

Міністерство освіти і науки України
Тернопільський національний технічний університет імені Івана Пулюя
(повне найменування вищого навчального закладу)

Факультет комп'ютерно-інформаційних систем і програмної інженерії
(назва факультету)

Кафедра кібербезпеки
(повна назва кафедри)

КВАЛІФІКАЦІЙНА РОБОТА

на здобуття освітнього ступеня

магістр

(освітній рівень)

на тему: "Розробка системи виявлення вторгнень на основі аналізу
аномалій у мережевому трафіку з використанням машинного
навчання"

Виконав: студент (ка)

Спеціальності:

125 «Кібербезпека та захист інформації»

(шифр і назва напрямку підготовки, спеціальності)

Поліщук Владислав Анатолійович

підпис

(прізвище та ініціали)

Керівник

Стадник М. А.

підпис

(прізвище та ініціали)

Нормоконтроль

Тимошук Д. І.

підпис

(прізвище та ініціали)

Завідувач кафедри

Загородна Н.В.

підпис

(прізвище та ініціали)

Рецензент

підпис

(прізвище та ініціали)

м. Тернопіль – 2024

Міністерство освіти і науки України
Тернопільський національний технічний університет імені Івана Пулюя

Факультет комп'ютерно-інформаційних систем і програмної інженерії

(повна назва факультету)

Кафедра кібербезпеки

(повна назва кафедри)

ЗАТВЕРДЖУЮ

Завідувач кафедри

Загородна Н.В.

(підпис)

(прізвище та ініціали)

«__» _____ 2024 р.

**ЗАВДАННЯ
НА КВАЛІФІКАЦІЙНУ РОБОТУ**

на здобуття освітнього ступеня Магістр

(назва освітнього ступеня)

за спеціальністю 125 Кібербезпека та захист інформації

(шифр і назва спеціальності)

Студенту Поліщуку Владиславу Анатолійовичу

(прізвище, ім'я, по батькові)

1. Тема роботи Розробка системи виявлення вторгнень на основі аналізу аномалій у мережевому трафіку з використанням машинного навчання

Керівник роботи Стадник Марія Андріївна, к.т.н., доцент

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

Затверджені наказом ректора від «20» 11 2024 року № 4/7-1112

2. Термін подання студентом завершеної роботи 20.12.2024

3. Вихідні дані до роботи Наукові публікації про системи виявлення вторгнень та алгоритми машинного навчання з метою виявлення аномалій в мережевому трафіку

4. Зміст роботи (перелік питань, які потрібно розробити)

Вступ 1. Сучасні тенденції мережевих атак та засоби протидії їм 1.1 Аналіз статистичних даних про мережеві атаки 1.2 Сучасні інструменти та технології для протидії мережевим атакам 1.3 Виклики та перспективи у протидії мережевим атакам 2. Системи виявлення вторгнень та застосування машинного навчання 2.1 Поняття про системи виявлення вторгнень 2.2 Огляд існуючих рішень для виявлення вторгнень 2.3 Методи аналізу аномалій у мережевому трафіку 2.4 Використання машинного навчання у виявленні аномалій 3. Розробка та тестування системи виявлення вторгнень 3.1 Вибір інструментів та технологій для побудови системи 3.2 Налаштування та розгортання Snort 3.3 Розробка системи машинного навчання для виявлення аномалій 3.4 Тестування системи шляхом імітації атак Висновки

5. Перелік графічного матеріалу (з точним зазначенням обов'язкових креслень, слайдів)

Вступ, мета, завдання. Наукова новизна та практичне значення дослідження. Поняття про мережеві атаки. Поняття про системи виявлення вторгнень та конкретні приклади існуючих програмних рішень. Розгортання Snort. Розробка моделі машинного навчання для виявлення аномалій. Налаштування кожної з систем на моніторинг трафіку. Тестування системи.

Висновки.

6. Консультанти розділів роботи

Розділ	Прізвище, ініціали та посада консультанта	Підпис, дата	
		завдання видав	завдання прийняв
Охорона праці	Осухівська Г.М., к.т.н., доцент	10.12.2024	
Безпека в надзвичайних ситуаціях	Клепчик В.М., проректор з адміністративно-господарської роботи та будівництва		

7. Дата видачі завдання 23.09.2024 р.

КАЛЕНДАРНИЙ ПЛАН

№ з/п	Назва етапів роботи	Термін виконання етапів роботи	Примітка
1.	Ознайомлення з завданням до кваліфікаційної роботи	23.09-25.09	Виконано
2.	Пошук наукових джерел і статей про імплементацію алгоритмів машинного навчання в системи виявлення вторгнень	26.09 – 12.10	Виконано
3.	Переклад і аналіз наукових джерел і статей про імплементацію машинного навчання в системи виявлення вторгнень	13.10 – 25.10	Виконано
4.	Оформлення розділу “Сучасні тенденції мережевих атак та засоби протидії їм”	26.10 – 08.11	Виконано
5.	Оформлення розділу “Системи виявлення вторгнень та застосування машинного навчання”	09.11 – 14.11	Виконано
6.	Виконання практичної частини щодо розробки системи виявлення вторгнень з імплементацією алгоритмів машинного навчання	15.11 – 24.11	Виконано
7.	Оформлення розділу “Розробка та тестування системи виявлення вторгнень”	25.11 – 29.11	Виконано
8.	Виконання завдання до підрозділу “Охорона праці”	29.11 – 01.12	Виконано
9.	Виконання завдання до підрозділу “Безпека в надзвичайних ситуаціях”	02.12 – 07.12	Виконано
10.	Оформлення кваліфікаційної роботи	08.12 – 11.12	Виконано
11.	Нормоконтроль	12.12 – 13.12	Виконано
12.	Перевірка на плагіат	14.12 – 16.12	Виконано
13.	Попередній захист кваліфікаційної роботи	17.12 – 20.12	Виконано
14.	Захист кваліфікаційної роботи	27.12.2024	

Студент

(підпис)

Поліщук В. А.

(прізвище та ініціали)

Керівник роботи

(підпис)

Стадник М. А.

(прізвище та ініціали)

АНОТАЦІЯ

Розробка системи виявлення вторгнень на основі аналізу аномалій у мережевому трафіку з використанням машинного навчання // ОР «Магістр» // Поліщук Владислав Анатолійович // Тернопільський національний технічний університет імені Івана Пулюя, факультет комп'ютерно-інформаційних систем і програмної інженерії, кафедра кібербезпеки, група СБм-62 // Тернопіль, 2024 // С. 82, рис. – 11, табл. – 1 , кресл. – __, додат. – 3.

Ключові слова: система виявлення вторгнень, Random Forest, аномалія.

Кваліфікаційна робота присвячена дослідженню ефективності використання машинного навчання в контексті аналізу мережевого трафіку на предмет зловмисної активності. Адже існуючі системи виявлення вторгнень, здебільшого, базуються на сигнатурному підході, який потребує періодичного втручання в конфігурацію для впровадження оновлень з метою підтримки рівня ефективності.

В ході дослідження розгорнуто систему виявлення вторгнень Snort з метою порівняння традиційної IDS із системою на основі машинного навчання. Також розроблено систему машинного навчання з використанням алгоритму класифікації Random Forest, яку було навчено на наборі даних CICIDS2017. Кожну з систем було налаштовано на моніторинг мережевого трафіку в локальній мережі в режимі реального часу з метою виявлення аномалій.

Робота спрямована на розробку інструменту для моніторингу мережевого трафіку для мінімізації хибно-позитивних спрацьовувань за рахунок попередньо навченої моделі машинного навчання на наборі даних з прикладами легітимного та зловмисного трафіку.

Отримані результати мають велике значення для аналітиків з кібербезпеки, адже пропонують альтернативний спосіб моніторингу трафіку в мережі, який є більш гнучким та має високу точність спрацьовувань.

ABSTRACT

Development of Intrusion Detection System based on anomaly analysis with Machine Learning algorithms // Thesis of educational level «Master»// Vladyslav Polishchuk // Ternopil Ivan Puluj National Technical University, Faculty of Computer Information Systems and Software Engineering, Department of Cybersecurity, group СБМ-62 // Ternopil, 2024 // P. 82, figs. 11, tables – 1, drawings – __, added. – 3.

Keywords: cyberattack, IDS, Random Forest, anomaly.

The qualifying work is devoted to the study of the effectiveness of using machine learning in the context of analyzing network traffic for malicious activity. Existing intrusion detection systems are mostly based on a signature approach, which requires periodic intervention in the configuration to implement updates in order to maintain the level of efficiency.

During the study, the Snort intrusion detection system was deployed to compare the traditional IDS with a system based on machine learning. A machine learning system was also developed using the Random Forest classification algorithm, which was trained on the CICIDS2017 dataset. Each of the systems was configured to monitor network traffic in the local network in real time in order to detect anomalies.

The work is aimed at developing a tool for monitoring network traffic to minimize false positives due to a pre-trained machine learning model on a dataset with examples of legitimate and malicious traffic.

The results obtained are of great importance to cybersecurity analysts, as they offer an alternative way of monitoring network traffic that is more flexible and has high accuracy of detections.

ЗМІСТ

ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ, СИМВОЛІВ, ОДИНИЦЬ, СКОРОЧЕНЬ І ТЕРМІНІВ	7
ВСТУП.....	8
РОЗДІЛ 1 СУЧАСНІ ТЕНДЕНЦІЇ МЕРЕЖЕВИХ АТАК ТА ЗАСОБИ ПРОТИДІЇ ЇМ	11
1.1 Аналіз статистичних даних про мережеві атаки.....	11
1.2 Сучасні інструменти та технології для протидії мережевим атакам.....	17
1.3 Виклики та перспективи у протидії мережевим атакам	21
РОЗДІЛ 2 СИСТЕМИ ВИЯВЛЕННЯ ВТОРГНЕНЬ ТА ЗАСТОСУВАННЯ МАШИННОГО НАВЧАННЯ	28
2.1 Поняття про системи виявлення вторгнень	28
2.2 Огляд існуючих рішень для виявлення вторгнень	31
2.3 Методи аналізу аномалій у мережевому трафіку	34
2.4 Використання машинного навчання у виявленні аномалій.....	37
РОЗДІЛ 3 РОЗРОБКА ТА ТЕСТУВАННЯ СИСТЕМИ ВИЯВЛЕННЯ ВТОРГНЕНЬ	43
3.1 Вибір інструментів та технологій для побудови системи	43
3.2 Налаштування та розгортання Snort	45
3.3 Розробка системи машинного навчання для виявлення аномалій.....	49
3.4 Тестування системи шляхом імітації атак.....	57
РОЗДІЛ 4 ОХОРОНА ПРАЦІ ТА БЕЗПЕКА В НАДЗВИЧАЙНИХ СИТУАЦІЯХ	61
4.1 Охорона праці.....	61
4.2 Безпека в надзвичайних ситуаціях	64
ВИСНОВКИ.....	69
СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ	71
ДОДАТОК А ПУБЛІКАЦІЯ	74
ДОДАТОК Б ЛІСТИНГ ФАЙЛУ mlIDS_teach.py	77
ДОДАТОК В ЛІСТИНГ ФАЙЛУ anomaly_monitoring.py	81

ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ, СИМВОЛІВ, ОДИНИЦЬ, СКОРОЧЕНЬ І ТЕРМІНІВ

IDS	—	Intrusion Detection System
IPS	—	Intrusion Prevention System
DoS	—	Denial of Service
DDoS	—	Distributed Denial of Service
MITM	—	Man-In-The-Middle
DBSCAN	—	Density-Based Spatial Clustering of Applications with Noise
SVM	—	Support Vector Machine
LSTM	—	Long Short-Term Memory

ВСТУП

У сучасну епоху цифрових технологій мережі відіграють ключову роль у забезпеченні комунікації, обміну даними та функціонуванні глобальних інформаційних систем. Проте зі зростанням залежності від мережевих технологій зростає й рівень загроз, що створюють хакери, злочинні угруповання та навіть державні структури. Серед найбільш поширених атак виділяються DDoS-атаки, атаки типу "людина посередині" (MITM), фішинг, а також експлуатація вразливостей у програмному забезпеченні. У цьому контексті забезпечення мережевої безпеки стає одним із головних викликів інформаційного суспільства.

Актуальність теми обумовлена не лише збільшенням кількості кібератак, але й еволюцією їхньої складності. Зловмисники використовують передові технології, включаючи штучний інтелект, автоматизацію та багатовекторний підхід, щоб обійти традиційні засоби захисту. Такий стан справ вимагає впровадження новітніх систем виявлення вторгнень (Intrusion Detection Systems, IDS), які здатні виявляти не лише відомі загрози, але й аномальні поведінкові патерни, які можуть свідчити про невідомі атаки.

Метою даної роботи є розробка ефективної системи виявлення вторгнень, що базується на аналізі аномалій у мережевому трафіку з використанням методів машинного навчання. Для досягнення цієї мети передбачається провести глибокий аналіз сучасних тенденцій мережевих атак, існуючих засобів протидії, а також визначити виклики та перспективи у цій сфері.

