мережі хробак може здійснювати такі атаки, як розподілена відмова в обслуговуванні (DDoS).

Трояни — шкідливий код або програмне забезпечення, яке видає себе за невинну програму, ховається в програмах, іграх або вкладеннях електронної пошти. Нічого не підозрюючи користувач завантажує троян, дозволяючи йому отримати контроль над його пристроєм.

Програми-вимагачі — користувачеві або організації заборонено доступ до власних систем або даних за допомогою шифрування. Зловмисник зазвичай вимагає сплати викупу в обмін на ключ розшифровки для відновлення доступу, але немає гарантії, що сплата викупу дійсно відновить повний доступ або функціональність.

Cryptojacking — зловмисники розгортають програмне забезпечення на пристрої жертви та починають використовувати її обчислювальні ресурси для генерації криптовалюти без її відома. Уражені системи можуть стати повільними, а набори для криптозлому можуть вплинути на стабільність системи.

Шпигунське програмне забезпечення — зловмисник отримує доступ до даних користувача, який нічого не підозрює, включаючи конфіденційну інформацію, таку як паролі та платіжні деталі. Шпигунське програмне забезпечення може впливати на настільні браузері, мобільні телефони та настільні програми.

Рекламне програмне забезпечення — активність веб-переглядача користувача відстежується для визначення моделей поведінки та інтересів, що дозволяє рекламодавцям надсилати користувачеві цільову рекламу. Рекламне програмне забезпечення відноситься до шпигунського програмного забезпечення, але не передбачає встановлення програмного забезпечення на пристрої користувача та не обов'язково використовується зі зловмисною метою, але воно може використовуватися без згоди користувача та порушувати його конфіденційність.

Безфайлове зловмисне програмне забезпечення — в операційній системі не встановлено програмне забезпечення. Власні файли, такі як WMI та PowerShell, редагуються, щоб увімкнути шкідливі функції. Цю приховану форму атаки важко виявити (антивірус не може ідентифікувати її), оскільки скомпрометовані файли розпізнаються як законні.

Руткити — програмне забезпечення, яке впроваджується в програми, мікропрограми, ядра операційної системи або гіпервізори, що забезпечує віддалений адміністративний доступ до комп'ютера. Зловмисник може запустити операційну систему в скомпрометованому середовищі, отримати повний контроль над комп'ютером і доставити додаткові шкідливі програми.

Атаки соціальної інженерії

Соціальна інженерія полягає в тому, щоб обманним шляхом змусити користувачів надати точку входу для зловмисного програмного забезпечення. Жертва надає конфіденційну інформацію або мимоволі встановлює зловмисне програмне забезпечення на свій пристрій, оскільки зловмисник видає себе за законного суб'єкта.

Ось деякі з основних типів атак соціальної інженерії:

Наманювання — зловмисник заманює користувача в пастку соціальної інженерії, зазвичай обіцяючи щось привабливе, наприклад безкоштовну подарункову картку. Жертва надає зловмиснику конфіденційну інформацію, наприклад облікові дані.

Перетекст — подібно до цькування, зловмисник змушує ціль надати інформацію під фальшивими приводами. Зазвичай це передбачає видавання себе за когось із повноваженнями, наприклад, офіцера податкової служби чи поліції, чия посада змусить жертву підкоритися.

Фішинг — зловмисник надсилає електронні листи, нібито надійшли з надійного джерела. Фішинг часто передбачає надсилання шахрайських електронних листів якомога більшої кількості користувачів, але також може бути більш цілеспрямованим. Наприклад, «фішинг» персоналізує електронну пошту для націлювання на конкретного користувача, тоді як «китобійний обіг» робить це ще далі, націлюючи на важливих осіб, таких як генеральні директори.

Вішинг (голосовий фішинг) — самозванець використовує телефон, щоб обманом змусити ціль розкрити конфіденційні дані або надати доступ до цільової системи. Vishing зазвичай націлений на людей похилого віку, але може бути використаний проти будь-кого.

Смішинг (SMS-фішинг) — зловмисник використовує текстові повідомлення як засіб обману жертви.

Підключення — авторизований користувач надає фізичний доступ іншій особі, яка «вилучає» облікові дані користувача. Наприклад, працівник може надати

доступ комусь, який видавав себе за нового працівника, який загубив свою облікову картку.

Перехоплення — неавторизована особа слідує за авторизованим користувачем у певне місце, наприклад, швидко прослизавши через захищені двері після того, як авторизований користувач їх відкрив. Ця техніка схожа на контрейлерство, за винятком того, що особа, яку слідкують, не підозрює, що її використовує інша особа.

Атаки на ланцюги поставок

Атаки на ланцюги поставок є новим видом загрози для розробників і постачальників програмного забезпечення. Його метою є зараження законних додатків і розповсюдження зловмисного програмного забезпечення за допомогою вихідного коду, створення процесів або механізмів оновлення програмного забезпечення.

Зловмисники шукають незахищені мережеві протоколи, серверну інфраструктуру та методи кодування та використовують їх, щоб скомпрометувати процес створення та оновлення, змінити вихідний код і приховати шкідливий вміст.

Атаки на ланцюг поставок є особливо серйозними, оскільки програми, які зламали зловмисники, підписані та сертифіковані надійними постачальниками. Під час атаки на ланцюг поставок програмного забезпечення постачальник програмного забезпечення не знає, що його програми чи оновлення заражені шкідливим програмним забезпеченням. Шкідливий код працює з тими самими довірою та привілеями, що й скомпрометована програма.

Типи атак на ланцюги поставок включають:

Порушення інструментів побудови або конвеєрів розробки

Порушення процедур підписання коду або облікових записів розробників

Шкідливий код, який надсилається як автоматичне оновлення апаратного забезпечення або компонентів мікропрограми

Шкідливий код, попередньо встановлений на фізичних пристроях

Атака "людина посередині".

Атака Man-in-the-Middle передбачає перехоплення зв'язку між двома кінцевими точками, такими як користувач і програма. Зловмисник може підслуховувати спілкування, викрадати конфіденційні дані та видавати себе за кожну сторону, яка бере участь у спілкуванні.

Приклади атак MitM:

Прослуховування Wi-Fi — зловмисник встановлює з'єднання Wi-Fi, видаючи себе за законного суб'єкта, наприклад компанію, до якого користувачі можуть підключитися. Шахрайський Wi-Fi дозволяє зловмиснику відстежувати активність підключених користувачів і перехоплювати такі дані, як дані платіжної картки та облікові дані для входу.

Викрадення електронної пошти — зловмисник підробляє електронну адресу законної організації, наприклад банку, і використовує її, щоб оманом змусити користувачів надати конфіденційну інформацію або переказати гроші зловмиснику. Користувач виконує інструкції, які, на його думку, надходять від банку, але насправді від зловмисника.

Підробка DNS — сервер доменних імен (DNS) підроблено, спрямовуючи користувача на зловмисний веб-сайт, який видає себе за законний сайт. Зловмисник може перенаправити трафік із законного сайту або викрасти облікові дані користувача.

IP-спуфінг — адреса Інтернет-протоколу (IP) з'єднує користувачів із певним веб-сайтом. Зловмисник може підробити IP-адресу, щоб видати її за веб-сайт і змусити користувачів подумати, що вони взаємодіють із цим веб-сайтом.

Підробка HTTPS — HTTPS зазвичай вважається більш безпечною версією HTTP, але також може використовуватися, щоб змусити браузер подумати, що шкідливий веб-сайт безпечний. Зловмисник використовує в URL-адресі «HTTPS», щоб приховати шкідливий характер веб-сайту.

Атака на відмову в обслуговуванні

Атака типу «відмова в обслуговуванні» (DoS) перевантажує цільову систему великим обсягом трафіку, що перешкоджає нормальному функціонуванню

системи. Атака із залученням кількох пристроїв відома як розподілена атака на відмову в обслуговуванні (DDoS).

Методи DoS-атаки включають:

HTTP flood DDoS — зловмисник використовує HTTP-запити, які здаються законними, щоб перевантажити програму або веб-сервер. Ця техніка не вимагає високої пропускної здатності або неправильно сформованих пакетів і зазвичай намагається змусити цільову систему виділяти якомога більше ресурсів для кожного запиту.

SYN flood DDoS — ініціювання послідовності підключення протоколу керування передачею (TCP) передбачає надсилення запиту SYN, на який хост повинен відповісти SYN-ACK, який підтверджує запит, а потім запитувач повинен відповісти ACK. Зловмисники можуть використовувати цю послідовність, зв'язуючи ресурси сервера, надсилаючи запити SYN, але не відповідаючи на запити SYN-ACK від хоста.

UDP flood DDoS — віддалений хост переповнюється пакетами протоколу дейтаграм користувача (UDP), надісланими на випадкові порти. Ця техніка змушує хост шукати програми на уражених портах і відповідати пакетами «Destination Unreachable», що використовує ресурси хоста.

ICMP flood — шквал пакетів ICMP Echo Request переповнює ціль, споживаючи як вхідну, так і вихідну смугу пропускання. Сервери можуть намагатися відповісти на кожен запит пакетом ICMP Echo Reply, але не можуть встигати за швидкістю запитів, тому система сповільнюється.

Посилення NTP — сервери протоколу мережевого часу (NTP) загальнодоступні, і зловмисники можуть використати їх для надсилення великих обсягів трафіку UDP на цільовий сервер. Це вважається атакою посилення через співвідношення запитів до відповідей від 1:20 до 1:200, що дозволяє зловмиснику використовувати відкриті NTP-сервери для виконання великих обсягів DDoS-атак із високою пропускною здатністю.

Ін'єкційні атаки

Ін'єкційні атаки використовують різноманітні вразливості, щоб безпосередньо вставити зловмисний вхід у код веб-програми. Успішні атаки можуть розкрити конфіденційну інформацію, виконати DoS-атаку або скомпрометувати всю систему.

Деякі з основних векторів ін'єкційних атак :

SQL-ін'єкція — зловмисник вводять SQL-запит у канал введення кінцевого користувача, наприклад у веб-форму чи поле для коментарів. Уразлива програма надсилає дані зловмисника до бази даних і виконає будь-які команди SQL, які були введені в запит. Більшість веб-додатків використовують бази даних на основі мови структурованих запитів (SQL), що робить їх уразливими до впровадження SQL. Новим варіантом цієї атаки є атаки NoSQL, націлені на бази даних, які не використовують реляційну структуру даних.

Впровадження коду — зловмисник може впровадити код у програму, якщо вона вразлива. Веб-сервер виконує шкідливий код так, ніби він є частиною програми.

Ін'єкція команд ОС — зловмисник може використати вразливість ін'єкції команд, щоб вводити команди для виконання операційною системою. Це дозволяє атаці викрадати дані ОС або заволодіти системою.

Ін'єкція LDAP — зловмисник вводять символи, щоб змінити запити LDAP. Система є вразливою, якщо вона використовує недезінфіковані запити LDAP. Ці атаки дуже серйозні, оскільки сервери LDAP можуть зберігати облікові записи користувачів і облікові дані для всієї організації.

XML eXternal Entities Injection — атака здійснюється за допомогою спеціально створених документів XML. Це відрізняється від інших векторів атак, оскільки використовує вразливі місця в застарілих аналізаторах XML, а не неперевірені дані користувача. Документи XML можна використовувати для проходження шляхів, віддаленого виконання коду та виконання підробки запитів на стороні сервера (SSRF).

Міжсайтовий сценарій (XSS) — зловмисник вводить рядок тексту, що містить шкідливий JavaScript. Браузер цілі виконує код, дозволяючи зловмиснику перенаправляти користувачів на зловмисний веб-сайт або викрадати файли cookie сеансу, щоб захопити сеанс користувача. Програма є вразливою до XSS, якщо вона не очищає введені користувачем дані для видалення коду JavaScript.



Malware Attacks

Viruses
Worms
Trojans
Ransomware
Cryptojacking

Spyware
Adware
Fileless malware
Rootkits



Social Engineering

Baiting
Pretexting
Phishing
Vishing (voice phishing)

Smishing
Piggybacking
Tailgating



Man-in-the-Middle

Wi-Fi eavesdropping
Email hijacking

DNS spoofing
IP spoofing
HTTPS spoofing



Denial-of-Service

HTTP flood DDoS
SYN flood DDoS
UDP flood DDoS
ICMP flood
NTP amplification



Injection

SQL injection
Code injection
OS command injection
LDAP injection

XML eXternal Entities (XXE)
Injection Cross-Site Scripting (XSS)

Рисю 1.1 Основні типи загроз

Безсумнівно з кожним роком кібератаки зростають. Кількість витоків даних продовжує зростати порівняно з попередніми роками, що вже було дуже страшно.

Також спостерігалось експоненціальне зростання складності та інтенсивності кібератак, таких як соціальна інженерія, програми-вимагачі та DDOS-атаки. Здебільшого це стало можливим завдяки хакерам, які використовували інструменти ШІ.

Останні кілька років спостерігається постійне зростання вартості порушень. Дозволяючи людям працювати з дому, компанії створили нові діри в безпеці, якими хакери можуть скористатися зі своїх домашніх офісів. Ці діри значно розширили зону кібератаки.

Крім того, поширеність зловмисного програмного забезпечення та хакерів у всіх комерційних галузях робить усіх, хто підключений до Інтернету, більш вразливими до злomu. Є занадто багато злочинних ворогів і занадто багато доступних точок входу, які можна приборкати та пом'якшити. На жаль, у 2024 році кіберстатистика залишатиметься тривожною.

У третьому кварталі 2024 року середня щотижнева кількість кібератак на організацію зросла до історичного максимуму в 1876, що означає приголомшливе зростання на 75% порівняно з тим самим періодом 2023 року та зростання на 15% порівняно з попереднім кварталом. Цей сплеск підкреслює тенденцію до того, що віртуальні загрози стають все більш частими та витонченими.

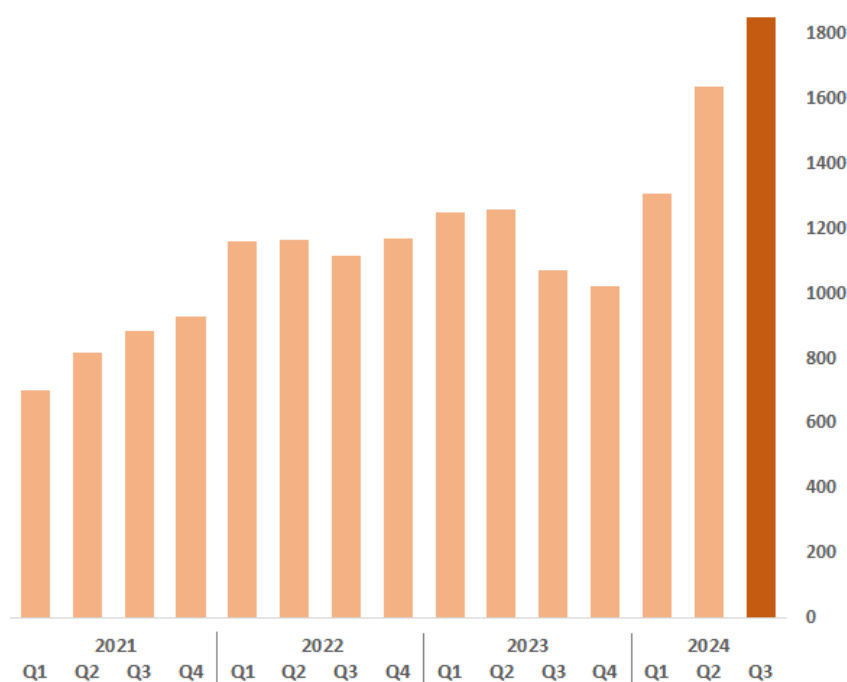


Рис. 1.2 Тенденції кібератак

1.2 Аналіз підходів до виявлення кіберінцидентів

Виявлення загроз на основі сигнатур було центральною фігурою в кібербезпеці з самого початку. Але коротка історія виявлення загроз на основі сигнатур у безпеці кінцевих точок і мережі також показує, що обмеження методів на основі сигнатур постійно спонукали практикуючих спеціалістів і постачальників до того, щоб зрештою повернутися до поведінкових методів. У результаті виявлення на основі сигнатур часто було зайвим на користь виявлення поведінкових загроз.

Протягом багатьох років підписи називали різними речами, включаючи «евристики» та «правила». Суть полягає в тому, що виявлення на основі сигнатур залежить від відповідності. Це може означати збіг відомої атаки, як-от IP-адреса чи файл. Або це може означати зіставлення фрагмента коду з відомими вірусами чи шкідливим програмним забезпеченням. Зрештою, виявлення на основі сигнатур намагається зіставити поточний трафік, поведінку чи активність зі списком «відомих шкідливих компонентів».

Сигнатури були присутні на значній частині ринку мережевої безпеки, що розвивався, хоча еволюція брандмауерів, рішень IDS та IPS, можливо, не була настільки добре відома чи розголошена. У 1994 році Checkpoint випустила найвідоміший комерційний міжмережевий екран (хоча технічно не перший) з Firewall-1. Початкові брандмауери могли блокувати трафік на основі порту, протоколу та IP-адреси. Рішення IDS стали популярними як додатковий інструмент у 2000-х роках для вирішення нових типів атак, таких як ін'єкції SQL і міжсайтовий сценарій. За своєю суттю рішення IDS виявляє експлойти проти програм у традиційних локальних мережах. IPS (класично кажучи), навпаки, вбудований у трафік брандмауера, мета якого полягає в тому, щоб фактично зупинити зловмисний трафік.

Поведінкові методи виявлення є важливим компонентом у системах кібербезпеки, які орієнтовані на виявлення аномальних дій і потенційних загроз всередині мережі. Ці методи базуються на спостереженні за поведінкою

користувачів, пристроїв і систем, дозволяючи ідентифікувати небажані або підозрілі активності, які можуть вказувати на присутність кіберінцидентів або зловмисників.

Основна ідея поведінкових методів полягає в аналізі нормальних шаблонів поведінки користувачів і пристроїв у мережі. Коли відбувається зміна в цьому шаблоні — наприклад, підвищення кількості запитів, зміни в часі доступу або доступ до заборонених ресурсів — система може класифікувати таку поведінку як підозрілу або аномальну. Це дозволяє своєчасно реагувати на потенційні загрози, не покладаючись виключно на знання про відомі шкідливі коди або сигнатури атак.

Методи виявлення поведінкових аномалій

Аналіз логів і потоку даних. Поведінковий аналіз часто починається з ретельного моніторингу мережевих логів і потоку даних. Це може включати аналіз пакетів мережевого трафіку, з'ясування взаємодії між різними системами та ідентифікацію потенційних загроз, таких як сканування портів, спроби несанкціонованого доступу або використання застарілих паролів. Системи можуть використовувати методи машинного навчання та штучного інтелекту для обробки великих обсягів даних у реальному часі, допомагаючи виявляти аномалії та реагувати на них швидше.

Поведінковий аналіз користувачів. Зокрема, аналіз поведінки користувачів включає вивчення звичних моделей доступу, використання ресурсів і взаємодії з додатками. Відхилення від цих нормальних моделей може вказувати на несанкціоновану діяльність, таку як неавторизований доступ до конфіденційної інформації або використання неадекватного програмного забезпечення. Системи можуть застосовувати машинне навчання для побудови поведінкових профілів користувачів і раннього виявлення змін, які можуть свідчити про загрозу.

Моніторинг пристроїв. Поведінковий аналіз також включає моніторинг пристроїв у мережі, включаючи сервери, робочі станції, мобільні пристрої та інші кінцеві точки. Системи можуть визначати аномальну активність пристроїв, таку як необґрунтоване зростання завантаження центрального процесора, аномальні

мережеві запити або доступ до неавторизованих ресурсів, що може свідчити про компрометацію пристрою або присутність зловмисного програмного забезпечення.

Поведінкові методи виявлення є важливим компонентом у системах кібербезпеки, які орієнтовані на виявлення аномальних дій і потенційних загроз всередині мережі. Ці методи базуються на спостереженні за поведінкою користувачів, пристроїв і систем, дозволяючи ідентифікувати небажані або підозрілі активності, які можуть вказувати на присутність кіберінцидентів або зловмисників.

Основна ідея поведінкових методів полягає в аналізі нормальних шаблонів поведінки користувачів і пристроїв у мережі. Коли відбувається зміна в цьому шаблоні — наприклад, підвищення кількості запитів, зміни в часі доступу або доступ до заборонених ресурсів — система може класифікувати таку поведінку як підозрілу або аномальну. Це дозволяє своєчасно реагувати на потенційні загрози, не покладаючись виключно на знання про відомі шкідливі коди або сигнатури атак.

У мережевій безпеці виявлення загроз штучного інтелекту зосереджено на моніторингу мережевого трафіку для виявлення незвичайних моделей або аномалій. Використовуючи машинне навчання та аналіз даних, системи штучного інтелекту можуть розпізнавати ознаки злому, витоку даних і зараження шкідливим програмним забезпеченням і надавати сповіщення в режимі реального часу. Це дозволяє групам безпеки швидко запровадити цілеспрямовану тактику реагування на інциденти.

Три широко використовувані підходи для виявлення загроз штучного інтелекту в системах мережевої безпеки:

Виявлення аномалій використовує штучний інтелект для виявлення незвичайної поведінки, яка може сигналізувати про потенційні загрози.

Системи виявлення вторгнень (IDS): моніторинг мережевого трафіку на предмет підозрілих дій

Системи запобігання вторгненням (IPS): тісно співпрацюйте з IDS, щоб блокувати та запобігати виявленням загрозам.

Системи виявлення загроз на основі штучного інтелекту стикаються з упередженням даних і етичними проблемами. Прозорість і постійний моніторинг важливі для забезпечення точності прогнозів і запобігання небажаним наслідкам. Особиста інформація також має бути захищена, і тут на допомогу приходять такі закони, як GDPR . Створюючи систему виявлення загроз ШІ, важливо враховувати захист прав людей на конфіденційність і етичне використання даних.

Дані та алгоритми штучного інтелекту для навчання моделей виявлення загроз штучного інтелекту повинні бути ретельно перевірені, щоб уникнути спотворених результатів. Різноманітні набори даних і безперервна оцінка проти упередженості необхідні для забезпечення справедливості в моделях ШІ та справедливих і точних результатів для різних демографічних груп і сценаріїв.

Майбутнє виявлення загроз за допомогою ШІ багатообіцяюче. Експерти передбачають, що це включатиме вдосконалення технологій глибокого навчання для більш детального розпізнавання образів, інтеграцію квантових обчислень для швидшої обробки даних і підвищення прозорості ШІ для кращого розуміння процесу прийняття рішень.

Ймовірно, це призведе до розробки прогнозової аналітики для проактивних дій команд безпеки, автономних систем реагування на інциденти та покращеної персоналізації. Загалом очікується, що майбутнє штучного інтелекту у виявленні загроз покращить його здатність адаптуватися до нових загроз у постійно мінливому та складному ландшафті загроз.

Використовуючи можливості штучного інтелекту, організації можуть досягти більш проактивного та цілісного підходу до кібербезпеки, зрештою зміцнюючи свій захист від постійно існуючих загроз.

Традиційні методи виявлення на основі сигнатур, хоч і пропонують базовий рівень захисту, явно не вистачають перед обличчям загроз, що розвиваються. Їхнє основне обмеження полягає в їх реактивній природі. Ці методи покладаються на

попередньо визначені сигнатури відомих шкідливих дій, що робить їх нездатними виявляти нові атаки або атаки нульового дня, які раніше не зустрічалися. Крім того, підтримання вичерпної бібліотеки актуальних підписів є трудомістким завданням, а статичність цих підписів дозволяє зловмисникам розробляти методи, які повністю уникають виявлення на основі підписів. Крім того, підходи на основі сигнатур часто створюють велику кількість хибних спрацьовувань, переповнюючи аналітиків безпеки сповіщеннями, які вимагають ручного дослідження, що призводить до втоми від сповіщень і потенційно затримує ідентифікацію справжніх загроз

Машинне навчання (ML) пропонує зміну парадигми у виявленні загроз, дозволяючи проактивну ідентифікацію аномалій і раніше невідомих шаблонів атак. Це досягається завдяки застосуванню алгоритмів, які можуть вивчати історичні дані, що містять мережевий трафік, системні журнали та відомі шкідливі дії.

Методи навчання під наглядом:

Машина підтримки векторів (SVM): ці алгоритми вчаться класифікувати точки даних, визначаючи гіперплощину, яка максимізує запас між різними класами (наприклад, звичайний мережевий трафік проти зловмисного). SVM чудово справляються з обробкою даних великого розміру та можуть бути ефективними для виявлення аномальної мережевої активності на основі історичних моделей законної поведінки мережі.

Випадкові ліси: ця методика ансамблевого навчання поєднує передбачення кількох дерев рішень, що призводить до більш надійної та точної класифікації. модель. Випадкові ліси можна навчати на історичних даних, що містять позначені приклади як звичайних, так і зловмисних дій, що дозволяє їм виявляти відхилення від нормальної поведінки мережі, які можуть означати потенційну загрозу. Використовуючи методи навчання під наглядом, організації можуть створювати моделі ШІ, які можуть ефективно розрізняти між собою законний мережевий трафік і зловмисна діяльність у режимі реального часу.

Підходи до навчання без контролю:

Алгоритми неконтрольованого навчання, на відміну від навчання під наглядом, не потребують позначених даних для навчання. Натомість вони визначають шаблони та зв'язки в немаркованих наборах даних. Це робить їх особливо придатними для виявлення аномалій, де метою є визначення відхилень від очікуваної поведінки.

Коефіцієнт локального викиду (LOF): цей алгоритм визначає точки даних, які суттєво відхиляються від їх локальної щільності. У контексті виявлення загроз LOF можна застосовувати для аналізу моделей мережевого трафіку та визначення аномальних дій, які можуть вказувати на потенційну кібератаку. Неконтрольована природа LOF дозволяє виявляти нові моделі атак, які можуть бути неочевидними в підходах до навчання під контролем.

Архітектури глибокого навчання. Архітектури глибокого навчання (DL) пропонують значний потенціал для вдосконаленого виявлення загроз завдяки своїй здатності обробляти складні та багатовимірні дані. Ці архітектури, натхненні структурою та функціями людського мозку, складаються з кількох шарів штучних нейронів, які можуть вивчати складні шаблони в даних.

Згорткові нейронні мережі (CNN): CNN особливо вправно аналізують зображення та дані мережевого трафіку. Їх можна навчити ідентифікувати шкідливий вміст у потоках мережевого трафіку шляхом розпізнавання шаблонів, характерних для зловмисного програмного забезпечення або інших векторів атак. Наприклад, CNN можуть бути ефективними у виявленні спроб фішингу, аналізуючи візуальні характеристики підозрілих електронних листів.

Повторювані нейронні мережі (RNN): RNN чудово розпізнають часові послідовності в даних. Це робить їх придатними для аналізу журналів мережевого трафіку та виявлення векторів атак, які розгортаються з часом. RNN можна навчити виявляти складні кампанії атак, які передбачають серію взаємопов'язаних кроків, таких як розвідка, проникнення та ексфільтрація даних. Впроваджуючи архітектури глибокого навчання в системи виявлення загроз, організації можуть отримати глибше розуміння моделей мережевого трафіку та визначити навіть найтонші аномалії, які можуть означати потенційну кібератаку.

Generative Adversarial Networks (GAN): Хоча генеративні змагальні мережі (GAN) не використовуються безпосередньо для виявлення загроз, вони можуть відігравати допоміжну роль у збільшенні даних. GAN складаються з двох конкуруючих нейронних мереж: генеративної моделі, яка вчиться створювати синтетичні дані, схожі на реальні дані, і дискримінаційної моделі, яка намагається розрізняти реальні та синтетичні дані. Цей змагальний навчальний процес можна використовувати для створення реалістичних, але нових прикладів зловмисних дій, які потім можна використовувати для збагачення навчальних наборів даних для контрольованих і неконтрольованих алгоритмів навчання. Використовуючи дані, створені GAN, організації можуть підвищити надійність і ефективність своїх систем виявлення загроз на основі ІІІ.

Оцінка ризиків на основі штучного інтелекту. Ландшафт загроз, що постійно розвивається, потребує проактивного підходу до оцінки ризиків, який базується на даних. Традиційні методи оцінки ризиків часто покладаються на ручне сканування вразливостей і суб'єктивні оцінки, що призводить до потенційних неточностей і неефективності. Проте ІІІ пропонує трансформаційний підхід до оцінки ризиків, використовуючи аналітику даних і алгоритми машинного навчання, щоб забезпечити більш повну та об'єктивну оцінку вразливостей системи.