Об'єктом дослідження є мережевий трафік, що формується у сучасних інформаційних системах, зокрема під час функціонування корпоративних та публічних мереж. Предметом дослідження є методи та алгоритми виявлення аномалій у мережевому трафіку з метою ідентифікації потенційних загроз.

Для реалізації мети роботи необхідно вирішити такі завдання:

- Проаналізувати статистичні дані про мережеві атаки та виявити ключові тренди у їхньому розвитку.

- Дослідити сучасні інструменти та технології для протидії мережевим атакам, а також їхні переваги та недоліки.
- Визначити виклики та перспективи у сфері виявлення вторгнень і боротьби з кіберзагрозами.
- Оцінити можливості використання методів машинного навчання для покращення систем виявлення вторгнень.
- Розробити та протестувати систему виявлення вторгнень із використанням аналізу аномалій у мережевому трафіку.

Наукова новизна роботи полягає у поєднанні методів машинного навчання з традиційними підходами до виявлення аномалій, яке досі не зустрічалося у наукових працях в контексті аналізу мережевого трафіку на предмет зловмисної активності. Більшість існуючих IDS орієнтовані на сигнатурний аналіз, що робить їх уразливими до нових загроз. Натомість інтеграція алгоритмів класифікації дозволяє створювати більш адаптивні системи, які здатні ефективно ідентифікувати невідомі атаки на основі аналізу поведінки, для цього у дослідженні наведено порівняння ефективності захоплення підозрілих подій однією із найпоширеніших систем виявлення вторгнень та новоствореною системою машинного навчання. Це має на меті наочно довести, що моделі машинного навчання можуть бути не менш дієвими у виконанні завдань систем виявлення вторгнень, а також можуть виступати як паралельно-працююча система, що дає змогу аналізувати інциденти з різних перспектив.

Проведене дослідження також акцентує увагу на необхідності впровадження багаторівневої архітектури безпеки, що включає збір даних, їхній попередній аналіз, класифікацію загроз та автоматичне реагування. Такий підхід дозволяє не лише зменшити ймовірність успішних атак, але й скоротити час реагування на інциденти.

Окрім цього, розробка нової системи виявлення вторгнень орієнтована на потреби малих і середніх підприємств (SMB), які найчастіше стають жертвами атак через обмеженість ресурсів. Запропонована модель забезпечує баланс між ефективністю та доступністю, що робить її придатною для широкого впровадження.

Результати роботи можуть бути використані для вдосконалення систем кібербезпеки як у корпоративних, так і у державних структурах. Впровадження запропонованого підходу дозволить знизити ризик втрати даних, забезпечити стабільну роботу мереж та підвищити загальний рівень інформаційної безпеки.

Таким чином, ця робота робить внесок у розвиток сучасних систем виявлення вторгнень, демонструючи ефективність використання методів машинного навчання у поєднанні з аналізом аномалій. Подальші дослідження можуть бути спрямовані на оптимізацію моделей та інтеграцію їх із хмарними платформами для забезпечення масштабованого захисту.

Апробація результатів магістерської роботи. Основні елементи роботи, зокрема мета, актуальність, об'єкт та завдання дослідження були презентовані на: XII науково-технічна конференції «Інформаційні моделі, системи та технології».

Публікації. Основні елементи роботи викладені у працях конференції (див. додаток А).

РОЗДІЛ 1 СУЧАСНІ ТЕНДЕНЦІЇ МЕРЕЖЕВИХ АТАК ТА ЗАСОБИ ПРОТИДІЇ ЇМ

1.1 Аналіз статистичних даних про мережеві атаки

У сучасних умовах стрімкого зростання обсягів даних та інтенсивності кібератак важливість систем виявлення вторгнень (IDS) стає дедалі очевиднішою. За даними звіту IBM "Cost of a Data Breach Report 2024", середня вартість витоку даних становить 4,88 мільйона доларів США, що на 10% більше порівняно з попереднім роком (рис. 1.1). З огляду на це, запровадження ефективних рішень для виявлення вторгнень у мережевому трафіку є не просто рекомендацією, а необхідністю для організацій, які прагнуть захистити свої інформаційні ресурси.[1]

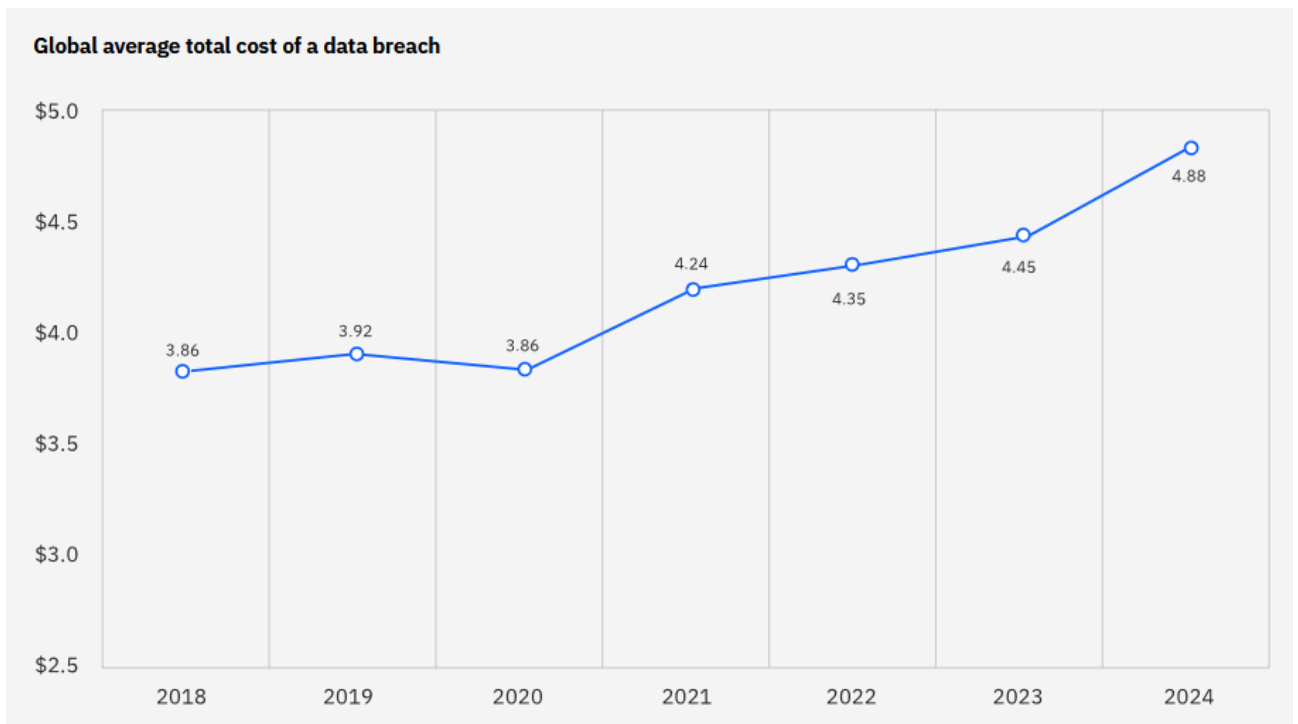


Рисунок 1.1 – Тенденція зростання середньої вартості витоку даних у проміжку останніх шести років

Згідно зі звітом від Fortinet перші дві сходинки у списку найбільш поширених кібер-атак займають DoS/DDoS та MITM. Атака типу «відмова в

обслуговуванні» (Denial-of-Service, DoS) спрямована на перевантаження ресурсів системи до такої міри, що вона не може відповідати на легітимні запити користувачів. Distributed Denial-of-Service (DDoS) має схожу мету – виснажити ресурси системи. DDoS-атака ініціюється за допомогою великої кількості заражених шкідливим програмним забезпеченням пристроїв, які перебувають під контролем зловмисника. Такі атаки називають «відмовою в обслуговуванні», оскільки сайт жертви не може надавати послуги тим, хто намагається отримати до нього доступ.[4]

Під час DoS-атаки цільовий сайт отримує величезну кількість нелегітимних запитів. Оскільки система має реагувати на кожен із них, її ресурси витрачаються на ці відповіді. У результаті сайт не може обслуговувати користувачів у звичайному режимі, що часто призводить до повного припинення його роботи.

DoS та DDoS-атаки відрізняються від інших типів кібератак, які дозволяють хакеру отримати доступ до системи або розширити вже існуючі повноваження. У таких випадках зловмисник отримує прямий зиск від своїх дій. Натомість метою DoS та DDoS-атак є порушення роботи цільового сервісу. Якщо зловмисник найнятий конкурентом, фінансова вигода від атаки може бути опосередкованою.

DoS-атака також може бути використана для створення вразливостей для інших типів атак. Успішна DoS або DDoS-атака змушує систему тимчасово припинити роботу, що робить її вразливою до подальших вторгнень. Одним із поширених способів запобігання таким атакам є використання брандмауера, який визначає, чи є запити до сайту легітимними. Підозрілі запити відсіюються, дозволяючи звичайному трафіку проходити без перешкод. Прикладом масштабної інтернет-атаки такого типу є атака на Amazon Web Services (AWS), що сталася у лютому 2020 року.

Отже DoS та DDoS-атаки залишаються одними з найбільш поширених і руйнівних загроз для компаній, які володіють веб-ресурсами. Їхня популярність серед зловмисників обумовлена відносною простотою реалізації та значним впливом на бізнес-операції. Успішна атака такого типу може не лише тимчасово

вивести з ладу веб-сайт або інфраструктуру, але й спричинити фінансові втрати, репутаційні збитки та втрату довіри клієнтів. Зокрема, підприємства, які надають електронну комерцію, фінансові послуги або працюють у галузі розваг, особливо вразливі до таких атак через їхню залежність від безперервного доступу до веб-ресурсів.

Враховуючи ці загрози, впровадження надійних засобів захисту від DoS та DDoS-атак є обов'язковою умовою для сучасного бізнесу. Передові рішення, такі як брандмауери веб-застосунків (WAF), розподілені системи захисту (CDN) та інтегровані засоби виявлення аномалій у трафіку, дозволяють ефективно ідентифікувати та блокувати підозрілі запити. Крім того, хмарні платформи захисту, які можуть автоматично масштабуватися, забезпечують додатковий рівень безпеки для організацій із великими обсягами трафіку. Таким чином, інвестиції в ефективний захист не лише мінімізують ризики, але й гарантують стабільну роботу веб-ресурсів навіть під час спроб атак.

Говорячи про атаку Man-In-The-Middle (MITM), вона однією з поширених форм кіберзагроз, яка спрямована на перехоплення та маніпуляцію даними між двома сторонами, що вважають, що вони спілкуються безпосередньо. У цій атаці зловмисник виступає як посередник, отримуючи доступ до переданої інформації або змінюючи її, залишаючись непоміченим для обох сторін. Це створює значні ризики для конфіденційності, цілісності та автентичності даних, що передаються у мережі.

MITM-атаки часто використовуються для крадіжки конфіденційної інформації, такої як паролі, банківські дані або особисті дані користувачів. Для успішної реалізації такої атаки зловмисник може використовувати кілька методів, зокрема підробку ARP (ARP Spoofing), перехоплення трафіку через відкриті точки Wi-Fi, DNS-спуфінг або використання шкідливих сертифікатів SSL. Наприклад, під час атаки через Wi-Fi хакер може створити підроблену точку доступу, до якої підключаються жертви, надаючи йому можливість переглядати та змінювати трафік, що проходить через неї.[5]

Особливу небезпеку MITM-атаки становлять для незашифрованого трафіку, оскільки передані дані легко читаються та аналізуються. Проте навіть

зашифровані з'єднання можуть бути вразливими, якщо зломисник вдається до використання фальшивих сертифікатів SSL або експлуатує вразливості протоколу HTTPS. Відомим прикладом є атака типу SSL Strip, під час якої зломисник примусово переводить зашифроване з'єднання на незашифроване, залишаючи користувача без належного захисту.

Для захисту від MITM-атак рекомендується використовувати кілька рівнів захисту. По-перше, шифрування трафіку за допомогою HTTPS і VPN значно ускладнює перехоплення даних. По-друге, автентифікація користувачів і серверів через надійні сертифікати SSL/TLS допомагає уникнути спроб підробки. По-третє, використання двофакторної автентифікації (2FA) додає додатковий рівень безпеки, навіть якщо зломиснику вдалося отримати доступ до облікових даних. Нарешті, усвідомленість користувачів щодо ризиків використання відкритих Wi-Fi мереж та перевірка автентичності веб-ресурсів також знижують ризики успішної атаки.

Таким чином, атака MITM є серйозною загрозою, яка потребує комплексного підходу до захисту, включаючи як технічні засоби, так і підвищення обізнаності користувачів. З огляду на постійний розвиток методів атак, важливо забезпечувати актуальність захисних механізмів і дотримуватися найкращих практик кібербезпеки.