Методи ІІІ для оцінки ризиків: ІІІ дає можливість організаціям проводити оцінку ризиків, орієнтовану на дані, що забезпечує більш детальне розуміння вразливостей безпеки. Ось кілька ключових методів штучного інтелекту, які використовуються в цій області:

Байєсовські мережі: ці ймовірнісні графічні моделі представляють залежності між змінними. У контексті кібербезпеки байєсовські мережі можна використовувати для моделювання взаємозв'язків між компонентами системи, уразливими місцями та потенційними загрозами.

Включаючи історичні дані про минулі інциденти та бази даних експлойтів, ці моделі можуть розрахувати комплексний показник ризиків для кожної вразливості, враховуючи ймовірність її використання та потенційний вплив на

організацію. Цей підхід, що керується даними, дозволяє командам із безпеки визначити пріоритети своїх зусиль і зосередитися на усуненні вразливостей, які становлять найвищий ризик.

Навчання з підкріпленням (RL): цей тип машинного навчання передбачає взаємодію агента із змодельованим середовищем, навчання на його діях і отриманні винагороди або штрафи. Алгоритми RL можна застосовувати для оцінки ризиків кібербезпеки шляхом імітації поведінки зловмисника. Дозволяючи агенту RL взаємодіяти з імітованою версією мережевого середовища організації, алгоритм може навчитися визначати вразливості, які можна використовувати, і визначати пріоритети на основі їхньої простоти використання та потенційної шкоди. Цей проактивний підхід до оцінки ризиків дає змогу командам безпеки усунути критичні вразливості до того, як ними зможуть скористатися реальні зловмисники

III для покращеного реагування на інциденти. Ефективність кібербезпеки організації залежить не лише від виявлення загроз та оцінки ризиків, але й від її здатності швидко й ефективно реагувати на інциденти безпеки. Традиційні процеси IR часто передбачають ручне розслідування та відновлення, що призводить до затримок у стримуванні та потенційної ескалації шкоди. Проте III може значно підвищити ефективність реагування на інциденти шляхом автоматизації рутинних завдань, прискорення розслідування та розширення можливостей прийняття обґрунтованих рішень.

Важливість ефективного реагування на інциденти. Швидко та добре скоординоване реагування на інциденти має першочергове значення для мінімізації впливу порушення безпеки. Кожна мить, коли зловмисник залишається непоміченим у мережевому середовищі, збільшує потенціал для викрадання даних, збою в системі та шкоди репутації. Автоматизуючи повторювані завдання та надаючи статистику в режимі реального часу, штучний інтелект може значно прискорити процес реагування на інциденти, дозволяючи командам безпеки ізолювати загрозу, викорінити присутність зловмисника та швидко відновити нормальну роботу.

1.3 Огляд сучасних рішень для ідентифікації кіберінцидентів

Сучасні кіберзагрози стають дедалі складнішими, і організації змушені покладатися на потужні інструменти для виявлення атак і мінімізації їхнього впливу. Серед таких інструментів особливе місце займають платформи виявлення вторгнень (IDS) та моніторингу мережі, як-от Suricata, Zeek і Snort. Ці системи забезпечують аналіз трафіку, виявлення аномалій і загроз, а також підтримку автоматизації та інтеграцію з методами штучного інтелекту (AI).

Suricata є потужним IDS/IPS/NSM інструментом із підтримкою аналізу як на рівні сигнатур, так і на рівні поведінки. Вона здатна обробляти великий обсяг мережевого трафіку завдяки багатопоточності.

Основні функції:

- Підтримка аналізу HTTP, TLS, DNS, FTP і SSH.
- Інтеграція з популярними базами сигнатур, такими як Emerging Threats.
- Журналювання на основі JSON для глибокого аналізу подій.

Переваги:

- Висока продуктивність.
- Можливість інтеграції з іншими системами (наприклад, Kibana для візуалізації).

Zeek, раніше відомий як Bro, орієнтований на глибокий аналіз мережевих подій, надаючи детальний контекст для виявлення загроз.

Основні функції:

- Аналіз додатків (HTTP, DNS, FTP, SSL).
- Інтеграція з мовами сценаріїв для налаштування правил.
- Підтримка великих розподілених середовищ.

Переваги

- Орієнтація на поведінковий аналіз.
- Гнучкість і розширюваність.

Snort — одна з найпоширеніших сигнатурних систем виявлення вторгнень, яка відома своєю простотою і стабільністю.

Основні функції:

- Аналіз трафіку в реальному часі.
- Виявлення відомих атак за допомогою сигнатур.
- Підтримка обробки пакетів на рівні протоколів.

Переваги:

- Широка спільнота користувачів.
- Багата база сигнатур.

Зі зростанням складності кіберзагроз автоматизація й інтеграція методів штучного інтелекту стають ключовими факторами успіху виявлення атак. Більшість сучасних систем виявлення інтегрують інструменти автоматизації для зменшення ручної роботи.

Наприклад:

Використання скриптів для автоматичного налаштування правил і оновлення баз сигнатур.

Автоматичне реагування на загрози, такі як блокування підозрілих IP-адрес або запуск ізоляції заражених вузлів.

Використання ШІ і машинного навчання значно розширює можливості існуючих платформ:

- Аналіз аномалій. Наприклад, ШІ може виявляти невідомі загрози, аналізуючи трафік, що не відповідає типовим патернам.
- Класифікація загроз. Навчені моделі можуть класифікувати події за рівнем ризику, пріоритизуючи серйозні атаки.
- Прогнозування загроз. Використовуючи алгоритми прогнозування, системи можуть попереджати про потенційні атаки до їхньої реалізації.

Сучасні реалізації, такі як використання Random Forest і нейронних мереж, у поєднанні з Suricata або Zeek, значно підвищують точність виявлення та зменшують кількість хибнопозитивних спрацювань.

Популярні платформи виявлення кіберзагроз, такі як Suricata, Zeek і Snort, мають свої сильні та слабкі сторони. Їхня ефективність залежить від конкретних сценаріїв використання, обсягу мережевого трафіку та вимог до контекстного аналізу. Інтеграція автоматизації й штучного інтелекту робить ці системи ще потужнішими, дозволяючи їм адаптуватися до швидко змінного ландшафту кіберзагроз.

Теоретичні переваги штучного інтелекту в кібербезпеці можна додатково підтвердити шляхом вивчення його практичного застосування в реальних сценаріях. Тематичні дослідження в різних галузях пропонують переконливі докази трансформаційного потенціалу штучного інтелекту в покращенні виявлення загроз, оцінки ризиків і реагування на інциденти. Виявлення загроз на основі штучного інтелекту:

Індустрія фінансових послуг: транснаціональний банк розгорнув систему виявлення загроз на основі штучного інтелекту, використовуючи алгоритми машинного навчання для аналізу моделей мережевого трафіку. Ця система успішно виявила складну фішингову кампанію, спрямовану на облікові записи електронної пошти співробітників. Модель штучного інтелекту, навчена на історичних спробах фішингу, виявила незначні аномалії в мові електронної пошти та поведінці відправника, що дозволило банку швидко втрутитися та запобігти потенційним фінансовим втратам.

Сектор охорони здоров'я: провідний медичний заклад реалізував архітектуру глибокого навчання для аналізу мережевого трафіку. Дана система, оснащена CNN ефективно ідентифікувала шкідливий вміст, вбудований у, здавалося б, законні файли медичних зображень. Модель CNN, навчена на величезному наборі даних зразків зловмисного програмного забезпечення, змогла виявити незначні варіації в шаблонах пікселів зображення, що призвело до своєчасної ізоляції заражених систем і захисту конфіденційних даних пацієнтів.

Впровадження реагування на інциденти за допомогою штучного інтелекту:

Сектор роздрібної торгівлі: велика роздрібна мережа розгорнула систему реагування на інциденти на основі штучного інтелекту, що включає обробку природної мови (NLP) і методи машинного навчання. Під час інциденту безпеки, пов'язаного з атакою програм-вимагачів, компонент NLP системи аналізував нотатки про викуп і системні журнали, вилучаючи ключову інформацію про вимоги зловмисника та масштаби злому. Ця інформація в поєднанні з автоматизованими процедурами стримування на основі машинного навчання сприяла швидкій ізоляції заражених систем і швидшому процесу відновлення.

Телекомунікаційна галузь: телекомунікаційна компанія впровадила чат-бота на основі ШІ, щоб допомагати користувачам під час DDoS атаки. Чат-бот, оснащений можливостями обробки природної мови, надавав інструкції в режимі реального часу користувачам, які зазнали збоїв у роботі служби, направляючи їх через кроки з усунення несправностей і пом'якшуючи потенційну паніку. Це рішення на основі штучного інтелекту не тільки зменшило навантаження на ІТ-персонал підтримки, але й покращило роботу користувачів під час критичного інциденту.

Дані тематичні дослідження ілюструють практичне застосування штучного інтелекту в різних галузях, демонструючи його ефективність у виявленні складних загроз, визначенні пріоритетів вразливостей і сприянні ефективному реагуванню на інциденти. Оскільки ландшафт кібербезпеки продовжує розвиватися, штучний інтелект, безсумнівно, відіграватиме все більш вирішальну роль у зміцненні організаційного захисту та захисті конфіденційної інформації.

Володіння передовими технологічними інструментами може покращити стратегію реагування на інциденти кібербезпеки. Цифрова сітка інструментів мережевого аналізу, виправлення виявлення кінцевих точок і платформ реагування на інциденти забезпечують надійний арсенал для швидкого та стійкого усунення порушень безпеки.

Технологічний ентузіаст рішуче виступить за застосування штучного інтелекту, блокчейну та машинного навчання. ШІ, машинне навчання та інструменти безпеки на основі штучного інтелекту можуть покращити реагування

на інциденти кібербезпеки, мінімізуючи людську взаємодію під час виявлення загроз. Це призводить до швидших і ефективніших відповідей. Штучний інтелект може аналізувати величезні обсяги даних для виявлення аномалій, позначаючи потенційні загрози з більшою точністю, ніж людський аналіз. Методи машинного навчання можна використовувати, щоб навчити системи розпізнавати шаблони, пов'язані зі зловмисними діями, зменшуючи помилкові спрацьовування та підвищуючи гнучкість процесу реагування.

Подібним чином блокчейн пропонує багатообіцяюче застосування для забезпечення цілісності даних, ключової частини реагування на інциденти кібербезпеки. Децентралізована природа блокчейну робить майже неможливим маніпулювання даними зловмисниками, захищаючи процеси відновлення даних і обмежуючи шкоду, завдану зловмисниками. Хмарна кібербезпека також відіграє важливу роль у реагуванні на інциденти. Ця технологія пропонує масштабовані, гнучкі та комплексні рішення, включаючи можливості швидкого виявлення та передову аналітику.

Допомога в пом'якшенні та управлінні етапами реагування на інциденти, досягнення в автоматизації та оркеструванні є незамінними. Інструменти моніторингу системи та мережі виявляють раптові зміни в мережевому трафіку або продуктивності системи в режимі реального часу, забезпечуючи аналітикам постійний огляд середовища, готовий помітити ознаки атаки. Після виявлення потенційної загрози автоматизовані інструменти оркестровки можуть виконувати заздалегідь визначену тактику реагування, забезпечуючи швидке та систематичне стримування та ліквідацію.

Платформи кібербезпеки, такі як SOAR об'єднують дані про загрози, керування інцидентами та інтерактивні інформаційні панелі в одному інтерфейсі. Він забезпечує обізнаність про ситуацію в режимі реального часу, розширене виявлення загроз, спрощені робочі процеси та швидкий час реагування.

Інфільтрувати реагування на інциденти кібербезпеки за допомогою технологій – це не просто розумно, це необхідна умова виживання. У цифровому середовищі, що постійно розвивається, жодне підприємство не може ризикувати

бути технологічно бездіяльним. Вони повинні залишатися попереду розвитку технологій, інакше ризикують стати жертвою злих намірів кіберзлочинців. Обговорювані елементи технології в поєднанні з ключовими компонентами та стратегією, згаданими раніше, пропонують потужний захист і прискорений механізм реагування на потенційні кіберзагрози.

Кібербезпеки – це не лише реакція, це цілодобове спостереження, де технології не лише допомагають реагувати, але й створюють міцну фортецю, у яку важко проникнути кіберзагрозам.

2. АНАЛІЗ МЕТОДІВ ТА ІНСТРУМЕНТІВ ДЛЯ РОЗСЛІДУВАННЯ КІБЕРІНЦИДЕНТІВ

2.1 Характеристика основних методів розслідування кіберінцидентів

Управління реагуванням на інциденти є невід’ємною частиною операцій з кібербезпеки. Спеціалісти з реагування на інциденти першими реагують на будь-який інцидент безпеки: вони допомагають організаціям виявити, локалізувати, викорінити та відновити інцидент. Обробники інцидентів допомагають створити плани управління інцидентами для процедур виявлення та відновлення.

Життєвий цикл реагування на інцидент — це серія процедур, які виконуються у разі інциденту безпеки. Дані кроки визначають робочий процес для загального процесу реагування на інциденти. Кожен етап передбачає певний набір дій, які організація повинна виконати.

Існує кілька способів визначення життєвого циклу реагування на інцидент. Національний інститут стандартів і технологій (NIST; Cichonski та ін., 2012) розробив структуру обробки інцидентів, яка є найбільш поширеною моделлю. Процес, описаний у структурі NIST, включає п'ять етапів:

- Підготовка
- Виявлення та аналіз
- Стимування
- Ерадикація та відновлення
- Діяльність після події

1. Підготовка

На цьому етапі бізнес створює план управління інцидентами, який може виявити інцидент у середовищі організації. Підготовчий етап включає, наприклад, ідентифікацію різних атак зловмисного програмного забезпечення та визначення їхнього впливу на системи. Це також передбачає забезпечення того, щоб організація мала інструменти для реагування на інцидент і відповідні заходи безпеки, щоб запобігти інциденту.

2. Виявлення та аналіз

Аналітик реагування на інциденти відповідає за збір і аналіз даних, щоб знайти будь-які підказки, які допоможуть визначити джерело атаки. На цьому етапі аналітики визначають характер атаки та її вплив на системи. Бізнес і спеціалісти з безпеки, з якими він працює, використовують інструменти та індикатори компрометації (ІОС), які були розроблені для відстеження атакованих систем.

3. Стимування, ліквідація та відновлення

Це основна фаза реагування на інцидент безпеки, під час якої респонденти вживають заходів, щоб зупинити подальшу шкоду. Цей етап складається з трьох кроків:

Стимування. На цьому етапі використовуються всі можливі методи запобігання поширенню шкідливих програм або вірусів. Дії можуть включати відключення систем від мереж, розміщення заражених систем на карантині (Landesman, 2021) або блокування трафіку до та з відомих шкідливих IP-адрес.

Ерадикація. Після виявлення проблеми безпеки зловмисний код або програмне забезпечення необхідно видалити з середовища. Це може включати використання антивірусних інструментів або методів ручного видалення (Williams, 2022). Це також включатиме забезпечення того, щоб усе програмне забезпечення безпеки було оновлено, щоб запобігти будь-яким інцидентам у майбутньому.

Відновлення. Після усунення зловмисного програмного забезпечення важливо відновити всі системи до стану до інциденту (Mazzoli, 2021). Це може включати відновлення даних із резервних копій, перебудову заражених систем і повторне ввімкнення вимкнених облікових записів.

Останньою фазою життєвого циклу реагування на інцидент є виконання посмертного аналізу всього інциденту. Це допомагає організації зрозуміти, як стався інцидент і що вона може зробити, щоб запобігти подібним інцидентам у майбутньому.

Розробка комплексного плану реагування на інциденти має вирішальне значення для того, щоб організації могли ефективно реагувати на кібератаки та відновлюватися після них. Основні елементи успішного плану реагування на кібератаки включають підготовку, виявлення та аналіз, стримування, ліквідацію та відновлення.

Розробка комплексного плану реагування на інциденти має вирішальне значення для того, щоб організації могли ефективно реагувати на кібератаки та відновлюватися після них. Основні елементи успішного плану реагування на кібератаки включають підготовку, виявлення та аналіз, стримування, ліквідацію та відновлення.

Крім того, він має включати чітку місію та цілі, ролі та обов'язки групи реагування на інциденти та документацію щодо підготовки до кіберзагроз.

Створення групи реагування на інцидент передбачає:

- Збирання осіб із відповідним досвідом і набором навичок, що охоплює як технічну (зазвичай частину червоних команд, так і синіх команд) і нетехнічні знання

- Чітке визначення ролей і обов'язків
- Пропонуємо необхідне навчання та ресурси
- Проведення планових навчань і тренувань
- Розробка індивідуальних комунікаційних протоколів для різних зацікавлених сторін

- Впровадження ефективних каналів обміну інформацією

Ці кроки мають вирішальне значення для створення ефективної групи реагування на інциденти.

Крім того, він має включати чітку місію та цілі, ролі та обов'язки групи реагування на інциденти та документацію щодо підготовки до кіберзагроз.

Створення групи реагування на інцидент передбачає:

Збирання осіб із відповідним досвідом і набором навичок, що охоплює як технічну (зазвичай частину червоних команд, так і синіх команд) і нетехнічні знання

- Чітке визначення ролей і обов'язків
- Пропонуємо необхідне навчання та ресурси
- Проведення планових навчань і тренувань
- Розробка індивідуальних комунікаційних протоколів для різних зацікавлених сторін
- Впровадження ефективних каналів обміну інформацією

Ці кроки мають вирішальне значення для створення ефективної групи реагування на інциденти.

Цифрова криміналістика схожа на аналіз ДНК кіберсвіту, відіграючи ключову роль у розслідуванні кіберзлочинів, запобіганні витоку даних і допомагаючи правоохоронним органам у встановленні місцезнаходження злочинців. Він передбачає ідентифікацію, збереження, аналіз і документування цифрових доказів для використання в суді.

Електронні дані, такі як комп'ютерні документи, електронні листи, текстові та миттєві повідомлення, транзакції, зображення та історії Інтернету з пристроїв, залучених до злочину, збираються як частина цифрових доказів. Збереження цих доказів передбачає збереження поточного стану пристрою, належне вимкнення пристрою та дублювання всіх відповідних пристроїв зберігання даних для підтримки цілісності доказів.

Потім докази аналізуються за допомогою методологій цифрової криміналістики, накопичувальних зображень і комплексних інструментів аналізу мережі.

2.2 Порівняння інструментів для аналізу логів та даних кібербезпеки для розслідування кіберінцидентів за допомогою штучного інтелекту

Сучасні системи забезпечення кібербезпеки активно використовують спеціалізовані інструменти для аналізу логів та інших даних, що генеруються мережевими пристроями й додатками. Такі інструменти забезпечують обробку

великих обсягів інформації, що дозволяє виявляти аномалії та потенційні загрози у кіберсередовищі.

Zeek відомий раніше як Bro, є потужним засобом для моніторингу мережових подій та аналізу даних. Його основна функція полягає у глибокому аналізі мережових протоколів, таких як HTTP, DNS та FTP, з подальшим логуванням подій. Zeek дозволяє формувати деталізовані журнали, що можуть слугувати базою для подальшого аналізу та автоматизованої обробки даних.

Використання спеціальних скриптів дозволяє адаптувати систему до потреб конкретного середовища, забезпечуючи високий рівень налаштовуваності. Завдяки своїй архітектурі Zeek виконує роль не лише інструменту для виявлення аномалій, але й платформи для інтеграції з іншими рішеннями, такими як системи управління інформацією та подіями (SIEM).

Logstash — це відкритий інструмент для обробки, перетворення та аналізу логів, який є частиною Elastic Stack. Його основне завдання — агрегувати дані з різних джерел, перетворювати їх у стандартизований формат і передавати в сховище для подальшого аналізу.

Архітектура Logstash базується на принципі модульності: дані проходять через етапи збору (input), обробки (filter) та передачі (output). Це дозволяє застосовувати складні правила для трансформації даних, включаючи парсинг, нормалізацію та маркування. Унікальним аспектом Logstash є його здатність обробляти нестандартні формати даних, що робить його універсальним інструментом для складних середовищ.

Splunk є комерційним рішенням, яке поєднує функціонал збору, зберігання та аналізу логів у масштабованій архітектурі. Його можливості включають індексацію великих обсягів даних у реальному часі, створення візуалізацій та виконання складних запитів для виявлення аномалій.

Унікальною особливістю Splunk є його інтерфейс, орієнтований на зручність використання, та підтримка мови запитів, яка дозволяє виконувати складні аналітичні операції. Splunk часто використовується у великих

організаціях, де потрібна інтеграція даних з різних джерел та їхнє швидке оброблення для прийняття оперативних рішень.

Для аналізу логів та кіберзагроз часто використовуються універсальні бібліотеки та інструменти, що надають засоби для автоматизованої обробки даних. Зокрема, Pandas є бібліотекою для роботи з табличними даними, яка забезпечує інструменти для фільтрації, групування та візуалізації. Scikit-learn, у свою чергу, надає широкий спектр алгоритмів машинного навчання, що дозволяють будувати моделі для класифікації, кластеризації та прогнозування.

ШІ здійснив революцію в галузі кібербезпеки, завдяки своїм можливостям машинного навчання, які дозволяють розширене виявлення та запобігання загрозам. Інструменти ШІ відіграють вирішальну роль у підвищенні безпеки даних, пропонуючи широкий спектр функцій. До них належать алгоритми машинного навчання, автоматизація кібербезпеки, керування вразливістю, автоматизація автентифікації, оцінка ризиків і точне виявлення загроз і аномалій.

Інструменти безпеки штучного інтелекту, такі як Crowdstrike, забезпечують надійний захист за допомогою хмарних функцій AntiVirus наступного покоління (NGAV) і функцій виявлення та EDR. Дані інструменти використовують штучний інтелект для аналізу величезних обсягів даних і виявлення потенційних загроз у режимі реального часу, забезпечуючи швидке й ефективне реагування.

Інший вартий уваги інструмент, Cognito від Vectra, спеціалізується на виявленні та реагуванні на кібератаки в реальному часі. Використовуючи штучний інтелект, Cognito може швидко виявляти та пом'якшувати кіберзагрози, мінімізуючи вплив потенційних зломів на безпеку даних організації.

IBM QRadar Advisor із Watson — це ще один потужний інструмент на базі штучного інтелекту, який допомагає визначати пріоритети високоточних сповіщень і допомагає усунути загрози. Завдяки використанню можливостей штучного інтелекту цей інструмент значно покращує ефективність виявлення загроз, дозволяючи командам безпеки зосередити свої зусилля на найбільш критичних загрозах.

Crowdstrike - антивірус наступного покоління, виявлення та реагування кінцевих точок.

Cognito by Vectra - виявлення та реагування на кібератаки в реальному часі.

IBM QRadar Advisor with Watson - пріоритетні сповіщення та усунення загроз.

Окрім цих інструментів, ШІ також використовується у військовій кібербезпеці через фреймворки, такі як BioHAIFCS. Дана структура використовує виявлення аномалій і зловмисного програмного забезпечення, кероване ШІ, для захисту конфіденційних військових даних від потенційних загроз.

Крім того, DefPloreX, набір інструментів криміналістики кіберзлочинів на базі штучного інтелекту, призначений для широкомасштабних розслідувань. Він використовує можливості ШІ для аналізу величезних обсягів даних, допомагаючи у виявленні та запобіганні кіберзлочинам.

Crowdstrike — це провідний інструмент безпеки штучного інтелекту, який пропонує покращений захист даних за допомогою хмарного антивірусу наступного покоління і можливостей виявлення та реагування на кінцеві точки (EDR). Ця потужна комбінація дозволяє організаціям завчасно виявляти кіберзагрози та реагувати на них, запобігаючи потенційним витокам даних і атакам.

Основні характеристики

- Хмарний антивірус нового покоління (NGAV)
- Виявлення кінцевої точки та відповідь (EDR)

Переваги:

- Проактивне виявлення загроз
- Швидке реагування на інцидент
- Постійні оновлення та масштабованість

Cognito від Vectra — це інструмент на базі штучного інтелекту, який чудово виявляє кібератаки в режимі реального часу та реагує на них, забезпечуючи проактивну підтримку безпеки даних. Завдяки вдосконаленим алгоритмам

машинного навчання Cognito відстежує мережевий трафік, поведінку користувачів і взаємодію пристроїв, щоб виявити потенційні загрози та аномалії.

Завдяки можливостям роботи в режимі реального часу Cognito може швидко виявляти та реагувати на різні кіберзагрози, такі як АРТ, зараження зловмисним програмним забезпеченням і спроби викрадання даних. Використовуючи алгоритми штучного інтелекту, він може виявляти шаблони та відхилення від нормальної поведінки, дозволяючи командам безпеки завчасно реагувати на нові загрози.

Переваги:

- Визначає аномальну мережеву активність, яка вказує на потенційні загрози.
- Виявляє відхилення від нормальної поведінки, що дозволяє швидко реагувати на нові загрози.
- Пріоритезує сповіщення на основі серйозності, дозволяючи командам безпеки зосередитися на критичних загрозах.
- Вмикає автоматичні дії для пом'якшення загроз і скорочення часу відповіді.

Використовуючи ШІ і машинне навчання, Cognito від Vectra пропонує організаціям підвищений рівень безпеки даних. Це дає змогу командам безпеки проактивно виявляти кіберзагрози та реагувати на них у режимі реального часу, захищаючи таким чином цінні активи даних.

IBM QRadar Advisor із Watson — це інструмент на базі ШІ, який покращує безпеку даних, визначаючи пріоритетність високоточних сповіщень і надаючи цінну допомогу в усуненні загроз. Завдяки своїм розширеним можливостям цей інструмент революціонізує спосіб виявлення організаціями загроз кібербезпеці та реагування на них.

Однією з ключових особливостей IBM QRadar Advisor із Watson є його здатність визначати пріоритетність сповіщень на основі їх серйозності та актуальності. Використовуючи алгоритми штучного інтелекту, він аналізує величезні обсяги даних із різноманітних джерел, включаючи мережеві журнали,

поведінку користувачів і канали розвідки про загрози, щоб визначити та виділити найбільш критичні сповіщення, які потребують негайної уваги. Це дає змогу групам безпеки зосередитися на найбільш серйозних загрозах і оперативно реагувати, тим самим скорочуючи час реагування та мінімізуючи вплив потенційних порушень.

Окрім встановлення пріоритетів сповіщень, IBM QRadar Advisor із Watson також пропонує усунення загроз за допомогою ШІ. Він надає аналітикам безпеки дієві рекомендації та покрокові вказівки щодо того, як ефективно реагувати на виявлені загрози. Це не тільки прискорює процес реагування на інциденти, але й допомагає організаціям покращити загальну безпеку, забезпечуючи своєчасне й ефективне реагування на загрози.

Основні функції IBM QRadar Advisor із Watson

- Надає пріоритет сповіщенням високої точності.
- Пропонує аналіз загроз на основі ШІ.
- Надає дієві рекомендації щодо усунення загроз.
- Прискорює процес реагування на інциденти.

Підсумовуючи, IBM QRadar Advisor із Watson — це потужний інструмент на основі ШІ, який значно підвищує безпеку даних. Розставляючи сповіщення за пріоритетністю та сприяючи ефективному виправленню загроз, це дає можливість організаціям проактивно виявляти кіберзагрози та реагувати на них, зрештою захищаючи їхні цінні дані та пом'якшуючи потенційні ризики.

БіоНАІFCS — це передова структура кібербезпеки, яка використовує ШІ для виявлення аномалій і захисту конфіденційних військових даних від кіберзагроз. Завдяки розширеним можливостям виявлення аномалій штучного інтелекту БіоНАІFCS гарантує виявлення та усунення потенційних порушень безпеки в режимі реального часу. Аналізуючи величезні масиви даних і застосовуючи алгоритми машинного навчання, БіоНАІFCS може ідентифікувати закономірності та аномалії, які вказують на зловмисну діяльність або спроби несанкціонованого доступу.

Однією з ключових переваг BioHAIFCS є його здатність адаптуватися та вдосконалювати свої можливості виявлення аномалій з часом. Оскільки він продовжує аналізувати дані та вивчати шаблони, він стає більш точним у визначенні нових і нових загроз, що робить його цінним активом для безпеки військових даних. Забезпечуючи безперервний моніторинг і аналіз, BioHAIFCS дозволяє проактивно виявляти загрози та реагувати на них, даючи військовим організаціям перевагу у виявленні та пом'якшенні кіберзагроз, перш ніж вони можуть завдати значної шкоди.