Однією з гучних реалізацій атаки «людина посередині» (MITM) стала спроба викрадення автомобілів Tesla за допомогою фішингу. Дослідники кібербезпеки виявили вразливість, яка дозволяє зломисникам перехоплювати з'єднання між додатком Tesla та серверами компанії, що використовуються для управління автомобілем. Застосовуючи фішинг і MITM, зломисники можуть отримати контроль над транспортним засобом, зокрема розблокувати двері, запустити двигун і навіть викрасти автомобіль.

Вразливість полягала у перехопленні авторизаційного токена, який генерується під час входу користувача до мобільного додатка Tesla (рис. 1.2). Для реалізації атаки зломисник створює фальшиву Wi-Fi-мережу, до якої підключається жертва. Через цю мережу зломисник здійснює MITM-атаку, спрямовану на перехоплення з'єднання між додатком Tesla і його сервером. У

ході атаки користувач, нічого не підозрюючи, вводить свої облікові дані або виконує дії у додатку, що призводить до передачі авторизаційного токена зловмиснику.

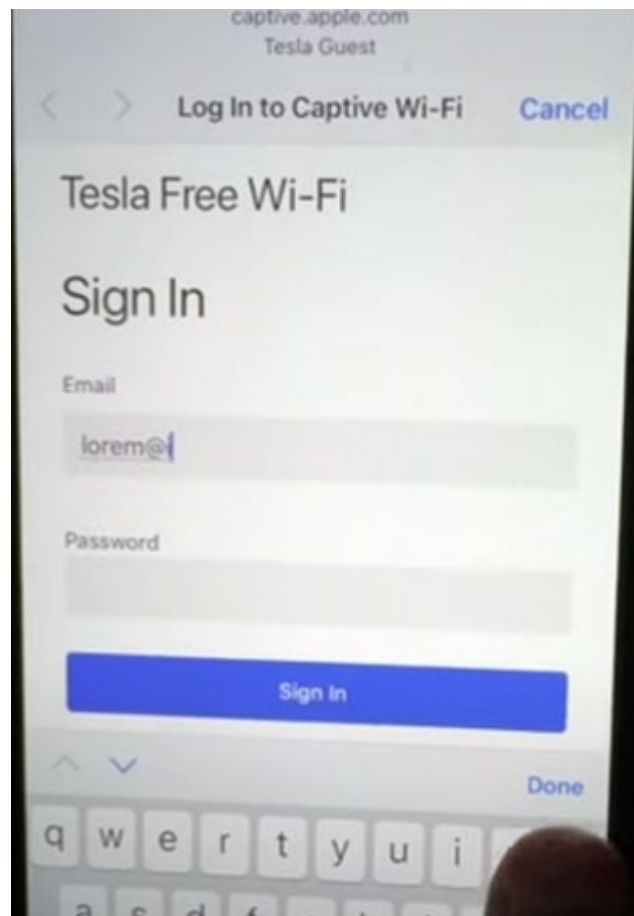


Рисунок 1.2 – Приклад вводу даних для входу в обліковий запис Tesla жертвою

Як тільки зловмисник отримує цей токен, він може використати його для автентифікації в системі Tesla від імені жертви. Це дозволяє йому віддалено виконувати різні дії, такі як розблокування дверей, доступ до функцій управління транспортним засобом або навіть запуск двигуна. Вразливість стає особливо критичною через те, що власники Tesla значною мірою покладаються на мобільний додаток для доступу до свого автомобіля, а безпека з'єднання є ключовим фактором у захисті їхніх даних.

Для захисту від подібних атак компанія Tesla рекомендує своїм користувачам вживати заходів, які знижують ризик MITM-атак. Серед них — уникання підключення до публічних Wi-Fi-мереж без захищеного з'єднання та

використання VPN для шифрування трафіку. Також важливо бути обачними щодо фішингових сайтів і регулярно оновлювати програмне забезпечення, щоб скористатися останніми патчами безпеки.

Цей інцидент демонструє, наскільки важливими є безпека мобільних додатків і шифрування даних у сучасних автомобілях, що значною мірою покладаються на технології підключення до інтернету. Він також нагадує про необхідність постійного вдосконалення захисту від атак МІТМ та підвищення обізнаності користувачів щодо кіберзагроз.

Важливою деталлю цієї атаки була частина перехоплення даних входу користувача, яка отримувалась потенційними зловмисниками за допомогою «фішингу». Тож для комплексного розуміння, розглянутого випадку, варто сформулювати уявлення і про цей вид кібератак.

Фішинг є однією з найпоширеніших і найефективніших форм кібератак, спрямованих на крадіжку конфіденційної інформації, такої як паролі, номери кредитних карток або особисті дані. Зловмисники використовують соціальну інженерію, щоб обманом змусити жертв поділитися своїми даними. Найчастіше атака реалізується через підроблені електронні листи, повідомлення або веб-сайти, які маскуються під легітимні компанії чи організації.

Типовий приклад фішингової атаки — підроблений електронний лист, який виглядає так, ніби надісланий вашим банком або популярною платформою, наприклад, PayPal чи Amazon. У повідомленні вас можуть попросити підтвердити інформацію про акаунт, змінити пароль або оновити платіжні дані. Посилання в листі веде на підроблений сайт, який майже повністю імітує справжній. Коли користувач вводить свої дані, вони передаються зловмисникам.

Зокрема, дослідники компанії IBM звертають увагу на зростання складності фішингових атак. Сучасні фішингові кампанії часто адаптовані під конкретні цільові аудиторії (spear-фішинг) або використовують автоматизацію для масштабування. Зловмисники створюють сайти, які є настільки реалістичними, що навіть досвідчені користувачі не завжди можуть їх відрізнити. Крім того, атаки часто експлуатують поточні події, наприклад, пандемію COVID-19, економічні кризи або великі знижки під час розпродажів.

Для захисту від фішингу рекомендується впроваджувати багаторівневі заходи безпеки. По-перше, регулярне навчання співробітників компаній щодо розпізнавання фішингових повідомлень допомагає мінімізувати ризики. По-друге, використання двофакторної автентифікації (2FA) забезпечує додатковий захист, навіть якщо облікові дані були скомпрометовані. Крім того, інструменти для фільтрації електронної пошти та системи виявлення загроз, такі як SIEM, можуть автоматично ідентифікувати підозрілі повідомлення та блокувати їх.

Цей приклад підкреслює важливість постійної обачності та впровадження технологій, які можуть запобігти успішним фішинговим атакам. Фішинг залишається значною загрозою для користувачів і організацій, але розуміння методів зловмисників та використання сучасних засобів захисту можуть суттєво знизити їхній вплив.

1.2 Сучасні інструменти та технології для протидії мережевим атакам

У сучасних умовах, коли мережеві атаки стають дедалі складнішими та багатогранними, використання єдиного інструменту захисту є недостатнім. Потреба в комплексному підході, що включає різні типи засобів безпеки, зумовлена необхідністю забезпечення надійного захисту від широкого спектра загроз, таких як витоки даних, шкідливе програмне забезпечення, спроби отримання несанкціонованого доступу до мережі та експлуатація вразливостей. Розуміння особливостей і переваг кожного типу інструментів є ключовим для створення багаторівневої системи безпеки, яка дозволить мінімізувати ризики та забезпечити стабільну роботу інформаційних систем.

Витоки даних є однією з найбільших загроз для малих і середніх бизнесів, оскільки вони можуть призвести до фінансових втрат, втрати довіри клієнтів та юридичних наслідків. Інструменти запобігання витокам даних (Data Loss Prevention, DLP) допомагають компаніям контролювати та захищати конфіденційну інформацію, як-от фінансові записи, особисті дані клієнтів чи інтелектуальну власність. Вони забезпечують моніторинг, виявлення і блокування спроб передавання таких даних за межі організації.[6]

Прикладом є рішення Symantec DLP, яке пропонує функції моніторингу електронної пошти, виявлення аномальної активності у файлах і контроль доступу до хмарних сервісів. Інший приклад — McAfee Total Protection for DLP, що дозволяє бізнесам відстежувати всі шляхи передачі даних, включаючи USB-пристрої, веб-трафік і хмарні сервіси. Такі рішення можуть автоматично блокувати передачу конфіденційної інформації на підозрілі сервери

Важливою функцією інструментів DLP є здатність визначати, які дані є найбільш критичними для бізнесу, та забезпечувати для них підвищений рівень захисту. Наприклад, платформа Forcero DLP надає організаціям можливість налаштовувати політики захисту даних відповідно до регуляторних вимог, таких як GDPR або HIPAA. Це особливо корисно для медичних установ або фінансових компаній.

Крім того, інструменти DLP інтегруються з іншими системами безпеки, як-от SIEM, для кореляції подій та запобігання витокам у режимі реального часу. Завдяки цьому компанії можуть швидше реагувати на загрози та запобігати потенційним інцидентам до їх реалізації. Таким чином, використання DLP є основою для забезпечення конфіденційності даних і захисту бізнес-операцій.

Антивірусні рішення є базовим компонентом кібербезпеки для корпоративних мереж, які допомагають захищати робочі станції, сервери та мережі від шкідливих програм. Вони забезпечують виявлення, блокування та видалення загроз, таких як віруси, трояни, програмне забезпечення-вимагач (ransomware) та шпигунські програми. Для компаній важливо вибирати антивірусні рішення, які забезпечують простоту управління та високу ефективність.

Прикладом ефективного антивірусного рішення є Norton Small Business, яке пропонує організаціям малого та середнього масштабу багаторівневий захист, включаючи сканування в реальному часі, фільтрацію веб-трафіку та контроль USB-пристроїв. Інше рішення, Kaspersky Small Office Security, забезпечує розширені функції, такі як захист онлайн-платежів та автоматичне резервне копіювання критичних файлів. Ці інструменти ідеально підходять для середніх та малих бізнесів, які працюють із фінансовими транзакціями.

Важливо зазначити, що сучасні антивірусні програми мають використовувати алгоритми штучного інтелекту та машинного навчання для виявлення загроз. Наприклад, Trend Micro Worry-Free Services Advanced аналізує поведінкові патерни програм, щоб визначати потенційно небезпечні дії. Це дозволяє виявляти навіть ті загрози, які не мають відомих сигнатур.

Антивірусні рішення також інтегруються з іншими інструментами кібербезпеки, як-от мережеві брандмауери, що забезпечує комплексний підхід до захисту даних. Компанії, які впроваджують такі рішення, мають більше шансів уникнути збоїв у бізнес-операціях, викликаних кібератаками.

Системи виявлення та запобігання вторгнень (IDS/IPS) є важливими інструментами, які забезпечують моніторинг мережевого трафіку з метою виявлення та блокування загроз. IDS виявляє аномалії у трафіку, тоді як IPS реагує на них автоматично, блокуючи небезпечні дії. Це дозволяє підприємствам мінімізувати ризик несанкціонованого доступу до мережі.

Одним із прикладів є система Snort, яка пропонує відкритий код і дозволяє налаштовувати правила виявлення загроз відповідно до своїх потреб. Інше популярне рішення — Suricata, яке використовує багатопотоковий аналіз для обробки великих обсягів трафіку у реальному часі. Ці інструменти забезпечують ефективний захист не вимагаючи великих інвестицій, як для комерційних продуктів.

Сучасні IDS/IPS також інтегруються з хмарними рішеннями для моніторингу трафіку між хмарними сервісами. Наприклад, Cisco Secure IPS дозволяє аналізувати взаємодію між локальними мережами та хмарними середовищами. Це особливо корисно для бізнесів, які активно використовують хмарні технології.

Інструменти IDS/IPS можуть бути комбіновані з SIEM-системами для підвищення ефективності виявлення загроз. Завдяки такій інтеграції системи компаній отримують можливість аналізувати загрози в контексті всієї мережевої активності, що дозволяє виявляти складні атаки на ранніх етапах.

Мережеві брандмауери є першочерговим інструментом для захисту бізнесів від зовнішніх загроз. Вони забезпечують контроль доступу до мережі,

блокуючи небажаний трафік і дозволяючи лише авторизовані з'єднання. Сучасні брандмауери нового покоління (NGFW) поєднують традиційні функції брандмауера з додатковими можливостями, такими як аналіз додатків і виявлення загроз у реальному часі.

Прикладом є Fortinet FortiGate, який пропонує комплексний захист мережі з використанням функцій IDS/IPS, фільтрації URL-адрес та контролю додатків. Інший популярний продукт — Palo Alto Networks NGFW, що забезпечує детальний аналіз додатків і дозволяє компаніям створювати політики доступу на основі ідентифікаторів користувачів.[6]

Брандмауери також інтегруються з іншими інструментами безпеки, такими як VPN, для забезпечення безпечного доступу до корпоративних мереж із віддалених місць. Це особливо важливо для організацій, які працюють у гібридному середовищі або використовують віддалені робочі станції.