Крім того, BioHAIFCS пропонує попередження та сповіщення в режимі реального часу, що дозволяє персоналу служби безпеки негайно вжити заходів у разі виявлення потенційної загрози. Це мінімізує час відповіді та забезпечує швидке локалізацію та усунення кіберінцидентів. Поєднуючи виявлення аномалій III зі швидким реагуванням на інциденти, BioHAIFCS допомагає військовим організаціям підтримувати високий рівень безпеки даних і захищати критично важливу інформацію від несанкціонованого доступу або компрометації.

Основні характеристики BioHAIFCS:

Виявлення аномалій: алгоритми на основі штучного інтелекту аналізують дані для виявлення аномальних моделей і поведінки.

Моніторинг у режимі реального часу: постійний моніторинг мережевого трафіку та даних для виявлення потенційних загроз у режимі реального часу.

Швидке реагування на інциденти: миттєві попередження та сповіщення дозволяють швидко реагувати та стримувати кіберінциденти.

Адаптивність машинного навчання: фреймворк вивчає нові шаблони та з часом адаптує свої можливості виявлення аномалій.

Завдяки можливостям виявлення аномалій і моніторингу в режимі реального часу на основі штучного інтелекту BioHAIFCS є ключовим компонентом для забезпечення найвищого рівня безпеки даних для військових операцій. Використовуючи технологію штучного інтелекту для виявлення кіберзагроз і реагування на них, структура дозволяє військовим організаціям бути

на крок попереду потенційних супротивників, захищаючи конфіденційні дані та критично важливі системи.

DefPloreX — це інноваційний набір інструментів на основі штучного інтелекту, який революціонує масштабну криміналістичну експертизу кіберзлочинів, надаючи слідчим розширені можливості захисту даних. Це сучасне рішення поєднує в собі потужність штучного інтелекту та передові криміналістичні методи для покращення процесу розслідування та ефективної боротьби з кіберзагрозами.

За допомогою DefPloreX дослідники можуть аналізувати та отримувати цінну інформацію з величезних обсягів складних даних, що дозволяє їм ідентифікувати закономірності, виявляти аномалії та виявляти приховані зв'язки. Набір інструментів використовує алгоритми машинного навчання для автоматичного виявлення потенційних кіберзлочинців і надання оперативної інформації для проактивного пом'якшення загроз.

Однією з ключових особливостей DefPloreX є його здатність проводити широкомасштабну криміналістику кіберзлочинів, що дозволяє слідчим ефективно обробляти величезні обсяги даних. Використовуючи автоматизацію на основі штучного інтелекту, набір інструментів оптимізує робочий процес розслідування, скорочуючи час і зусилля, необхідні для виявлення важливих доказів і створення серйозних справ проти кіберзлочинців.

Ключові характеристики DefPloreX

- Розширений аналіз даних і можливості візуалізації для комплексних судових розслідувань
- Виявлення загроз у реальному часі та реагування на них для швидкого вирішення нових кіберзагроз
- Інтелектуальні алгоритми виявлення аномалій для виявлення підозрілих дій і потенційних порушень
- Потужна кореляція даних і аналіз зв'язків для виявлення прихованих з'єднань і мереж

- Автоматизоване створення звітів для спрощеного документування та представлення доказів

Використовуючи можливості ІІІ, DefPloreX дозволяє слідчим залишатися на крок попереду кіберзлочинців, забезпечуючи безпеку конфіденційних даних і захищаючи організації від нових кіберзагроз. Цей набір інструментів на основі ІІІ є цінним активом у боротьбі з кіберзлочинністю, надаючи необхідні інструменти та інформацію для захисту даних і підвищення загальної цифрової безпеки.

Однією з ключових переваг використання інструментів ІІІ для захисту даних є їх здатність запобігати несанкціонованому доступу. Ці інструменти аналізують поведінку користувачів, відстежують мережеву діяльність і впроваджують багатофакторну автентифікацію, щоб гарантувати доступ до важливої інформації лише авторизованому персоналу. Постійно навчаючись на шаблонах і аномаліях, інструменти штучного інтелекту можуть виявляти підозрілі дії, такі як несанкціоновані спроби входу або незвичні запити на доступ до даних, і негайно вживати заходів для запобігання можливим порушенням.

Крім того, системи контролю доступу на основі ІІІ можуть адаптуватися до мінливих обставин і динамічно коригувати привілеї користувачів на основі оцінки ризиків. Цей інтелектуальний підхід мінімізує ризик розкриття даних через людські помилки або скомпрометовані облікові дані, зміцнюючи загальну безпеку організацій.

Інструменти штучного інтелекту відрізняються проактивним виявленням загроз і реагуванням на них, пропонуючи додатковий рівень захисту від кібератак. Завдяки здатності аналізувати величезні обсяги даних у режимі реального часу ці інструменти можуть визначати закономірності, аномалії та ознаки зламу, які можуть залишитися непоміченими традиційними рішеннями безпеки.

Використовуючи алгоритми машинного навчання, інструменти ІІІ можуть постійно відстежувати мережевий трафік, дії кінцевих точок і системні журнали, щоб виявляти потенційні загрози, такі як зараження зловмисним програмним забезпеченням, спроби фішингу або внутрішні загрози. Вони також можуть співвідносити різні події безпеки, щоб забезпечити цілісне уявлення про

ландшафт загроз, дозволяючи командам безпеки визначати пріоритети та ефективно реагувати на інциденти.

Joblib — це бібліотека Python, яка надає високорівневий інтерфейс для розпаралелювання циклів і викликів функцій, що полегшує прискорення складних обчислювальних завдань. У контексті моделей ШІ Joblib часто використовується для прискорення таких завдань, як попередня обробка даних, навчання моделі та налаштування гіперпараметрів. Використовуючи можливості розпаралелювання Joblib, моделі ШІ можуть ефективніше обробляти великі набори даних, що призводить до швидшого навчання та підвищення загальної продуктивності. Це особливо важливо для моделей, які потребують великомасштабних обчислень, таких як моделі глибокого навчання.

Порівняно з іншими методами зберігання та завантаження моделей машинного навчання використання Joblib має ряд переваг. Оскільки дані зберігаються як рядки байтів, а не як об'єкти, їх можна швидко й легко зберігати в меншому обсязі, ніж традиційне маринування. Крім того, він автоматично виправляє помилки під час читання або запису файлів, що робить його більш надійним, ніж маринування вручну. І останнє, але не менш важливе: використання Joblib дає змогу зберігати численні ітерації однієї моделі, полегшуючи їх порівняння та визначення найточнішої.

Joblib забезпечує багатопроцесорну роботу на кількох машинах або ядрах на одній машині, що дозволяє програмістам розпаралелювати завдання на багатьох машинах. Це полегшує використання розподілених обчислювальних ресурсів, таких як кластери або графічні процесори, для прискорення процесу навчання моделі.

Joblib може бути корисним у розробці, пропонуючи інструменти для налагодження та профілювання, забезпечуючи повторюваність, прискорюючи процедуру тестування та дозволяючи експериментувати. За допомогою цих функцій ви можете створювати більш суттєві та складні програми з більшою продуктивністю та ефективністю.

Joblib пропонує розробникам фреймворків машинного навчання на основі Python, таких як scikit-learn і TensorFlow, ефективний підхід до миттєвого збереження та завантаження вивчених моделей без необхідності повторювати трудомісткий і дорогий процес навчання з нуля кожного разу, коли вони потрібні.

Набори інструментів SciPy або scikits широко використовуються для машинного навчання. Scikit — це спеціальний набір інструментів, який використовується для певних цілей, таких як машинне навчання або обробка зображень. Scikit-learn і Scikit-image є двома спеціалізованими пакетами, які використовуються для цих цілей. Пакет містить набори зручних алгоритмів, які використовуються для обробки процесів машинного навчання та обробки зображень.

Scikits надзвичайно популярні серед програмістів і розробників програмного забезпечення. Scikit-learn можна навіть вважати одним із стовпів машинного навчання за допомогою Python. можливо використовувати це для створення різноманітних моделей, підготовки та оцінки даних або навіть для аналізу після моделі. Пакет scikit-learn має багато швидких службових функцій, які допомагають його можливостям машинного навчання.

Random forests — це популярний контрольований алгоритм машинного навчання, який може обробляти як завдання регресії, так і класифікації. Нижче наведено деякі з основних характеристик випадкових лісів:

- Random forests призначені для контрольованого машинного навчання, де є позначена цільова змінна.
- Random forests можна використовувати для вирішення задач регресії (числова цільова змінна) і класифікації (категоріальна цільова змінна).
- Random forests — це ансамблевий метод, тобто вони поєднують прогнози з інших моделей.
- Кожна з менших моделей у випадковому лісовому ансамблі є деревом рішень.

У класифікації Random forests кілька дерев рішень створюються з використанням різних випадкових підмножин даних і ознак. Кожне дерево рішень схоже на експерта, який надає свою думку про те, як класифікувати дані. Прогнози робляться шляхом обчислення прогнозу для кожного дерева рішень, а потім береться найпопулярніший результат. (Для регресії в прогнозах замість цього використовується техніка усереднення.)

На діаграмі нижче (Рис.2.1) відображено випадковий ліс із n дерев рішень, і відображено перші 5 разом із їхніми прогнозами («Собака» або «Кіт»). Кожне дерево піддається різній кількості функцій і різному зразку вихідного набору даних, і тому кожне дерево може відрізнятися. Кожне дерево робить передбачення.

Дивлячись на перші 5 дерев, відображено, що 4/5 передбачили, що вибірка була Котом. Зелені кружечки вказують на гіпотетичний шлях, який пройшло дерево, щоб прийняти рішення. Random forests буде підраховувати кількість прогнозів із дерев рішень для Cat для Dog вибиратиме найпопулярніший прогноз.

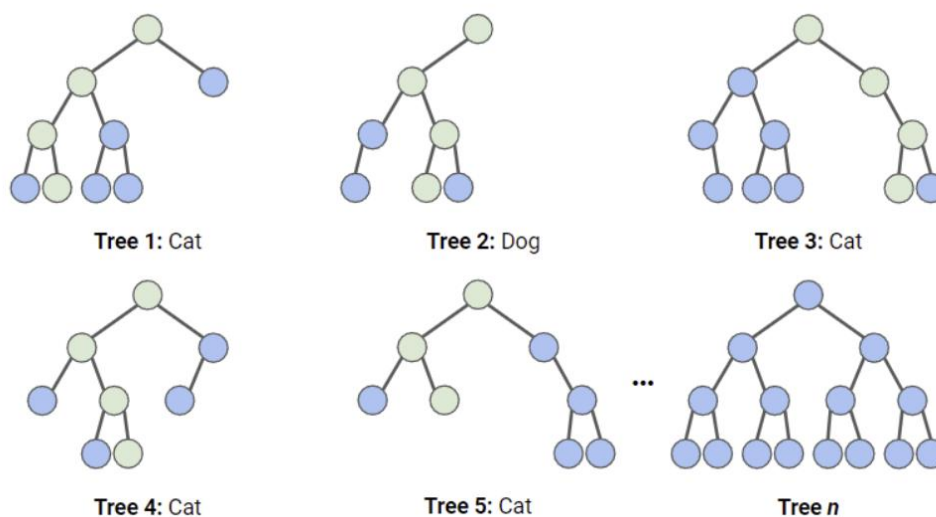


Рис.2.1 Ілюстрація, як працює випадкова класифікація Random forests.

ІІІ значно змінив кібербезпеку, використовуючи машинне навчання для покращення виявлення та запобігання загрозам. Вищезазначені інструменти безпеки ІІІ разом з іншими дають можливість організаціям захищати свої

конфіденційні дані, запобігаючи несанкціонованому доступу, виявляючи кібератаки та реагуючи на них, а також покращуючи загальну безпеку даних.

2.2 Вибір оптимального підходу до автоматизації розслідування за допомогою штучного інтелекту

У динамічному ландшафті цифрової криміналістики інтеграція ШІ і машинного навчання виступає як трансформаційна технологія, готова підвищити ефективність і точність цифрових криміналістичних розслідувань. Однак використання машинного навчання та ШІ в цифровій криміналістиці все ще перебуває на стадії зародження.

За останні роки галузь цифрової криміналістики швидко розширилася, покладаючись на технологію збору та аналізу цифрових доказів під час кримінальних розслідувань. Оскільки використання цифрових доказів у кримінальних розслідуваннях продовжує зростати, існує більша потреба в ефективних і ефективних стратегіях розслідування злочинів. Машинне навчання (ML) і ШІ є двома потужними технологіями, які мають потенціал зробити революцію в цифровій криміналістиці, дозволяючи аналітикам швидко й точно обробляти величезні обсяги даних, виявляючи таким чином важливі докази.

Застосування та впровадження ШІ, наприклад, як методи ШІ можуть бути застосовані у сфері реагування на інциденти і в контексті реагування на інциденти в обмеженому середовищі. Примітно, що використання штучного інтелекту в кримінальних розслідуваннях має важливе значення, особливо з огляду на все більшу поширеність технологій і кіберзлочинності. Численні дослідження показали, що електронні кіберзлочини становлять переважну більшість правопорушень, що підкреслює важливість цифрового рішення.

Незважаючи на те, що бази даних для зберігання вирішених, невирішених і незавершених справ зростають, необхідно підтримувати цю інформацію в Інтернеті задля доступності та безпеки. З цієї причини цілком природно

використовувати програми ШІ та машинного навчання для навчання наборів даних, які слідчі цифрової криміналістики можуть широко використовувати.

У сучасному суспільстві використання ШІ і машинного навчання у кримінальних розслідуваннях стає все більш вирішальним. У цифровій криміналістиці, яка передбачає отримання, обробку та аналіз величезної кількості цифрових даних для кримінальних розслідувань, ШІ та ML виявилися особливо корисними. Розгортання ШІ у цифровій криміналістиці вимагає ретельного розгляду достовірності даних, які аналізуються та обробляються. Однак визначення достовірності даних представляє серйозну проблему, оскільки небагато дослідників поділилися ефективними методами перевірки даних.

В останні роки спостерігається значний інтерес до використання ШІ і машинного навчання (ML) у цифровій криміналістиці. Це пов'язано з тим, що поточні процедури людської експертизи займають багато часу, схильні до помилок і нездатні обробляти величезну кількість криміналістичних даних, які генерують сучасні цифрові пристрої.

Крім того, завдяки автоматизації та оптимізації різноманітних дій, пов'язаних із переглядом цифрових доказів, таких як аналіз даних, обробка зображень і відео та розпізнавання образів, цифрові судові слідчі можуть швидко аналізувати величезні обсяги даних, визначати відповідну інформацію та встановлювати зв'язки, які неможливо помітити звичайними людськими методами.