Сучасні брандмауери також підтримують функції автоматичного оновлення баз даних загроз, що дозволяє малим бізнесам залишатися захищеними від новітніх атак. Крім того, вони можуть інтегруватися з хмарними платформами, такими як AWS або Azure, для захисту багатохмарних середовищ.

Інструменти сканування на предмет вразливостей допомагають ідентифікувати слабкі місця у системах безпеки. Вони дозволяють виявляти помилки у конфігураціях, незахищені протоколи чи застаріле програмне забезпечення, що може стати ціллю для атак. Завдяки цим інструментам бізнеси отримують змогу своєчасно виправляти вразливості та знижувати ризик кібератак.

Прикладом є інструмент Nessus, який дозволяє виконувати сканування мережі для ідентифікації потенційних загроз, таких як незахищені порти чи вразливі служби. Інший приклад — OpenVAS, що пропонує відкритий код і широкий набір функцій для оцінки стану безпеки. Обидва рішення є популярними серед малих компаній завдяки своїй простоті у використанні та високій ефективності.

Сканери вразливостей також дозволяють проводити оцінку відповідності стандартам безпеки, таким як PCI DSS або ISO 27001. Наприклад, Qualys

Vulnerability Management інтегрує можливості сканування із функціями звітності для оцінки ризиків та створення рекомендацій щодо усунення виявлених проблем. Це особливо корисно для бізнесів, які працюють у фінансовій чи медичній галузях.

Крім того, деякі інструменти сканування надають функції для автоматичного оновлення бази даних вразливостей. Це забезпечує бізнеси різного масштабу актуальною інформацією про новітні загрози та дозволяє оперативно реагувати на них. Використання таких рішень допомагає не лише виявляти, але й запобігати експлуатації вразливостей, що є критичним для забезпечення кібербезпеки.

1.3 Виклики та перспективи у протидії мережевим атакам

Протидія мережевим атакам стає дедалі складнішою через швидкий розвиток технологій, які одночасно відкривають нові можливості для захисту та створюють додаткові виклики. Одним із ключових факторів у цій динаміці є зловживання передовими технологіями, зокрема штучним інтелектом (ШІ), який зловмисники використовують для підвищення ефективності своїх атак. Використання ШІ в кіберзагрозах відкриває нові горизонти для атакуючих, що значно ускладнює процес виявлення та протидії таким діям.

Одна з основних загроз, пов'язана з використанням генеративного штучного інтелекту, полягає у його здатності спрощувати створення шкідливих інструментів і коду. Генеративний ШІ дозволяє розробляти складні скрипти або програмні засоби, які можуть бути використані для виконання шкідливих операцій у комп'ютерних мережах (CNO). Така технологія не лише демократизує доступ до складних обчислювальних можливостей, але й значно знижує поріг входу для менш досвідчених зловмисників, дозволяючи їм реалізовувати атаки, які раніше вимагали високого рівня технічної підготовки.

Крім цього, генеративний ШІ ефективно підтримує соціальну інженерію та інформаційні операції. Завдяки здатності створювати переконливий текст і контент, технологія полегшує реалізацію фішингових атак або кампаній із

дезінформації. Наприклад, автоматизація написання електронних листів, що виглядають як автентичні повідомлення від надійних джерел, значно підвищує успішність подібних атак.

Використання генеративного ШІ в зловмисних цілях поки що спостерігається рідко, але його потенціал викликає серйозне занепокоєння. У 2023 році було зафіксовано лише кілька випадків, коли такі технології прямо підтримували розробку або виконання шкідливих операцій. Причини цього можуть включати недостатню кількість прямих спостережень, відсутність чітких індикаторів використання генеративного ШІ або зусилля атакуючих приховати факт застосування цієї технології. Проте тенденція до зростання таких випадків є очевидною, що вимагає від захисників більш глибокого аналізу і розробки нових методів протидії.[1]

Особливої уваги потребують IoT-пристрої, які часто мають обмежені обчислювальні ресурси і слабкий захист. Наприклад, пристрої «розумного дому» або медичні пристрої, інтегровані до мережі, можуть бути використані як точка входу для зловмисників. Атаки на IoT пристрої вже неодноразово призводили до масштабних інцидентів, таких як атака бот-мережі Mirai у 2016 році, яка зруйнувала роботу численних сервісів. Без ефективного контролю і моніторингу безпека таких пристроїв стає слабкою ланкою всієї мережі.

Хмарні середовища також додають нових викликів. Оскільки багато організацій переходять на моделі, що базуються на хмарних обчисленнях, кіберзлочинці активно шукають вразливості у взаємодії між локальними і хмарними мережами. Наприклад, конфігураційні помилки в хмарних системах часто стають причиною витоків даних. У 2021 році було зафіксовано кілька інцидентів, пов'язаних із неналежним налаштуванням Amazon S3-бакетів, що спричинило витoki конфіденційної інформації.

Використання 5G-технологій відкриває нові горизонти для комунікацій, але разом із цим приносить і нові загрози. Збільшення пропускної здатності і зменшення затримок дозволяють реалізовувати складніші атаки, які раніше були неможливі через технічні обмеження. Наприклад, зловмисники можуть

використовувати високошвидкісні з'єднання для реалізації більш масштабних атак типу DDoS.

Організаціям необхідно інтегрувати рішення для багаторівневого моніторингу та захисту мереж. Наприклад, використання SIEM-систем, які здатні аналізувати дані з усіх компонентів мережі, дозволяє швидше виявляти аномалії. Крім того, важливо забезпечувати регулярний аудит конфігурацій хмарних і IoT-платформ, щоб мінімізувати ризики.

Малі та середні підприємства стикаються з унікальними викликами у сфері кібербезпеки через обмеженість ресурсів. Забезпечення захисту сучасних мережевих середовищ вимагає значних фінансових інвестицій у апаратне та програмне забезпечення, навчання персоналу, а також у моніторинг та оперативне реагування. Невеликі бізнеси часто не можуть дозволити собі впровадження передових рішень, таких як SIEM, IDS/IPS або Zero Trust-архітектури, що залишає їх уразливими перед сучасними загрозами.

Застаріле обладнання також стає серйозною проблемою для організацій. Багато малих компаній продовжують використовувати старі системи, які не мають підтримки оновлень безпеки від виробників. Це створює сприятливе середовище для атак, оскільки зловмисники часто експлуатують відомі вразливості у таких системах. Наприклад, атаки на застарілі версії операційних систем або апаратного забезпечення продовжують бути ефективними через відсутність ресурсів на модернізацію.

Недостатність фінансування також впливає на доступ до висококваліфікованих фахівців із кібербезпеки. У малих організаціях часто немає можливості утримувати команди, що займаються безпекою, і ці функції покладаються на IT-відділ, який може не мати відповідного досвіду. Це значно підвищує ризик того, що загрози залишаться непоміченими.

Ще одним обмеженням є складність інтеграції нових технологій із наявною інфраструктурою. Навіть якщо організація має ресурси для впровадження сучасних рішень, такі інтеграції можуть бути дорогими та потребувати багато часу. Наприклад, встановлення SIEM-систем вимагає налаштування збору даних

із різних джерел і навчання персоналу для ефективного використання цього інструменту.

Для подолання цих викликів компаніям рекомендується шукати економічно ефективні рішення, такі як хмарні платформи безпеки або сервіси «безпека як послуга» (Security-as-a-Service). Ці моделі дозволяють малим організаціям отримувати доступ до передових інструментів безпеки без необхідності великих початкових витрат. Крім того, партнерство з зовнішніми постачальниками послуг безпеки може допомогти організаціям захистити свої мережі навіть за обмежених ресурсів.

Серед перспектив точно потрібно виділити застосування штучного інтелекту (ШІ) у протидії мережевим атакам, що відкриває нові можливості для автоматизації та підвищення ефективності захисту. Згідно з даними IBM, організації, які впроваджують ШІ та автоматизацію в кібербезпеку, скорочують середній час виявлення і реагування на інциденти з 277 до 74 днів, що значно знижує потенційні збитки. Це особливо важливо для малих і середніх підприємств, які стикаються з браком ресурсів для своєчасного реагування.

Однією з ключових переваг ШІ є його здатність аналізувати великі обсяги даних у реальному часі. Це дозволяє виявляти приховані загрози, які неможливо розпізнати за допомогою традиційних методів. Наприклад, системи на основі машинного навчання можуть автоматично ідентифікувати аномалії у мережевому трафіку, які можуть свідчити про початок DDoS-атаки або спробу вторгнення.

ШІ також допомагає знизити навантаження на фахівців із кібербезпеки, автоматизуючи рутинні процеси. За даними IBM, організації, які використовують автоматизацію, знижують середню вартість інциденту на 3.05 мільйона доларів США. Це дозволяє зосередити зусилля персоналу на складніших завданнях, таких як стратегічне планування або розслідування складних атак.

Крім того, ШІ забезпечує ефективну інтеграцію між різними системами кіберзахисту. Наприклад, використання технологій на основі ШІ дозволяє SIEM-системам не лише збирати дані про події, але й автоматично виявляти та

блокувати загрози на ранніх стадіях. Це знижує ризики для організацій і забезпечує більш швидке відновлення після інцидентів.

Зростання інвестицій у ШІ для кібербезпеки свідчить про значний інтерес до цієї технології. Згідно з дослідженням IBM, 60% організацій уже використовують або планують впроваджувати ШІ-рішення для захисту своїх мереж. Це свідчить про визнання ефективності ШІ у боротьбі з постійно зростаючими загрозами.

Архітектура Zero Trust стає однією з ключових перспектив у сфері кібербезпеки, оскільки вона пропонує підхід, що мінімізує ризики доступу до мережі. Принцип «нульової довіри» передбачає постійний моніторинг і перевірку кожного запиту на доступ, незалежно від того, чи надходить він із внутрішньої мережі, чи ззовні. За даними IBM, організації, які впровадили Zero Trust, скорочують середню вартість інцидентів на 1,76 мільйона доларів.

Zero Trust забезпечує контроль доступу на основі багатфакторної аутентифікації (MFA) і сегментації мережі. Це означає, що навіть у разі успішного злому одного сегмента зловмисник не отримає доступ до інших частин мережі. Наприклад, використання політик доступу на основі ролей (RBAC) дозволяє надавати доступ лише до необхідних ресурсів, мінімізуючи потенційні ризики.

Впровадження Zero Trust особливо актуальне для організацій, які працюють у хмарних середовищах. Традиційні методи захисту, що покладаються на периметрову безпеку, вже не є ефективними в умовах хмарних обчислень і віддаленої роботи. Zero Trust дозволяє забезпечувати безпеку незалежно від місця розташування користувачів і пристроїв. Інтеграція Zero Trust в сучасні системи моніторингу, такі як SIEM і EDR (Endpoint Detection and Response), розширює можливості виявлення загроз. Наприклад, SIEM може автоматично аналізувати поведінкові патерни користувачів і сигналізувати про підозрілі дії, а EDR забезпечує глибокий контроль над пристроями в мережі.

Організації, які впроваджують Zero Trust, отримують значні конкурентні переваги. Згідно з IBM, у 2023 році 42% компаній уже використовували цей підхід, а ще 30% планують його впровадження в найближчі роки. Це свідчить

про зростаюче визнання ефективності Zero Trust як фундаментальної стратегії кіберзахисту.

Ефективна протидія сучасним кіберзагрозам вимагає колективних зусиль. Співпраця між організаціями у сфері кібербезпеки стає важливою перспективою, яка дозволяє швидко реагувати на нові загрози і обмінюватися даними про кіберінциденти. Такий підхід не лише знижує ризики для окремих організацій, але й сприяє побудові стійкішої глобальної екосистеми безпеки.

Одним із ключових напрямів співпраці є обмін інформацією про загрози через платформи Threat Intelligence Sharing. Ці платформи дозволяють організаціям ділитися даними про шкідливі домени, IP-адреси або інші індикатори компрометації (IoC). Наприклад, використання стандарту STIX/TAXII спрощує обмін даними між різними платформами і дозволяє організаціям інтегрувати отриману інформацію у свої системи захисту, такі як SIEM або IDS/IPS.

Міжнародні ініціативи, такі як програмні альянси Європейського агентства з кібербезпеки (ENISA) або програми INTERPOL Cybercrime Initiative, демонструють, наскільки важливою є співпраця на глобальному рівні. Ці ініціативи забезпечують обмін найкращими практиками, технологіями і методами реагування на загрози. Зокрема, INTERPOL у 2023 році допоміг виявити і нейтралізувати кілька міжнародних кіберзлочинних угруповань завдяки спільній роботі з державними й приватними партнерами.

Також розвивається співпраця між підприємствами та державними структурами. Наприклад, державні агентства в США, такі як CISA (Cybersecurity and Infrastructure Security Agency), активно співпрацюють із приватним сектором, надаючи інструменти для оцінки ризиків і доступ до ресурсів для реагування на інциденти. Подібні ініціативи сприяють підвищенню кібербезпеки в критично важливих секторах, таких як енергетика, транспорт і фінанси.

Ще одним аспектом співпраці є кіберстрахування, яке об'єднує організації, страховиків і постачальників послуг кібербезпеки. Страхові компанії аналізують дані про загрози, зібрані з різних джерел, і надають рекомендації для зниження

ризиків. Це стимулює компанії до впровадження передових рішень і створює загальну основу для зниження впливу кібератак на бізнес.

Таким чином, співпраця між організаціями стає ключовим компонентом сучасної стратегії кіберзахисту. Вона дозволяє поєднувати ресурси, технології та знання для ефективнішої боротьби зі зростаючими кіберзагрозами.

РОЗДІЛ 2 СИСТЕМИ ВИЯВЛЕННЯ ВТОРГНЕНЬ ТА ЗАСТОСУВАННЯ МАШИННОГО НАВЧАННЯ

2.1 Поняття про системи виявлення вторгнень

Перед розробкою і впровадженням обов'язково має бути пройдений етап ознайомлення з визначенням систем виявлення вторгнень (IDS), їх видів та основних функцій. Також важливим є виділити переваги та недоліки існуючих рішень IDS щоб якомога повніше задовільнити потреби кожної окремої інформаційно-комп'ютерної системи.

IDS є спеціалізованими програмно-апаратними засобами, призначеними для моніторингу, аналізу та виявлення підозрілої або зловмисної активності в мережевому трафіку. Основна мета таких систем — своєчасно ідентифікувати потенційні загрози та попередити їхнє розповсюдження, забезпечуючи тим самим цілісність, конфіденційність і доступність інформаційних ресурсів.[7]

Класифікація IDS зазвичай здійснюється за кількома критеріями. Залежно від методу аналізу даних, IDS поділяють на:

- сигнатурні;
- поведінкові;
- гібридні.

Сигнатурні системи використовують базу даних зразків атак, або сигнатур, для порівняння з мережею трафіку. Поведінкові системи, навпаки, орієнтовані на виявлення аномалій у поведінці мережі, які відхиляються від встановлених норм. Гібридні системи поєднують переваги обох підходів, забезпечуючи більш комплексний аналіз і високу точність виявлення загроз.

Принцип роботи IDS базується на зборі й аналізі даних із мережі або хостів. Системи збирають пакети даних, аналізують їх на основі заданих правил чи моделей і реагують у разі виявлення потенційної загрози. У разі сигнатурного аналізу система порівнює вхідний трафік із наявною базою відомих атак. Якщо трафік відповідає певній сигнатурі, система генерує сигнал тривоги. У поведінкових системах використовується аналіз патернів: система виявляє

аномальні відхилення від нормальної поведінки, наприклад, різке збільшення обсягу трафіку або підозрілу активність користувача.

Уявлення про розгортання IDS в комп'ютерній мережі має важливе значення для захисту інформаційних систем у різних галузях. У корпоративному середовищі ці системи забезпечують захист мережі від внутрішніх і зовнішніх загроз, допомагаючи уникнути витоків даних та зупинити атаки на ранніх етапах (рис. 2.1). Наприклад, у банківському секторі IDS застосовуються для запобігання шахрайству, моніторингу фінансових операцій і захисту клієнтських даних. У державних установах системи виявлення вторгнень використовуються для забезпечення захищеності критичної інфраструктури, включаючи енергетичні системи, які за останні 2 роки набули чималої важливості і, як наслідок, варті особливої уваги.[9]

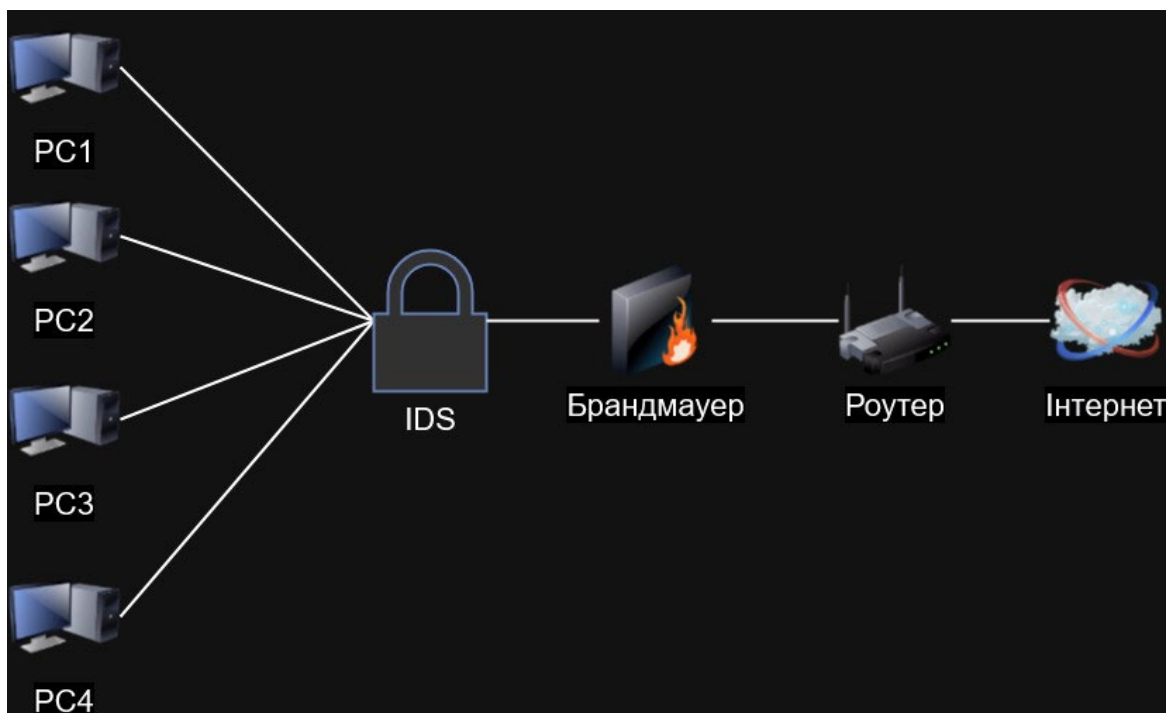


Рисунок 2.1 – Схематичне зображення впровадженої IDS у найпростішу комп'ютерну мережу, як одної з ланок на шляху передавання даних

Однією з головних переваг IDS є можливість інтеграції з іншими системами безпеки, такими як системи управління інформаційною безпекою (SIEM). Це дозволяє збирати й аналізувати дані з різних джерел, забезпечуючи більш глибоке розуміння загальної картини кіберзагроз. Така інтеграція також

сприяє автоматизації реагування на інциденти, що значно скорочує час між виявленням загрози та її усуненням.

Окрім традиційних підходів, у сучасних IDS активно впроваджуються технології штучного інтелекту та машинного навчання. Це дозволяє системам автоматично адаптуватися до нових загроз, навчатися на основі попередніх інцидентів і значно підвищувати точність виявлення. Наприклад, використання кластеризації або нейронних мереж дозволяє ідентифікувати складні атаки, які можуть залишитися непоміченими у традиційних системах. Таким чином, системи виявлення вторгнень є незамінним інструментом для забезпечення кібербезпеки у сучасному світі. Їхнє впровадження дозволяє не лише підвищити рівень захисту інформаційних ресурсів, а й мінімізувати ризики, пов'язані з кібератаками, завдяки інтеграції інноваційних технологій та автоматизації процесів.

Основна мета впровадження IDS полягає у виявленні підозрілої або зловмисної активності в мережі, що дозволяє своєчасно попередити кібератаки чи інші інциденти. Сучасні IDS здатні аналізувати великий обсяг трафіку в реальному часі, використовуючи як сигнатурні підходи, так і аналіз аномалій. Вони забезпечують багаторівневий захист, що охоплює як периферійні, так і внутрішні сегменти мережі, дозволяючи виявляти навіть складні атаки, зокрема "0-day". Згідно з даними IBM, організації, які впроваджують комплексні рішення з кібербезпеки, скорочують середню тривалість виявлення і стримування загроз на 29 днів, що істотно знижує потенційні втрати.[1]

Особливу увагу заслуговує роль штучного інтелекту (ШІ) у підвищенні ефективності систем виявлення вторгнень. ШІ дозволяє автоматизувати аналіз трафіку, виявляти складні патерни та прогнозувати можливі атаки з використанням алгоритмів машинного навчання. Наприклад, технології ШІ можуть допомогти скоротити час виявлення загрози з кількох днів до кількох годин або навіть хвилин. Крім того, за допомогою ШІ системи IDS здатні адаптуватися до нових загроз без необхідності ручного оновлення сигнатур. IBM у своєму звіті зазначає, що компанії, які використовують ШІ в кібербезпеці,

знижують загальні витрати на витоки даних у середньому на 1,76 мільйона доларів США.[1]

Переваги впровадження IDS, підсилених ШІ, очевидні: це не лише підвищення швидкості реагування на загрози, але й можливість точнішого прогнозування ризиків та автоматизація процесів, що раніше вимагали значних людських ресурсів. У результаті, організації отримують більш гнучкий і стійкий підхід до захисту своїх інформаційних активів, що є критично важливим у сучасному кіберсередовищі.

2.2 Огляд існуючих рішень для виявлення вторгнень

Системи виявлення вторгнень (IDS) є важливою складовою сучасної стратегії кіберзахисту. Їхня основна функція полягає у виявленні підозрілої активності у мережевому трафіку або на рівні хостів. Ці системи працюють як попереджувальний механізм, що дозволяє адміністраторам своєчасно реагувати на загрози, зменшуючи можливий збиток від атак. Ефективність IDS залежить від їхньої здатності точно ідентифікувати загрози, зберігаючи мінімальний рівень хибних спрацювань.

Snort, як одна з найпопулярніших IDS, отримала визнання завдяки своїй відкритій архітектурі та великій базі сигнатур. Вона працює на основі аналізу сигнатур відомих атак, що дозволяє швидко виявляти добре задокументовані загрози. Однак Snort може поступатися іншим системам, як-от Suricata, у сценаріях із високим навантаженням через обмеження одночасного оброблення потоків трафіку.

Suricata, розроблена як альтернатива Snort, пропонує багатопотокову (multi-threading) архітектуру. Багатопотоковість дозволяє системі виконувати паралельну обробку мережевих пакетів за допомогою декількох потоків (threads). Тобто в цьому режимі Suricata використовувати всі ядра процесора, які доступні на сервері, на відміну від систем, які працюють лише в одному потоці. Це робить її ідеальним вибором для корпоративних середовищ із високою інтенсивністю передачі даних. Крім того, Suricata підтримує інтеграцію з

сучасними протоколами та технологіями, забезпечуючи ширші можливості аналізу в порівнянні з більш традиційними IDS.

Zeek, який раніше називався Bro, представляє зовсім інший підхід до виявлення загроз, орієнтуючись на поведінковий аналіз. Його основна перевага полягає у здатності аналізувати поведінкові патерни та виявляти складні аномалії. Такий підхід є надзвичайно ефективним для боротьби з новими загрозами, проте Zeek менш підходить для оперативного реагування через повільніший процес детекції.

Комерційні рішення, як-от McAfee Network Security Platform, пропонують додаткові функції, які недоступні у відкритих системах. До таких можливостей належать автоматичне блокування загроз, інтеграція з SIEM-системами та розширені інструменти для аналізу звітів. Однак ці платформи часто є дорогими і орієнтовані на великі організації з розвинутою IT-інфраструктурою.

Таблиця 2.1 – Порівняльна таблиця розглянутих IDS

Характеристика	Snort	Suricata	Zeek (Bro)	McAfee Network Security Platform
Основний підхід	Сигнатурний аналіз	Сигнатурний аналіз	Поведінковий аналіз	Сигнатурний + інтеграція з корпоративними SIEM
Сильні сторони	Відкритий код, активна спільнота, гнучкість	Висока швидкість обробки, багатопотоковість	Аналіз складних аномалій, гнучкість інтеграції	Автоматизація, підтримка корпоративного рівня
Слабкі сторони	Складно працює з великим навантаженням	Вимога ресурсів і складність налаштування	Непридатний для реального часу	Висока вартість, орієнтація на великі компанії

Характеристика	Snort	Suricata	Zeek (Bro)	McAfee Network Security Platform
Найкраще використання	Невеликі і середні організації	Великі мережі з високими навантаженнями	Наукові та дослідницькі середовища	Великі корпоративні екосистеми
Вартість	Відкрите (безкоштовне)	Відкрите (безкоштовне)	Відкрите (безкоштовне)	Комерційне (дороге)

Snort залишається найпоширенішим вибором для невеликих і середніх організацій завдяки своїй доступності та гнучкості. Можливість створювати власні сигнатури дозволяє налаштовувати систему під унікальні потреби користувачів. Крім того, активна спільнота забезпечує постійну підтримку та оновлення баз даних загроз

Suricata, у свою чергу, є вигідним рішенням для організацій із великими обсягами трафіку. Її здатність до багатопотоковості надає перевагу у сценаріях, де швидкість аналізу має критичне значення. Проте для її налаштування та обслуговування можуть знадобитися висококваліфіковані фахівці.