Шляхом навчання цифрових криміналістичних інструментів розпізнаванню конкретних шаблонів або характеристик, які вказують на певні типи поведінки, алгоритми машинного навчання можуть зменшити кількість помилкових спрацьовувань і підвищити загальну точність. Крім того, деякі алгоритми ШІ та машинного навчання, наприклад, можуть виявляти закономірності та аномалії, які можуть бути не відразу помітні людському оку, що може бути особливо корисним при виявленні прихованих або замаскованих доказів, що призводить до більш точних і надійних результатів.

Автоматизація розслідування кіберзагроз, особливо за участю ІІІ, забезпечує швидший і точніший аналіз інцидентів. Оптимальний підхід до інтеграції таких рішень потребує ретельного планування, вибору інструментів і постійного контролю за їх ефективністю.

ІІІ у реагуванні на інциденти працює за допомогою машинного навчання та автоматизації для виявлення, аналізу та реагування на інциденти безпеки в режимі реального часу. На відміну від традиційних методів, системи ІІІ постійно вивчають дані, щоб підвищити точність і скоротити час відгуку. Ось розбивка того, як ІІІ функціонує в цьому контексті:

Приймання та нормалізація даних: системи ІІІ збирають дані з багатьох джерел, включаючи журнали, мережевий трафік і канали аналізу загроз. Потім ці дані нормалізуються для забезпечення узгодженості для аналізу.

Виявлення аномалій: ІІІ використовує розпізнавання образів і алгоритми машинного навчання для виявлення відхилень від нормальної поведінки, сигналізуючи про потенційні ризики безпеки. Це допомагає в ранньому виявленні інцидентів.

Кореляція подій: ІІІ корелює дані з різних джерел, щоб виявити шаблони, які вказують на складні багатоетапні атаки. Це дозволяє більш точно виявляти загрози.

Автоматизоване сортування інцидентів: ІІІ автоматизує процес сортування, класифікуючи інциденти на основі їх серйозності, зменшуючи навантаження на служби безпеки та забезпечуючи своєчасне реагування.

Аналіз першопричини: ІІІ прискорює виявлення першопричини інцидентів шляхом аналізу даних і точного визначення джерела проблеми. Це скорочує час вирішення та запобігає інцидентам у майбутньому.

Автоматизація реагування: системи ІІІ автоматично виконують заздалегідь визначені дії реагування на інциденти, такі як ізоляція скомпрометованих систем або блокування шкідливих IP-адрес, покращуючи ефективність реагування.

Постійне вдосконалення: моделі ШІ постійно навчаються на основі минулих інцидентів, щоб покращити можливості виявлення та вдосконалити процеси реагування, сприяючи більш проактивним заходам безпеки.

Реагування на інциденти, кероване ШІ, можна застосовувати на різних етапах обробки інцидентів безпеки. Ці випадки використання демонструють, як ШІ покращує процеси виявлення, реагування та навчання в управлінні інцидентами.

1. Виявлення та сповіщення

Системи штучного інтелекту постійно відстежують мережевий трафік, журнали та поведінку користувачів, дозволяючи в режимі реального часу виявляти аномалії та потенційні загрози. Алгоритми машинного навчання виявляють ненормальні шаблони на ранній стадії, запускаючи автоматичні сповіщення, щоб сповістити команди безпеки про можливі інциденти.

2. Аналіз першопричини (RCA)

ШІ прискорює процес аналізу першопричини, швидко аналізуючи дані та корелюючи події між системами. Цей автоматизований підхід допомагає групам безпеки швидше та з більшою точністю визначати основні причини інцидентів, що дозволяє швидше виправляти та запобігати майбутнім інцидентам.

3. Вирішення інцидентів та автоматизація

Інструменти, керовані ШІ, автоматизують вирішення інцидентів, виконуючи заздалегідь визначені дії, такі як ізоляція скомпрометованих систем, блокування зловмисного трафіку або застосування виправлень. Можливість значно скоротити час реагування та звести до мінімуму потребу в людському втручанні під час критичних етапів обробки інцидентів, використовуючи автоматизацію реагування на інциденти за допомогою Blink і Panther.

4. Аналіз і навчання після інциденту

Системи ШІ забезпечують постійний аналіз після інцидентів, дозволяючи командам безпеки вчитися на минулих подіях. Аналізуючи дані про інциденти, штучний інтелект виявляє закономірності та прогалини в поточному захисті,

сприяючи постійному вдосконаленню стратегій реагування на інциденти та допомагаючи організаціям активно зміцнювати свою безпеку.

Реагування на інциденти на основі ШІ пропонує численні переваги, які підвищують ефективність операцій безпеки. Нижче наведено найважливіші переваги.

Покращене виявлення інцидентів і реагування: ШІ значно покращує час виявлення та реагування, аналізуючи великі обсяги даних безпеки в режимі реального часу, виявляючи загрози, які можна пропустити традиційними методами.

Швидший час реагування: системи ШІ автоматизують процеси реагування на інциденти, дозволяючи організаціям миттєво реагувати на інциденти безпеки та мінімізувати потенційну шкоду.

Прискорений аналіз першопричин: ШІ прискорює виявлення першопричини інцидентів безпеки, забезпечуючи швидше їх вирішення та знижуючи ймовірність подібних інцидентів у майбутньому.

Покращена точність: ШІ мінімізує помилкові спрацьовування та негативні результати, аналізуючи моделі та аномалії з високою точністю, гарантуючи, що групи безпеки зосереджуються на реальних загрозах.

Масштабованість: системи на основі ШІ можуть масштабуватися для обробки великих обсягів даних і зростаючої кількості інцидентів, не вимагаючи збільшення людських ресурсів, що робить їх ідеальними для розширення інфраструктури.

Економія коштів і оптимізація ресурсів: шляхом автоматизації завдань і підвищення ефективності ШІ допомагає зменшити операційні витрати, одночасно оптимізуючи використання доступних ресурсів, дозволяючи командам безпеки зосередитися на більш важливих завданнях.

3. ТЕХНОЛОГІЯ РОЗСЛІДУВАННЯ КІБЕРІНЦИДЕНТІВ З ВИКОРИСТАННЯМ МЕТОДІВ ШТУЧНОГО ІНТЕЛЕКТУ

3.1 Порядок виявлення та розслідування кіберінцидентів за допомогою застосування штучного інтелекту

1. Встановлення операційної системи:

- Завантаження та встановлення Ubuntu 22.04 з офіційного сайту.
- Під час встановлення налаштовані базові параметри мережі та встановлено SSH-доступ для роботи через термінал.

2. Налаштування Python-середовища:

- Встановлено Python 3.12 за допомогою команди:

```
sudo apt update
```

```
sudo apt install python3.12 python3.12-venv python3-pip
```



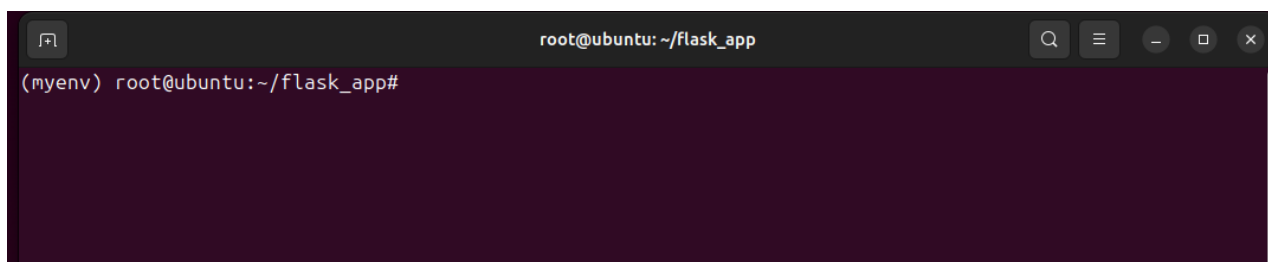
```
root@ubuntu: ~/flask_app
(myenv) root@ubuntu:~/flask_app# sudo apt install python3.12 python3.12-venv python3-pip
Reading package lists... Done
Building dependency tree... Done
Reading state information... Done
python3.12 is already the newest version (3.12.3-1ubuntu0.3).
python3.12 set to manually installed.
python3.12-venv is already the newest version (3.12.3-1ubuntu0.3).
python3-pip is already the newest version (24.0+dfsg-1ubuntu1.1).
0 upgraded, 0 newly installed, 0 to remove and 7 not upgraded.
(myenv) root@ubuntu:~/flask_app#
```

Рис. 3.1. Встановлення пакетів

- Створено та активовано віртуальне середовище:

```
python3 -m venv myenv
```

```
source myenv/bin/activate
```



```
root@ubuntu: ~/flask_app
(myenv) root@ubuntu:~/flask_app#
```

Рис. 3.2. Інтерфейс віртуального середовища

Встановлено необхідні бібліотеки:

```
pip install flask joblib pandas scikit-learn
```

```

root@ubuntu: ~/flask_app
(myenv) root@ubuntu:~/flask_app# pip install flask joblib pandas scikit-learn
Requirement already satisfied: flask in /root/myenv/lib/python3.12/site-packages (3.1.0)
Requirement already satisfied: joblib in /root/myenv/lib/python3.12/site-packages (1.4.2)
Requirement already satisfied: pandas in /root/myenv/lib/python3.12/site-packages (2.2.3)
Requirement already satisfied: scikit-learn in /root/myenv/lib/python3.12/site-packages (1.5.2)
Requirement already satisfied: Werkzeug>=3.1 in /root/myenv/lib/python3.12/site-packages (from flask) (3.1.3)
Requirement already satisfied: Jinja2>=3.1.2 in /root/myenv/lib/python3.12/site-packages (from flask) (3.1.4)
Requirement already satisfied: itsdangerous>=2.2 in /root/myenv/lib/python3.12/site-packages (from flask) (2.2.0)
Requirement already satisfied: click>=8.1.3 in /root/myenv/lib/python3.12/site-packages (from flask) (8.1.7)
Requirement already satisfied: blinker>=1.9 in /root/myenv/lib/python3.12/site-packages (from flask) (1.9.0)
Requirement already satisfied: numpy>=1.26.0 in /root/myenv/lib/python3.12/site-packages (from pandas) (2.1.3)
Requirement already satisfied: python-dateutil>=2.8.2 in /root/myenv/lib/python3.12/site-packages (from pandas) (2.9.0.post0)
Requirement already satisfied: pytz>=2020.1 in /root/myenv/lib/python3.12/site-packages (from pandas) (2024.2)
Requirement already satisfied: tzdata>=2022.7 in /root/myenv/lib/python3.12/site-packages (from pandas) (2024.2)
Requirement already satisfied: scipy>=1.6.0 in /root/myenv/lib/python3.12/site-packages (from scikit-learn) (1.14.1)
Requirement already satisfied: threadpoolctl>=3.1.0 in /root/myenv/lib/python3.12/site-packages (from scikit-learn) (3.5.0)
Requirement already satisfied: MarkupSafe>=2.0 in /root/myenv/lib/python3.12/site-packages (from Jinja2>=3.1.2->flask) (3.0.2)
Requirement already satisfied: six>=1.5 in /root/myenv/lib/python3.12/site-packages (from python-dateutil>=2.8.2->pandas) (1.16.0)
(myenv) root@ubuntu:~/flask_app#

```

Рис. 3.3. Встановлення бібліотек штучного інтелекту та машинного навчання

Flask — це мікрофреймворк для створення веб-застосунків і API на Python, який відомий своєю легкістю і гнучкістю. Flask не нав'язує жорстких структур і дозволяє розробникам налаштовувати рішення під конкретні потреби проєкту.

У нашій реалізації Flask використовується для:

1. Побудови REST API:

Flask обробляє HTTP-запити, приймає дані логів у форматі JSON або CSV, передає їх на обробку та повертає результати аналізу.

2. Маршрутизації:

Створення маршрутів для отримання (GET) та додавання (POST) логів. Наприклад, /logs дозволяє переглянути всі доступні логи, а /add_log — додати нові.

3. Інтеграції з ШІ:

Flask викликає функції, які обробляють дані за допомогою моделей машинного навчання, і повертає результат у вигляді структурованої відповіді.

Цей фреймворк також полегшує інтеграцію з іншими модулями, такими як бібліотеки для роботи з файлами, базами даних і ШІ.

Joblib — це бібліотека Python, призначена для ефективного збереження та завантаження складних Python-об'єктів, зокрема моделей машинного навчання.

Дана реалізація Joblib використовується для:

1. Збереження моделі машинного навчання:

Треновану модель RandomForestClassifier зберігається у файл model.joblib, щоб уникнути повторного навчання при кожному запуску програми.

2. Завантаження моделі під час роботи Flask:

При запуску застосунку Flask файл model.joblib завантажується, і модель стає доступною для обробки нових даних логів.

Переваги використання Joblib включають швидкість і зручність роботи з великими об'єктами Python.

Pandas — це потужна бібліотека для обробки та аналізу даних у табличному форматі. Pandas забезпечує зручну роботу з великими обсягами даних і підтримує різні формати файлів (CSV, Excel тощо).

Дана реалізація Pandas використовується для:

1. Обробки логів:

Дані логів зчитуються у форматі DataFrame, який дозволяє зручно фільтрувати, аналізувати і змінювати структуру даних.

1. Підготовки даних для III:

Логи, які користувач додає через API, конвертуються в DataFrame і обробляються, щоб вони відповідали формату, необхідному для моделі.

2. Збереження та зчитування даних:

Логи зберігаються у файлі log_data.csv для забезпечення їхньої доступності навіть після перезапуску програми.

Pandas є ключовим елементом для роботи з великими даними та підготовки їх до аналізу III.

Scikit-learn — це основна бібліотека для реалізації алгоритмів машинного навчання в Python. Scikit-learn забезпечує широкий спектр інструментів для класифікації, регресії, кластеризації та інших завдань III.

Дана реалізація Scikit-learn використовується для:

1. Тренування моделі машинного навчання:

Використання RandomForestClassifier — алгоритм ансамблевого навчання, який відмінно працює для класифікації даних логів.

2. Аналізу логів у реальному часі:

Нові дані, які надходять через API, обробляються за допомогою завантаженої моделі, яка визначає, чи є логи нормальними або потенційно загрозливими.

1. Оцінки результатів:

Модель забезпечує швидкий і точний аналіз даних, дозволяючи адміністраторам отримувати результати в реальному часі.

Scikit-learn є ядром інтелектуальної обробки даних у цьому проєкті, забезпечуючи високу точність класифікації і можливість інтеграції з іншими модулями Python.

3.1.2 Налаштування системи збору логів

1. Встановлення та налаштування Zeek:

- Встановлено Zeek за допомогою пакетного менеджера:

```
sudo apt update
```

```
sudo apt install zeek
```

- Налаштовано файл конфігурації для зберігання логів у вигляді CSV у папці (myenv) root@ubuntu:/flask_app# :

```
zeekctl deploy
```

- Переконалися, що лог-файли зберігаються у файлі log_data.csv.

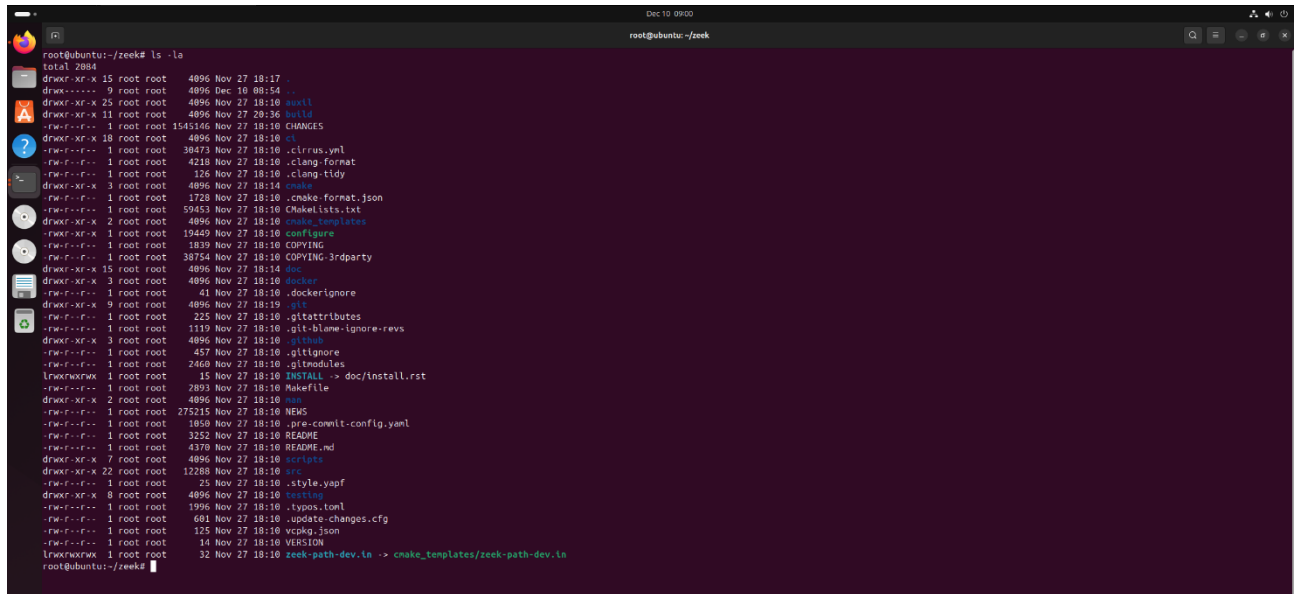


Рис. 3.4. Встановлення Zeek

3.1.3 Розробка веб-застосунку

1. Створення Flask-додатку:

- У папці (myenv) root@ubuntu:/flask_app# створено файл app.py

Файл app.py реалізує основну логіку веб-застосунку.

Імпортуються всі необхідні бібліотеки

Завантажуються моделі:

Маршрути:

/add_log: Приймає POST-запити з новимилогами. Логи аналізуються, результат додається у пам'ять.

/logs: Повертає список усіх проаналізованих логів.

Аналіз логів

Лог отримується через запит і конвертується у формат, придатний для моделі:

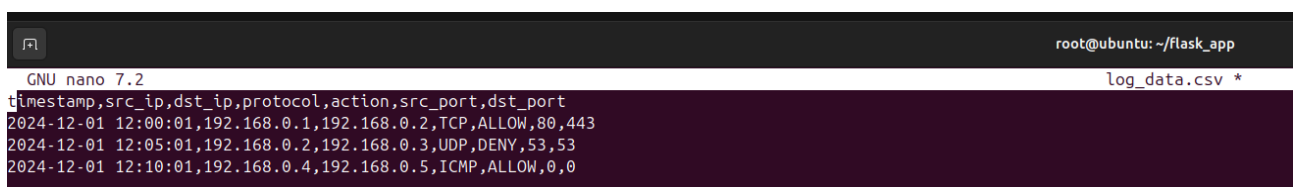


Рис. 3.5. Коректне відображення логування

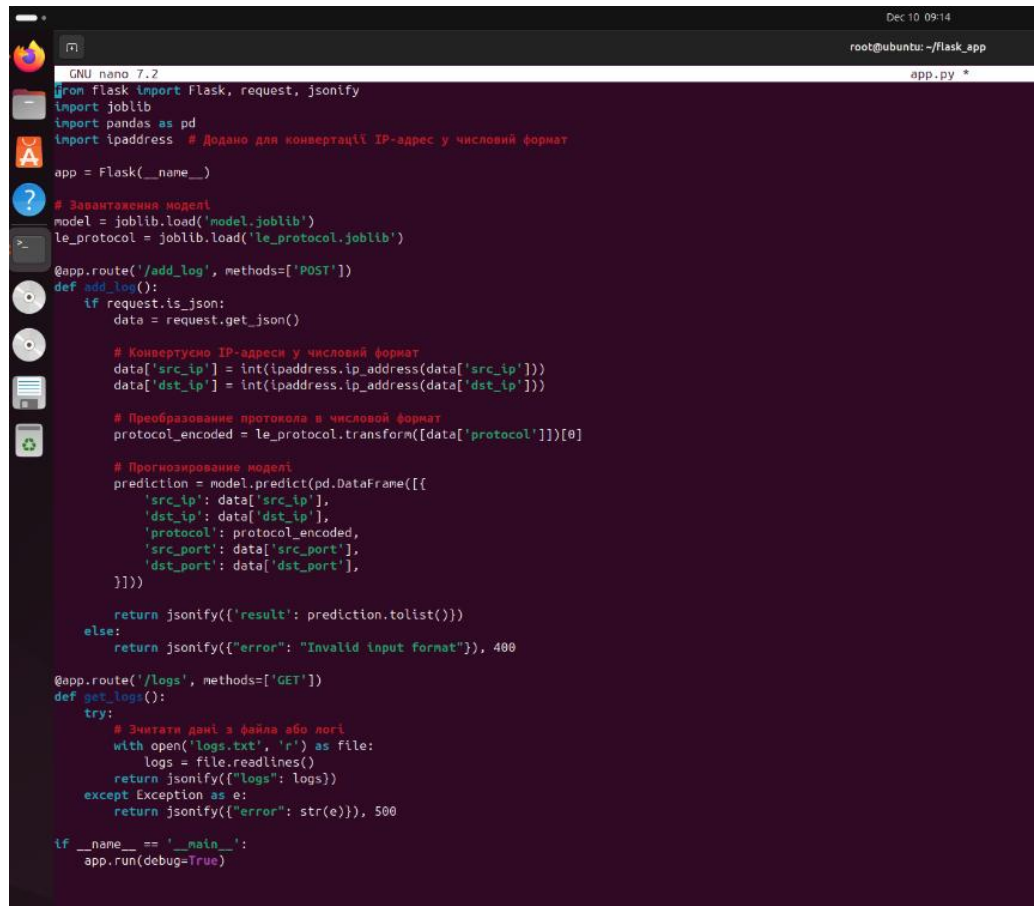
Результат додається у список

Результати:

Усі проаналізовані логи можна переглянути через API:

Відповідно даний скрипт виконує такі функції

- Обробка вхідних запитів для аналізу логів.
- Аналіз отриманих логів за допомогою попередньо натренованої моделі.
- Зберігання та відображення результатів.



```

GNU nano 7.2
from flask import Flask, request, jsonify
import joblib
import pandas as pd
import ipaddress # Додано для конвертації IP-адрес у числовий формат
app = Flask(__name__)

# Завантаження моделі
model = joblib.load('model.joblib')
le_protocol = joblib.load('le_protocol.joblib')

@app.route('/add_log', methods=['POST'])
def add_log():
    if request.is_json:
        data = request.get_json()

        # Конвертуємо IP-адреси у числовий формат
        data['src_ip'] = int(ipaddress.ip_address(data['src_ip']))
        data['dst_ip'] = int(ipaddress.ip_address(data['dst_ip']))

        # Преобразование протокола в числовой формат
        protocol_encoded = le_protocol.transform([data['protocol']])[0]

        # Прогнозирование модели
        prediction = model.predict(pd.DataFrame([[
            'src_ip': data['src_ip'],
            'dst_ip': data['dst_ip'],
            'protocol': protocol_encoded,
            'src_port': data['src_port'],
            'dst_port': data['dst_port'],
        ]]))

        return jsonify({'result': prediction.tolist()})
    else:
        return jsonify({"error": "Invalid input format"}), 400

@app.route('/logs', methods=['GET'])
def get_logs():
    try:
        # Зчитати дані з файла або логів
        with open('logs.txt', 'r') as file:
            logs = file.readlines()
        return jsonify({'logs': logs})
    except Exception as e:
        return jsonify({"error": str(e)}), 500

if __name__ == '__main__':
    app.run(debug=True)
  
```

Рис. 3.6. Основний скрипт системи

2. Інтеграція ШІ-моделі

Завантажено модель RandomForestClassifier, збережену у файлі model.joblib.

Для обробки текстових значень протоколу використано енкодер le_protocol.joblib.

3. Створення REST API:

- Маршрут /add_log: приймає POST-запити з даними логів.
- Маршрут /logs: повертає список оброблених логів у форматі JSON.

Технологія аналізу

1. Підготовка даних:

- Вхідні дані подаються у форматі JSON. Приклад вхідних даних:

json

Copy code

```
{
  "src_ip": "192.168.1.1",
  "dst_ip": "192.168.1.2",
  "protocol": "TCP",
  "src_port": 80,
  "dst_port": 443
}
```

- Значення поля protocol перетворюється на числове за допомогою енкодера le_protocol.joblib.
- 2. Робота моделі штучного інтелекту:
 - Модель RandomForestClassifier аналізує мережеві логи та визначає, чи є трафік безпечним або аномальним.
- 3. Обробка результатів:
 - Результат аналізу додається до масиву logs, який доступний через API /logs.

Візуалізація процесу роботи

У даному розділі представлено результати функціонування розробленого застосунку, що демонструють інтеграцію штучного інтелекту.

1. Запуск застосунку

Веб-застосунок, створений на базі Flask, успішно запускається локально на сервері. Після старту застосунку можна отримати доступ до веб-інтерфейсу за адресою `http://127.0.0.1:5000`.

2. Вивід логів у браузері

При зверненні до маршруту /logs користувач отримує список логів, що зберігаються у файлі `log_data.csv`. Інформація структурована у вигляді таблиці, яка відображає ключові поля для аналізу.

3. Додавання логів через REST API

За допомогою HTTP-запиту POST до маршруту `/add_log` можна додати нові дані.

Після успішної обробки логів результати автоматично зберігаються у систему, а модель машинного навчання визначає їх статус (нормальні чи аномальні).

4. Робота з результатами аналізу

Усі результати, отримані після обробки нових логів, відображаються в браузері. Користувач може переглядати попередні логи разом із висновками моделі штучного інтелекту.

3.2 Етапи демонстрації

1. Відображення роботи app.py

Підтверджує успішний запуск застосунку Flask. У терміналі видно інформацію про активний сервер, маршрути та оброблені HTTP-запити (GET і POST).

2. Відображення веб-інтерфейсу

Показує сторінку /logs, на якій відображаються наявні дані логів, включно з результатами аналізу моделі.

3. Відображення API-запиту

Приклад додавання нового логу через API. Після запиту в браузері оновлюються результати аналізу.

Представлена система демонструє, як сучасні інструменти Python можуть інтегруватися для створення ефективного застосунку з використанням методів штучного інтелекту. Ця реалізація є наочним прикладом автоматизації аналізу даних і дозволяє:

- розслідування кіберінцидентів з використанням методів штучного інтелекту;
- скоротити час на обробку великої кількості логів;
- мінімізувати помилки, пов'язані з людським фактором;
- оперативно реагувати на потенційні загрози в мережевому середовищі;

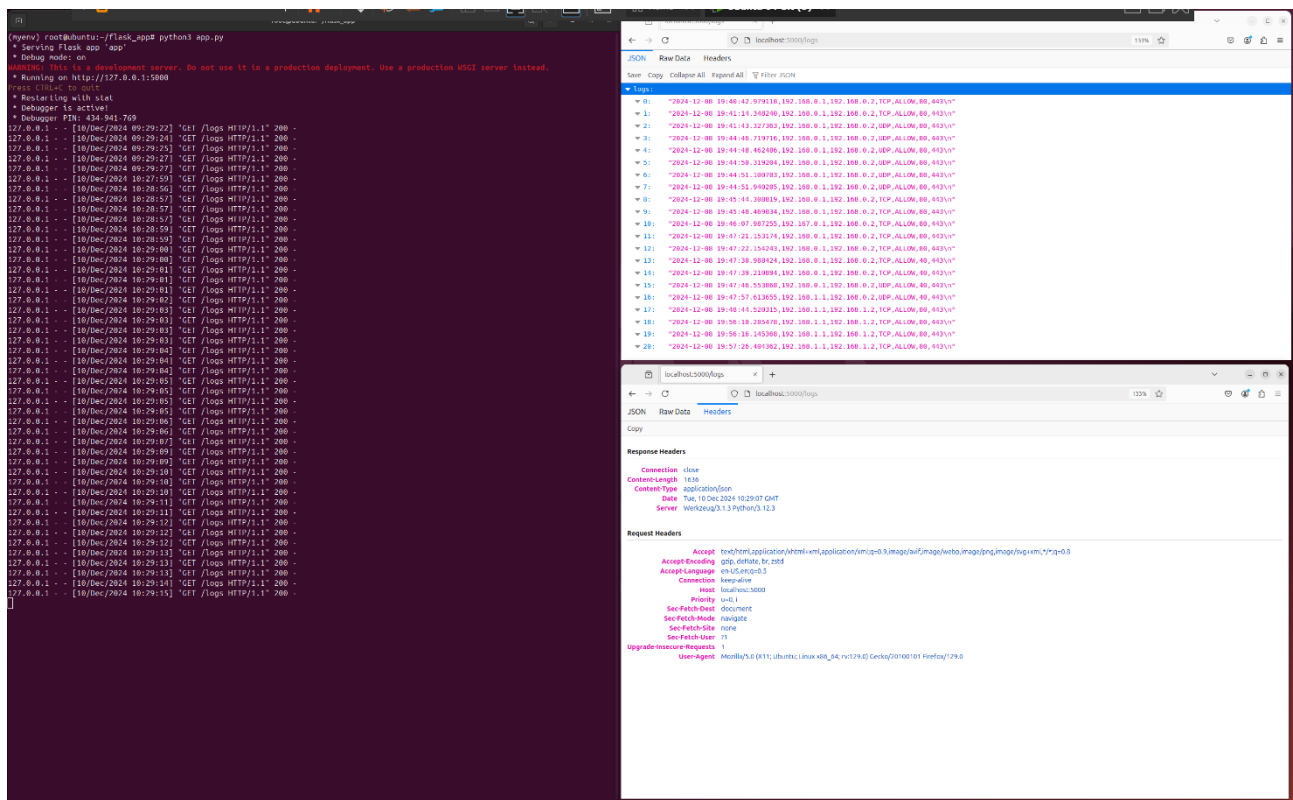


Рис. 3.7. Відображення коректної роботи системи

3.2 Рекомендації щодо виявлення та аналізу загроз для подальшого розслідування кіберінцидентів

Розроблена система аналізу мережевих логів використовує технології ШІ для автоматизації виявлення аномалій у даних для подальшого розслідування кіберінцидентів. Основою роботи є модель машинного навчання, яка інтегрована в веб-додаток на базі Flask. Використання ШІ значно підвищує ефективність роботи адміністраторів мережі та дозволяє забезпечити вищий рівень безпеки. Однак для ефективного використання системи слід врахувати її можливості, обмеження та перспективи розвитку.