Zeek знаходить своє застосування в дослідницьких середовищах, де аналіз поведінки трафіку має вирішальне значення. Унікальні можливості системи з обробки складних аномалій роблять її незамінною для організацій, які прагнуть зрозуміти більш глибокі закономірності у своїх мережах.

Рішення McAfee Network Security Platform є ідеальними для великих корпоративних екосистем, де потрібен високий рівень автоматизації та інтеграції. Висока вартість таких систем виправдана їхньою продуктивністю та здатністю адаптуватися до нових загроз. Водночас такі рішення можуть бути надмірними для невеликих організацій.

Таким чином, вибір IDS залежить від конкретних потреб організації. Snort пропонує баланс між ціною та функціональністю, Suricata орієнтована на швидкодію, а Zeek надає глибокий поведінковий аналіз. Комерційні рішення, такі як McAfee Network Security Platform, заповнюють прогалину для великих корпоративних середовищ, забезпечуючи інтеграцію з іншими інструментами кібербезпеки.

2.3 Методи аналізу аномалій у мережевому трафіку

Виявлення аномалій у мережевому трафіку є важливим інструментом для забезпечення кібербезпеки сучасних інформаційних систем. Цей підхід базується на ідентифікації відхилень від звичайної поведінки мережі, які можуть вказувати на спроби вторгнень або інші загрози. Аномалії, на відміну від стандартних атак, не завжди відповідають заздалегідь відомим патернам, що ускладнює їх виявлення. Технології аналізу аномалій здатні ефективно доповнювати традиційні сигнатурні системи, забезпечуючи більш комплексний захист мереж.

Основою виявлення аномалій є розуміння нормальної поведінки мережі, яка визначається за допомогою побудови основи (baseline). Вона є еталонним зразком для оцінки майбутніх змін у трафіку. Наприклад, якщо для певного пристрою середня кількість підключень становить 100 за годину, то раптове збільшення цієї кількості до 1000 може бути розцінено як аномалія.

На рисунку 2.2 зображений наочний приклад графіку, що відповідає вхідним запитам на стороні сервера за період у два дні. Бачимо те, що для системи є властивим плавне зростання кількості запитів ближче до 12:00 на позначку приблизно 300 тисяч запитів. Також добре видно аномальні зростання кількості отриманих запитів сервером за короткі відрізки часу, їх виділено червоними маркерами. На додачу, варто розуміти, що базова лінія має регулярно оновлюватися, оскільки поведінка мережі може змінюватися залежно від часу, подій або інших зовнішніх факторів.

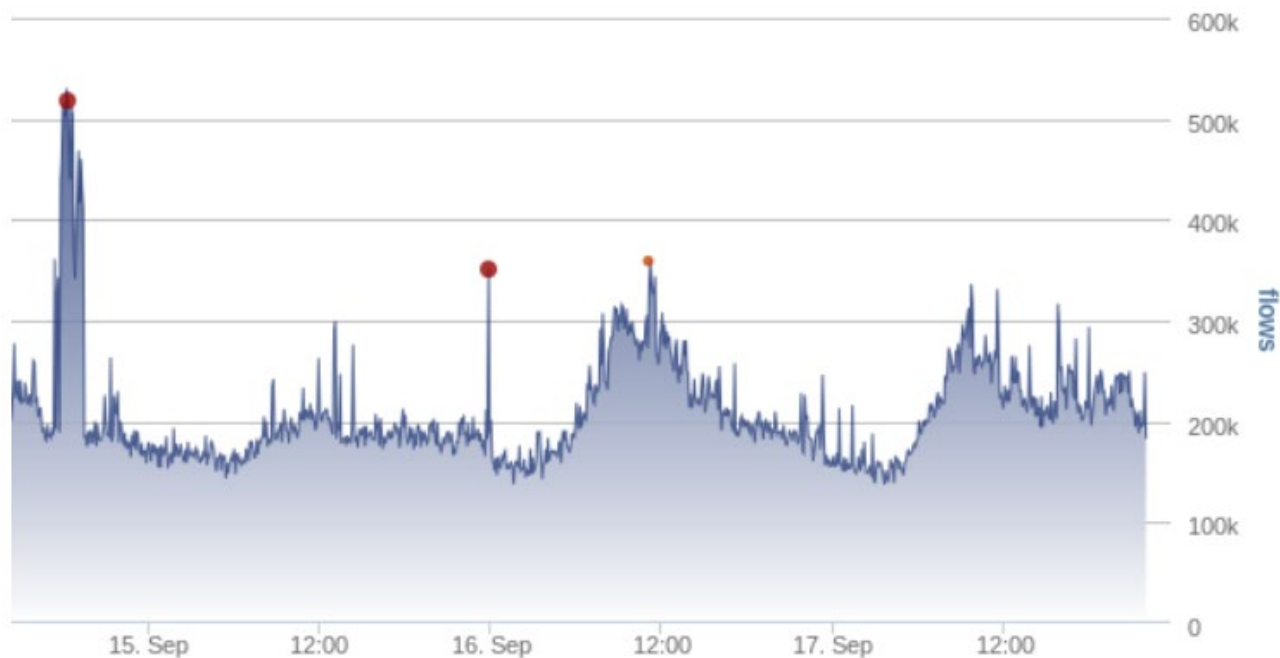


Рисунок 2.2 – Графічне зображення аномалій в мережевому трафіку

Серед підходів виділяється статистичний аналіз, який є основою для багатьох систем виявлення аномалій. Цей підхід включає розрахунок основних статистичних показників, таких як середнє значення, медіана або стандартне відхилення. Відхилення від цих значень може бути розцінено як аномалія. Наприклад, якщо середня тривалість з'єднання в мережі становить 5 хвилин, а поточний сеанс триває понад 2 години, це може свідчити про загрозу.

Інтеграція платформ для виявлення аномалій із системами управління інформацією та подіями безпеки (SIEM) є ще одним важливим елементом. SIEM забезпечують збір і кореляцію даних із різних джерел, що дозволяє поєднувати результати аналізу аномалій із іншими сигналами загроз. Наприклад, виявлення підозрілої активності у мережевому трафіку може бути підтверджено подібною аномалією у логах операційної системи або застосунків.

Сучасні системи виявлення аномалій також інтегруються з інструментами автоматизації. Це дозволяє системам автоматично реагувати на виявлені загрози, наприклад, блокуючи підозрілі IP-адреси або сповіщаючи адміністраторів. Завдяки автоматизації можна значно скоротити час реагування на інциденти, що є критично важливим у середовищах з високою динамікою атак.

Один із викликів у впровадженні систем виявлення аномалій полягає у потребі в якісних даних для побудови моделей. Дані, що містять помилки, нерепрезентативні або неповні, можуть суттєво вплинути на точність аналізу. Крім того, великий обсяг даних створює додаткове навантаження на обчислювальні системи, що потребує значних ресурсів для їхньої обробки.

Часові відрізки (timelines) є важливими компонентами у аналізі аномалій, оскільки вони дозволяють враховувати динамічні зміни у трафіку. Наприклад, різке зростання кількості підключень у короткий проміжок часу може свідчити про атаку типу DDoS. Аналіз часових відрізків дозволяє виявляти такі аномалії у реальному часі, що робить його незамінним інструментом у мережах із високим трафіком.

Додатково, аналіз аномалій може включати ідентифікацію специфічних типів поведінки трафіку, які зазвичай не відповідають звичайному функціонуванню. Наприклад, некоректні пакети, що надсилаються у великій кількості, або підключення до невідомих доменів можуть бути ознаками кіберзагроз.

Ще одним важливим підходом є побудова профілів активності для кожного вузла мережі. Ці профілі враховують типові обсяги передачі даних, часові інтервали активності та використовувані порти. Виявлення відхилень у цих параметрах дозволяє більш точно визначати аномалії.

Розширене використання поведінкових моделей у сучасних системах забезпечує додаткову гнучкість у роботі з аномаліями. Поведінковий підхід дозволяє зосередитися не лише на окремих подіях, але й на їхніх взаємозв'язках, що є важливим для розуміння складних загроз, таких як багатофазні атаки. Наприклад, сукупність невеликих відхилень у трафіку може вказувати на спробу злому в декілька етапів.

Також, ефективним є застосування кореляційного аналізу між подіями в різних сегментах мережі. Наприклад, одночасне збільшення навантаження на сервер бази даних і аномальна активність у сегменті зовнішніх з'єднань можуть бути взаємопов'язаними і вказувати на складну атаку.

Таким чином, виявлення аномалій у мережевому трафіку є складним і багатогранним процесом, який потребує інтеграції різних підходів та інструментів. Поєднання статистичного аналізу, поведінкових моделей та автоматизації дозволяє створювати системи, здатні ефективно ідентифікувати загрози в реальному часі.

2.4 Використання машинного навчання у виявленні аномалій

Машинне навчання є однією з ключових технологій, яка використовується для виявлення аномалій у мережевому трафіку. Його значення обумовлене здатністю аналізувати великі обсяги даних, виділяти приховані патерни та ідентифікувати потенційні загрози, які неможливо виявити за допомогою традиційних підходів. Застосування машинного навчання дозволяє автоматизувати процес виявлення, зменшуючи залежність від людського фактора і підвищуючи точність виявлення загроз.

Алгоритми машинного навчання є основою для багатьох сучасних рішень із виявлення аномалій. Зокрема, методи з наглядом, такі як Random Forest або метод опорних векторів (SVM), використовуються для класифікації поведінки як нормальної чи аномальної на основі попередньо маркованих даних. Алгоритми без нагляду, як-от K-means, дозволяють групувати схожі зразки даних, ідентифікуючи ті, що суттєво відрізняються від основної маси. Використання глибокого навчання, наприклад, автоенкодерів, значно підвищує точність аналізу багатовимірних даних.

Методи з наглядом (supervised learning) є найпоширенішими у системах виявлення аномалій. Вони вимагають наявності маркованих наборів даних, які містять приклади нормальної та аномальної поведінки. Серед популярних алгоритмів можна виділити: Random Forest і метод опорних векторів (SVM).

«Random forest – ансамблевий метод машинного навчання для класифікації, регресії та інших завдань, які оперують за допомогою побудови численних дерев прийняття рішень під час тренування моделі і продукують моду

для класів (класифікацій) або усереднений прогноз (регресія) побудованих дерев (рис. 2.3).»[14]

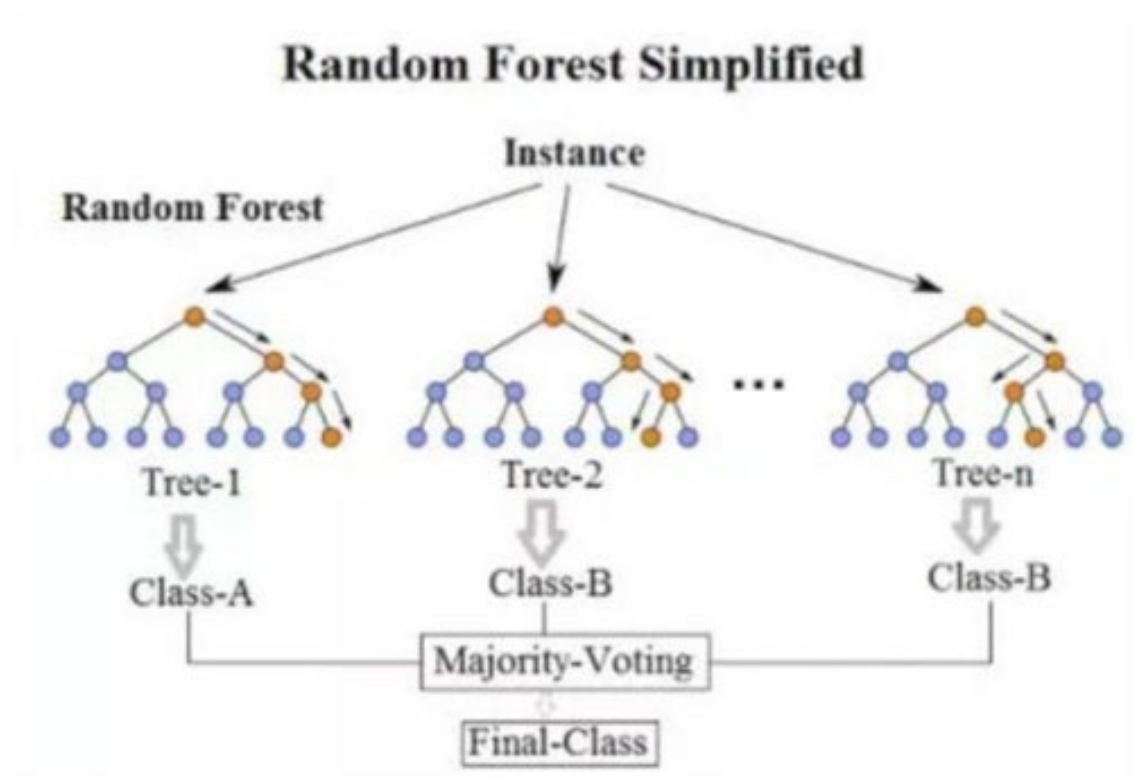


Рисунок 2.3 – Графічне зображення алгоритму Random Forest

У стандартному дереві рішень для поділу вузла аналізуються всі доступні ознаки, після чого обирається та, яка найкраще розділяє значення у вузлах. У випадковому лісі процес відбору ознаки обмежується випадковою підмножиною доступних ознак. Це сприяє більшій варіативності серед дерев у моделі, що знижує кореляцію між ними та підвищує їх різноманітність.