1. Переваги реалізації:

Інтелектуальна обробка даних:

Завдяки застосуванню бібліотеки Scikit-learn, система використовує модель машинного навчання для автоматичного класифікування логів. Це дозволяє ідентифікувати аномалії або загрози, які могли б залишитися непоміченими при ручному аналізі.

Автоматизація рутинних процесів:

Система значно скорочує час, необхідний для аналізу великих обсягів логів, завдяки автоматизованому підходу. Використання ІІІ зменшує кількість хибнопозитивних результатів, які часто зустрічаються при використанні традиційних методів.

Постійне вдосконалення моделі:

Модель машинного навчання може донавчатися на нових даних, що підвищує її адаптивність і точність в умовах змінюваних мережових середовищ.

Модульність і масштабованість:

Архітектура системи дозволяє легко інтегрувати нові моделі ІІІ, такі як нейронні мережі, а також масштабувати систему для роботи з великими даними за допомогою хмарних рішень або розподілених обчислень.

Простота інтеграції:

Flask забезпечує простий REST API для взаємодії з іншими системами, що дозволяє легко інтегрувати рішення в існуючу інфраструктуру безпеки.

2. Недоліки реалізації:

Залежність від навчальних даних:

Модель машинного навчання залежить від якості та обсягу навчального набору даних. Якщо лог-файли містять нестандартну інформацію або нові формати, система може некоректно класифікувати події.

Обмежена адаптивність:

На відміну від глибоких нейронних мереж, застосування класичних алгоритмів машинного навчання (наприклад, Random Forest) може бути недостатнім для аналізу дуже складних або нелінійних залежностей у даних.

Ресурсозалежність:

Для обробки великих обсягів даних або високочастотного оновлення логів може знадобитися потужніше обладнання або оптимізація коду.

3. Подальші дії:

Розширення функціоналу:

- Впровадження більш складних моделей ІІІ, таких як нейронні мережі (наприклад, LSTM для аналізу тимчасових даних).

- Додавання функції самонавчання моделі для автоматичного оновлення на основі нових даних.
- Інтеграція з системами оповіщення через месенджери або електронну пошту для миттєвого інформування про критичні події.

Візуалізація результатів:

Розробка дашбордів для графічного відображення логів і результатів аналізу допоможе адміністраторам швидко ідентифікувати аномалії.

Масштабування:

Інтеграція системи з хмарними сервісами, такими як AWS або Google Cloud, дозволить обробляти величезні обсяги даних. Використання баз даних, таких як Elasticsearch, забезпечить швидкий доступ до великих масивів логів.

4. Масштаб застосування:

Система може ефективно працювати з невеликими й середніми обсягами даних (до 1 мільйона записів). Для великих корпоративних мереж, які генерують мільйони логів щодня, знадобиться оптимізація алгоритмів та використання розподілених обчислень.

Реалізована система демонструє, як сучасні технології ШІ можуть використовуватися для аналізу мережевих логів. Її основою є гнучкий та легкий у використанні інструментарій Flask, Pandas, Scikit-learn і Joblib. Вона автоматизує рутинні процеси та забезпечує швидке реагування на потенційні загрози. Розширення функціоналу, інтеграція більш потужних моделей ШІ та масштабування на великі дані можуть значно підвищити її ефективність і розширити сфери застосування.

3.3 Оптимізація та перспективи розвитку автоматизованого аналізу з подальшим розслідуванням кіберінцидентів

Автоматизований аналіз мережевих логів є невід'ємною частиною сучасних систем кібербезпеки, і його ефективність значною мірою залежить від методів і технологій, що використовуються для підвищення точності моделей і масштабованості системи. У цьому розділі розглянемо основні методи

оптимізації, роль хмарних рішень у масштабуванні аналізу та можливості інтеграції з іншими системами кібербезпеки, включаючи застосування ШІ.

Методи підвищення точності моделей

Для підвищення точності моделей машинного навчання, які використовуються для аналізу мережевих логів, необхідно застосовувати різні методи оптимізації. Один з таких методів є гіперпараметрична оптимізація, яка включає підбір оптимальних параметрів моделей з метою досягнення найкращих результатів. Використання алгоритмів, таких як Grid Search або Random Search, дозволяє автоматизувати процес вибору параметрів, що є ключовим для забезпечення високої точності моделей. Крім того, гіперпараметрична оптимізація допомагає уникнути перенавчання моделей, знижуючи ризики неправильної класифікації та збільшуючи загальну ефективність автоматизованого аналізу.

Ансамблеві методи також мають значний вплив на підвищення точності моделей. Використання методів, таких як Random Forest, Gradient Boosting, або XGBoost, дозволяє об'єднати прогнози кількох базових моделей, що покращує загальну точність. Ансамблі знижують ймовірність помилок, пов'язаних з недооцінкою або переоцінкою даних, що є важливим фактором для автоматизованого аналізу мережевих логів.

Роль хмарних рішень у масштабуванні автоматизованого аналізу

Хмарні рішення відіграють ключову роль у масштабуванні автоматизованого аналізу мережевих логів. Використання хмарних платформ, таких як AWS, Microsoft Azure або Google Cloud, дозволяє значно розширити можливості для обробки великих обсягів даних у реальному часі. Ці платформи пропонують гнучке розширення ресурсів, що дозволяє налаштувати обробку даних під конкретні потреби організації, знижуючи витрати на інфраструктуру і підвищуючи ефективність аналізу. Крім того, хмарні рішення забезпечують зберігання та обробку даних, що робить їх ідеальними для аналізу мережевих логів на масштабних рівнях.

ШІ значно підсилює можливості хмарних рішень у масштабуванні автоматизованого аналізу. Використання ШІ алгоритмів, таких як машинне

навчання та глибоке навчання, дозволяє швидко обробляти дані, визначати аномалії та прогнозувати потенційні інциденти. Це не тільки підвищує точність виявлення загроз, але й значно зменшує час реагування на нові типи атак. ШІ також дозволяє автоматизувати багато процесів аналізу, що робить систему більш ефективною та знижує навантаження на операторів.

Можливості інтеграції з іншими системами кібербезпеки

Інтеграція автоматизованого аналізу мережевих логів з іншими системами кібербезпеки створює єдину екосистему, яка забезпечує комплексний захист організації. Дана інтеграція дозволяє обмінюватися інформацією про загрози та автоматично реагувати на виявлені інциденти. Наприклад, інтеграція з SIEM дозволяє централізовано обробляти дані про події та загрози, що покращує загальну картину кібербезпеки.

Інші можливості включають інтеграцію з AI/ML системами, такими як IBM Watson або Google AI Platform, що дозволяє створювати більш просунуті моделі для аналізу та реагування на інциденти. Використання цих систем допомагає підвищити швидкість і точність ідентифікації загроз, а також автоматизувати процеси усунення проблем. Важливий аспект для організацій, які прагнуть до ефективного управління кібербезпекою та знижують витрати на реагування на інциденти.

Крім того, інтеграція з автоматизованими системами попередження та захисту, такими як FireEye, Palo Alto Networks або Cisco Umbrella, створює можливості для оперативного реагування на нові загрози та зловмисні активності. Це дозволяє швидко адаптуватися до змінюваних умов і забезпечити єдину екосистему захисту, яка відповідає на нові виклики у кіберпросторі.

Таким чином, використання різноманітних методів, хмарних рішень для масштабування, а також інтеграції з іншими системами кібербезпеки. ШІ значно посилює ці можливості, створюючи більш інтелектуалізовану та ефективну систему захисту.

ВИСНОВКИ

В роботі досліджено технологію розслідування кіберінцидентів з використанням методів штучного інтелекту, визначено його мету та завдання. Важливість розслідування кіберінцидентів полягає в тому, що з кожним днем організації стикаються з різними загрозами безпеці, які можуть поставити під загрозу цілісність даних, конфіденційність і доступність інформаційних ресурсів організації.

Визначено існуючі підходи до виявлення та розслідування кіберінцидентів. Проаналізовано методи та засоби виявлення та розслідування кіберінцидентів. Відмічено, що методи штучного інтелекту відкривають значні перспективи у сфері кібербезпеки, їх застосування у розслідуванні кіберінцидентів супроводжується низкою суттєвих проблем. Розуміння цих викликів є важливим для розробки надійних та ефективних рішень у сфері кібербезпеки.

Визначено призначення, основні функції і цим самим підкреслює життєздатність і практичну застосовність штучного інтелекту в сфері цифрової криміналістики. Інтеграція штучного інтелекту у процеси розслідування кіберінцидентів не лише оптимізує реагування на загрози, а й дозволяє створити проактивні механізми захисту, здатні попереджати атаки до їх реалізації.

На основі досліджень проведених в роботі запропоновано порядок розслідування кіберінцидентів з використанням методів штучного інтелекту. Розроблено рекомендації фахівцям з кібербезпеки щодо виявлення та розслідування кіберінцидентів з використанням методів штучного інтелекту.

Таким чином, правильна реалізація технології розслідування кіберінцидентів з використанням методів штучного інтелекту має забезпечити ефективний моніторинг, виявлення та аналіз загроз у реальному часі, а також зменшення часу на реагування. Завдяки автоматизації процесів обробки великих обсягів даних та інтеграції інтелектуальних моделей, організації можуть отримувати більш точні прогнози щодо можливих загроз і своєчасно приймати рішення щодо їх нейтралізації.

ПЕРЕЛІК ПОСИЛАНЬ

1. Analysis of cyber threats in the context of the rapid development of information technologies, March 2024, Oleg Haiduk. URL: https://www.researchgate.net/publication/379627207_ANALYSIS_OF_CYBER_THREATS_IN_THE_CONTEXT_OF_RAPID_DEVELOPMENT_OF_INFORMATION_TECHNOLOGY
2. Cyber Incidents Risk Assessments Using Feature Analysis. URL: <https://link.springer.com/article/10.1007/s42979-023-02199-w>
3. AI-Enabled System for Efficient and Effective Cyber Incident Detection and Response in Cloud Environments. URL: <https://arxiv.org/abs/2404.05602>
4. AI for Predictive Cyber Threat Intelligence. URL: <https://www.ijsdcs.com/index.php/IJMESD/article/view/590>
5. Transactions on Latest Trends in Artificial Intelligence. URL: <https://ijsdcs.com/index.php/TLAI/article/view/396>
6. A Comprehensive Review on Cyber-Attacks in Power Systems: Impact Analysis, Detection, and Cyber Security. URL: <https://ieeexplore.ieee.org/abstract/document/10418207>
7. Cyber Security Incident Response: The Effectiveness of Open-Source Detection Tools in DLL Injection Detection. URL: <https://journals.nauss.edu.sa/index.php/JISCR/article/view/2819>
8. An innovative GPT-based open-source intelligence using historical cyber incident reports. URL: <https://journals.nauss.edu.sa/index.php/JISCR/article/view/2819>
9. Classification and Prediction of Significant Cyber Incidents (SCI) Using Data Mining and Machine Learning (DM-ML). URL: <https://ieeexplore.ieee.org/abstract/document/10053846>
10. Enhancing Cyber-Resilience for Small and Medium-Sized Organizations with Prescriptive Malware Analysis, Detection and Response. URL: <https://www.mdpi.com/1424-8220/23/15/6757>
11. Artificial intelligence for cybersecurity: Literature review and future research directions. URL: <https://www.sciencedirect.com/science/article/pii/S1566253523001136>

12. The Ransomware Epidemic: Recent Cybersecurity Incidents Demystified. URL: https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4717250
13. Tactics, Techniques and Procedures of Cybercrime: A Methodology and Tool for Cybercrime Investigation Process. URL: <https://dl.acm.org/doi/abs/10.1145/3600160.3605013>
14. Cyber Incident Response in Public Cloud. URL: <https://www.theseus.fi/handle/10024/803156>
15. Cyber-conflict – Moving from speculation to investigation. URL: <https://journals.sagepub.com/doi/abs/10.1177/00223433231219441>
16. Шевчук Владислав Ігорович. Технологія розслідування кіберінцидентів з використанням методів штучного інтелекту. Всеукраїнська наукова конференція «Актуальні проблеми кібербезпеки». Державний університет інформаційно-комунікаційних технологій. 25 жовтня 2024. Тези доповідей. С. 129-130. URL: https://duikt.edu.ua/uploads/p_2661_48963150.pdf

ДЕМОНСТРАЦІЙНІ МАТЕРІАЛИ (Презентація)