«Метод опорних векторів (Support Vector Machine) застосовується для пошуку аномалій в системах, де нормальна поведінка видається тільки одним класом. Даний метод визначає межу регіону, в якому знаходяться екземпляри нормальних даних. Для кожного досліджуваного екземпляра визначається, чи знаходиться він в певному регіоні. Якщо екземпляр виявляється поза регіоном, він позначається як аномальний (рис. 2.4).»[14]

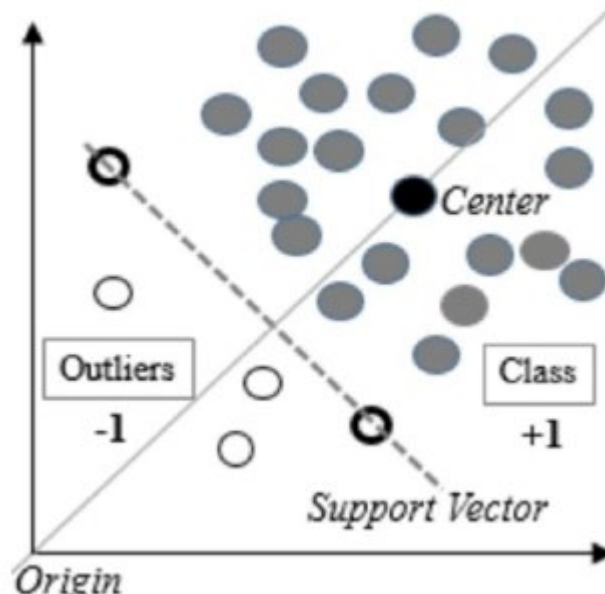


Рисунок 2.4 – Виявлення аномалій за допомогою SVM

Алгоритми без нагляду (unsupervised learning) є ефективними у випадках, коли марковані дані недоступні. Кластеризація, наприклад, алгоритм K-means, дозволяє групувати дані за схожістю, виділяючи кластери аномалій.

Метод k-середніх (K-Means) є популярним підходом, який спрямований на зменшення середньої квадратичної відстані між точками всередині одного кластера. Його головна ідея полягає у формуванні кластерів таким чином, щоб мінімізувати загальну варіацію в межах кластерів, тобто зменшити сумарну відстань між точками та їхнім центроїдом. Хоча існує кілька алгоритмів реалізації цього методу, стандартний алгоритм визначає загальну варіацію як суму квадратичних евклідових відстаней між точками та центроїдами їхніх кластерів.

$$W(C_k) = \sum_{x_i \in C_k} (x_i - \mu_k)^2 \quad (1.1)$$

де x_i – точка даних, що належить до кластера C_k , μ_k – середнє значення точок, що присвоєні кластеру C_k . [14]

Алгоритм вимагає заздалегідь визначити кількість кластерів. Спочатку центроїди кластерів задаються випадковим чином. На кожній наступній ітерації ці центроїди оновлюються відповідно до середнього значення точок у кожному кластері, сформованому на попередньому етапі. Далі точки знову розподіляються між кластерами залежно від того, до якого з оновлених центроїдів вони є ближчими за обраною метрикою. Процес завершується, коли на певній ітерації розподіл точок між кластерами перестає змінюватися.

Інший підхід — DBSCAN, який використовує щільність точок для ідентифікації ізольованих даних, що можуть свідчити про аномалії.

«Ідея методу DBSCAN (Density Based Spatial Clustering of Application with Noise) заснована на гіпотезі, що елементи одного кластера формують область, щільність об'єктів усередині якого, за деяким заданим порогом ϵ , перевищує щільність за його межами. У кластері повинна бути певна «безперервність» даних, іншими словами, якщо точки лежать недалеко один від одного, то і значення функції в них не повинно сильно відрізнитися. Таким чином, алгоритм групує разом точки, які щільно розташовані, позначаючи як викиди точки, які знаходяться самотньо в областях з малою щільністю.»[14]

Щоб створити оптимальну модель, необхідно налаштувати гіперпараметри. Це досягається шляхом запуску алгоритму з різними параметрами та вибору найкращого варіанта. Процес підбору гіперпараметрів є завданням оптимізації, де цільова функція — це показник точності або рівень помилки моделі. Завдання оптимізації полягає у знаходженні максимального значення (для точності) або мінімального (для помилок).

Для алгоритму DBSCAN основними гіперпараметрами є:

- ϵ — радіус гіперсфери або максимальна відстань між сусідніми точками, яка визначає межі кластеру;
- minPts — мінімальна кількість точок, необхідна для формування кластеру; якщо кількість сусідів точки менша за це значення, точка вважається шумом;
- distFunc — метрика для обчислення відстані між точками. Це може бути евклідова відстань, Манхеттенська метрика чи ортодромічна відстань.

Оптимальне значення параметра `minPts` залежить від цілей аналізу. Воно повинно відповідати мінімальному розміру, який вважається значущим кластером. Наприклад, якщо `minPts` дорівнює 20, усі кластери меншого розміру будуть або класифіковані як шум, або об'єднані із сусідніми кластерами. Збільшення цього параметра може призвести до злиття менших кластерів, тоді як його зменшення дозволить виділити більше кластерів, але зменшить їхню значущість.

Значення `minPts` також впливає на обчислення кластерної відстані. Кластерна відстань визначається як дистанція, необхідна для включення точки до кластеру відповідно до заданого `minPts`. Якщо вибрано велике значення `minPts`, кластерна відстань буде більшою, і точки на межах кластерів можуть залишитися незв'язаними. Натомість при меншому значенні `minPts` кластерна відстань зменшується, і точки легше об'єднуються в кластер, що може збільшити кількість кластерів, але знизити їхню якість.

Глибоке навчання (deep learning) є ще одним перспективним підходом для аналізу аномалій. Використання нейронних мереж, зокрема автоенкодерів, дозволяє обробляти багатовимірні та складні набори даних. Автоенкодери працюють шляхом стиснення даних до меншого розміру, з подальшим відновленням початкового вигляду. Значні відхилення у процесі відновлення сигналізують про можливу аномалію.

Рекурентні нейронні мережі (RNN) і довготривала короткочасна пам'ять (LSTM) є ефективними для аналізу часових рядів у мережевому трафіку. Вони здатні враховувати послідовність подій і передбачати наступні значення. Наприклад, LSTM дозволяє моделювати динамічні зміни у трафіку, такі як:

- Різке збільшення кількості підключень до одного сервера.
- Зміни у розподілі запитів, що можуть свідчити про атаку типу DDoS.

Використання ансамблів алгоритмів, як-от Isolation Forest, є ще одним ефективним підходом для виявлення аномалій. Isolation Forest працює, виділяючи точки даних, які легко ізолюються у просторі. Цей підхід є швидким та ефективним для великих наборів даних, що робить його популярним у реальному застосуванні.

Попри переваги, використання машинного навчання у виявленні аномалій стикається з низкою викликів. Одним із головних є потреба у великих обсягах якісних даних для навчання моделей. Інші виклики включають:

- Неповні або некоректні дані, які можуть суттєво знижувати точність алгоритмів.
- Складність адаптації моделей до швидкозмінних мережових середовищ, що може призводити до появи хибних спрацювань.

Переваги застосування машинного навчання у виявленні аномалій включають підвищену точність, здатність обробляти великі обсяги даних і адаптивність до нових типів загроз. Алгоритми машинного навчання дозволяють ідентифікувати патерни, які неможливо виявити за допомогою традиційних підходів, що робить їх незамінними у сучасній кібербезпеці.

Інтеграція машинного навчання з іншими системами, такими як SIEM (Security Information and Event Management), розширює можливості виявлення аномалій. SIEM забезпечують централізований збір даних, а моделі машинного навчання використовують ці дані для побудови точних прогнозів. Таке поєднання дозволяє не лише виявляти аномалії, а й здійснювати їхню контекстуальну оцінку.

Таким чином, машинне навчання є основою сучасних систем виявлення аномалій у мережевому трафіку. Завдяки використанню різних алгоритмів, таких як Random Forest, SVM, автоенкодери та LSTM, можливо ефективно протидіяти загрозам, які постійно еволюціонують. Подальший розвиток цієї технології сприятиме покращенню точності, швидкості та адаптивності у боротьбі з кіберзагрозами.

РОЗДІЛ 3 РОЗРОБКА ТА ТЕСТУВАННЯ СИСТЕМИ ВИЯВЛЕННЯ ВТОРГНЕНЬ

3.1 Вибір інструментів та технологій для побудови системи

Ефективна система виявлення вторгнень (IDS) вимагає використання сучасних технологій, які здатні забезпечити високу точність, адаптивність і масштабованість. Для розробки такої системи необхідно ретельно обрати інструменти, які не лише відповідають вимогам проєкту, але й забезпечують можливість інтеграції різних компонентів. У цьому підрозділі буде обґрунтовано вибір технологій для побудови системи.

Одним із основних компонентів є система виявлення вторгнень Snort. Snort є відкритим програмним забезпеченням, яке вже зарекомендувало себе як ефективний інструмент для аналізу мережевого трафіку та виявлення вторгнень. Його популярність пояснюється підтримкою великої кількості правил для ідентифікації загроз, можливістю гнучкого налаштування та великою спільнотою користувачів, які забезпечують своєчасне оновлення бази правил.

Ключовою перевагою Snort є його здатність до інтеграції з іншими системами, зокрема системами машинного навчання. Це дозволяє не лише працювати з сигнатурними методами виявлення загроз, але й застосовувати алгоритми аналізу аномалій. Такий підхід є особливо важливим для виявлення нових, раніше невідомих загроз.

Для реалізації компонента машинного навчання було обрано Python як основну платформу розробки. Python пропонує широкий набір бібліотек для аналізу даних, побудови моделей і автоматизації процесів, таких як Scikit-learn, TensorFlow, Keras тощо. Ці інструменти забезпечують гнучкість у виборі алгоритмів і дозволяють ефективно працювати з великими обсягами даних.

Ще одним важливим інструментом є платформа для зберігання та обробки логів. Для цього доцільно використовувати бази даних, такі як MySQL або PostgreSQL, які забезпечують високу продуктивність і підтримують інтеграцію з іншими компонентами системи. Зберігання логів дозволяє проводити

ретроспективний аналіз, що є важливим для розслідування інцидентів та вдосконалення системи.

Для аналізу та візуалізації даних можна використовувати такі інструменти, як Grafana або Kibana. Вони дозволяють створювати інтерактивні дашборди для моніторингу мережевого трафіку, аналізу продуктивності системи та відстеження виявлених загроз у реальному часі. Такий підхід підвищує зручність роботи з системою та забезпечує оперативність у прийнятті рішень.

Для підготовки та тестування даних обрано стандартні набори, такі як CICIDS2017 і KDDCup99. Ці набори містять як звичайний трафік, так і дані про різні типи атак, що дозволяє ефективно навчати моделі та оцінювати їхню точність. Використання таких стандартів забезпечує порівнюваність результатів із іншими дослідженнями.

Для моделювання атак і тестування системи доцільно використовувати інструменти, такі як Metasploit або Scapy. Вони дозволяють створювати реалістичні сценарії атак, включаючи DDoS, SQL Injection, MITM, фішинг тощо. Це забезпечує всебічну перевірку ефективності розробленої системи.

Також важливим елементом є вибір інструментів для сповіщення про загрози. Наприклад, використання Syslog для централізованого збирання логів або інтеграція з електронною поштою для оперативного інформування адміністраторів про інциденти. Це дозволяє мінімізувати час реакції на загрози.

На додаток до технічних інструментів, у системі буде враховано регуляторні вимоги, такі як GDPR чи ISO 27001. Це забезпечує відповідність стандартам безпеки та підвищує довіру до системи з боку клієнтів і користувачів.

Таким чином, вибір інструментів та технологій для побудови системи базується на їхній ефективності, сумісності та здатності до інтеграції. Використання цих технологій дозволить створити сучасну, масштабовану та адаптивну систему виявлення вторгнень, що забезпечить захист від широкого спектра загроз.

3.2 Налаштування та розгортання Snort

Для розгортання системи виявлення вторгнень було обрано операційну систему Ubuntu версії 24.04.1. Це рішення зумовлене низкою факторів, які роблять Ubuntu оптимальним вибором для побудови інфраструктури кібербезпеки. Насамперед, Ubuntu є одним із найбільш популярних дистрибутивів Linux, що забезпечує широке ком'юніті підтримки, доступ до великої кількості документації та регулярні оновлення. Крім того, версія 24.04.1 є останнім довгостроковим випуском (LTS), що гарантує підтримку на рівні безпеки та стабільності протягом кількох років.

Ubuntu також надає потужний пакетний менеджер APT, який спрощує встановлення й оновлення необхідного програмного забезпечення, зокрема таких інструментів, як Snort, які використовуються для аналізу мережевого трафіку. Інтеграція з сучасними апаратними платформами, підтримка контейнеризації через Docker і Kubernetes, а також зручність налаштування мережевих параметрів роблять Ubuntu 24.04.1 ідеальним вибором для експериментів із системами виявлення вторгнень у локальних та віртуальних середовищах. Цей дистрибутив також забезпечує високу сумісність із віртуалізаційними платформами, такими як VirtualBox, що полегшує створення тестових середовищ для імітації атак і перевірки роботи системи.

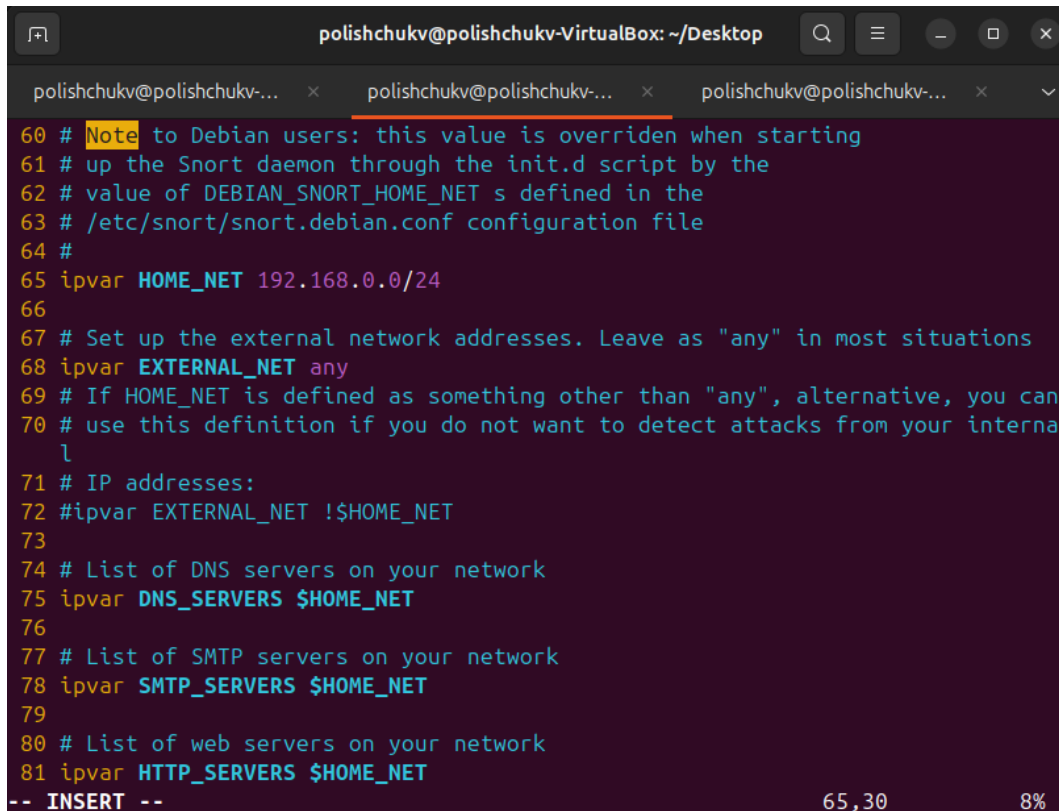
Snort є одним із найпопулярніших інструментів для виявлення вторгнень, що використовується для аналізу мережевого трафіку та виявлення аномалій чи загроз. Процес встановлення Snort у середовищі Linux є досить простим завдяки пакетному менеджеру APT. Для встановлення Snort на систему було використано команду: `sudo apt install snort`.

Ця команда завантажує та встановлює всі необхідні компоненти Snort із офіційних репозиторіїв. Прапорець `-y` автоматично підтверджує встановлення, спрощуючи процес. Після завершення інсталяції система створює базові конфігураційні файли та структуру каталогів для подальшої роботи Snort.

Одним із перших кроків налаштування Snort було переведення мережевого адаптера у `promiscuous` режим командою `ip link set eth1 promisc on`. Цей режим

дозволяє мережевому адаптеру захоплювати весь трафік у мережі, а не лише трафік, адресований конкретному хосту. Така функціональність є необхідною для ефективної роботи Snort, оскільки він повинен аналізувати весь мережевий трафік для виявлення потенційних загроз. Переведення адаптера у цей режим забезпечило можливість аналізувати дані, які проходять через віртуальну мережу, зокрема у середовищі VirtualBox.

Після цього було змінено базові параметри конфігурації у файлі `/etc/snort/snort.conf`. Одним із найважливіших налаштувань є визначення змінної `HOME_NET`, яка вказує Snort на мережу, що буде моніторитися. За замовчуванням значення цієї змінної виставлено, як “any”, проте для моніторингу визначеної мережі, параметр було змінено, як на рисунку 3.1.

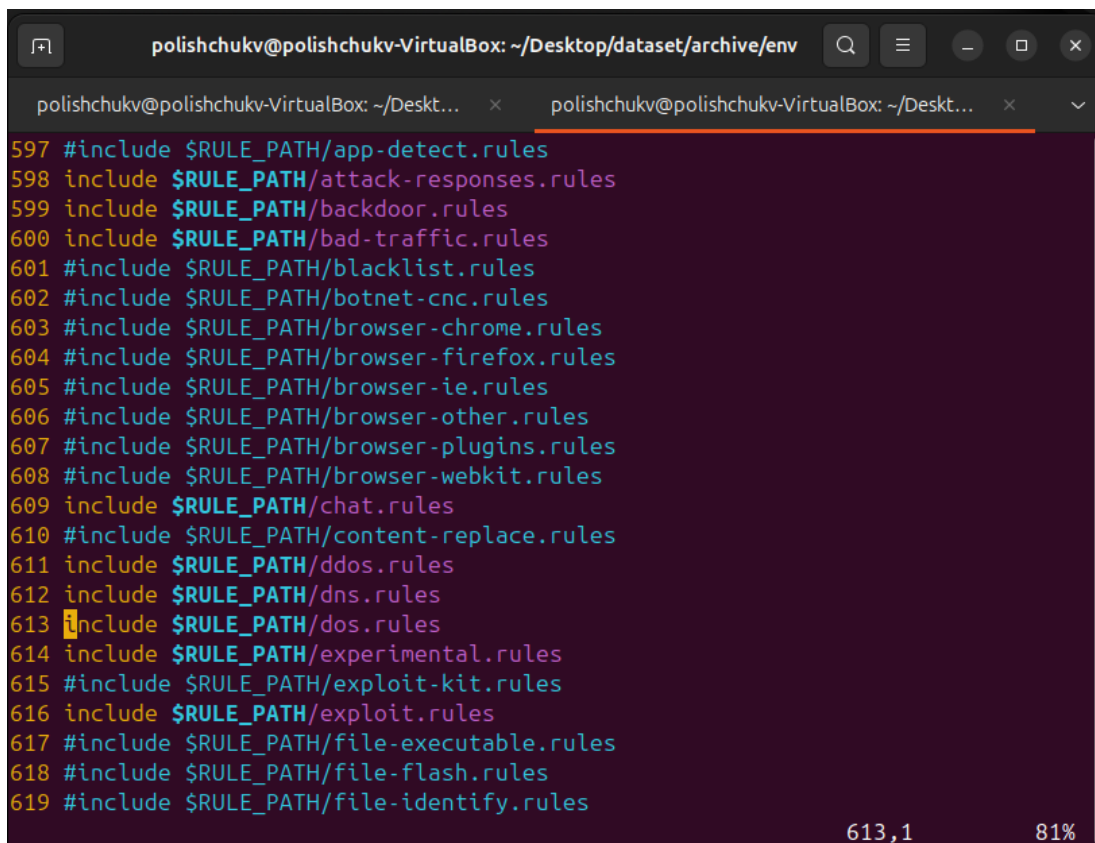


```
polishchukv@polishchukv-VirtualBox: ~/Desktop
polishchukv@polishchukv-... x polishchukv@polishchukv-... x polishchukv@polishchukv-... x
60 # Note to Debian users: this value is overridden when starting
61 # up the Snort daemon through the init.d script by the
62 # value of DEBIAN_SNORT_HOME_NET s defined in the
63 # /etc/snort/snort.debian.conf configuration file
64 #
65 ipvar HOME_NET 192.168.0.0/24
66
67 # Set up the external network addresses. Leave as "any" in most situations
68 ipvar EXTERNAL_NET any
69 # If HOME_NET is defined as something other than "any", alternative, you can
70 # use this definition if you do not want to detect attacks from your interna
71 # IP addresses:
72 #ipvar EXTERNAL_NET !$HOME_NET
73
74 # List of DNS servers on your network
75 ipvar DNS_SERVERS $HOME_NET
76
77 # List of SMTP servers on your network
78 ipvar SMTP_SERVERS $HOME_NET
79
80 # List of web servers on your network
81 ipvar HTTP_SERVERS $HOME_NET
-- INSERT --
65,30 8%
```

Рисунок 3.1 – Виставлення мережі для моніторингу за допомогою Snort

Це відповідає мережевому діапазону віртуальних машин VirtualBox, які використовуються на пристрої. Вказівка цього параметра дозволяє Snort правильно ідентифікувати трафік, який надходить до цієї мережі, та розпізнавати загрози, що можуть бути спрямовані на неї.

Після внесення змін до конфігураційного файлу Snort було перезапущено для застосування оновлених налаштувань. На цьому етапі система готова до базового моніторингу мережевого трафіку та виявлення загроз. Це початкове налаштування створило базу для подальшого розширення функціональності Snort, зокрема активація правил для виявлення специфічних атак. За замовчуванням, ці правила присутні у конфігураційному файлі але у деактивованому стані. Щоб перевести їх у активний стан потрібно прибрати символ коментування “#”, як на рисунку 3.2. На ньому бачимо, що після внесення змін, правила у стрічках 600, 611, 613 та 616 стали активними, на відміну від правил 601-608, які є закоментованими, а отже – пропускаються під час перегляду списку правил програмою.



```
polishchukv@polishchukv-VirtualBox: ~/Desktop/dataset/archive/env
polishchukv@polishchukv-VirtualBox: ~/Deskt... x polishchukv@polishchukv-VirtualBox: ~/Deskt... x
597 #include $RULE_PATH/app-detect.rules
598 include $RULE_PATH/attack-responses.rules
599 include $RULE_PATH/backdoor.rules
600 include $RULE_PATH/bad-traffic.rules
601 #include $RULE_PATH/blacklist.rules
602 #include $RULE_PATH/botnet-cnc.rules
603 #include $RULE_PATH/browser-chrome.rules
604 #include $RULE_PATH/browser-firefox.rules
605 #include $RULE_PATH/browser-ie.rules
606 #include $RULE_PATH/browser-other.rules
607 #include $RULE_PATH/browser-plugins.rules
608 #include $RULE_PATH/browser-webkit.rules
609 include $RULE_PATH/chat.rules
610 #include $RULE_PATH/content-replace.rules
611 include $RULE_PATH/ddos.rules
612 include $RULE_PATH/dns.rules
613 include $RULE_PATH/dos.rules
614 include $RULE_PATH/experimental.rules
615 #include $RULE_PATH/exploit-kit.rules
616 include $RULE_PATH/exploit.rules
617 #include $RULE_PATH/file-executable.rules
618 #include $RULE_PATH/file-flash.rules
619 #include $RULE_PATH/file-identify.rules
```

Рисунок 3.2 – Список правил в конфігураційному файлі snort.conf

Для моніторингу трафіку у мережі за допомогою Snort в режимі реального часу потрібно почати з правильної конфігурації цього інструменту. Спочатку потрібно вибрати мережевий інтерфейс, через який Snort буде аналізувати трафік. Команда `snort -i enp0s3` вказує на моніторинг інтерфейсу, який

налаштований на локальну мережу віртуальних машин VirtualBox на основній машині, однак його можна замінити на інший потрібний інтерфейс, залежно від конфігурації системи. Важливо впевнитись, що обраний інтерфейс має права адміністратора для перехоплення мережевого трафіку, що зазвичай потребує запуску Snort з підвищеними привілеями за допомогою команди `sudo` на початку команди.

Наступним кроком є вказівка шляху до конфігураційного файлу, який задає всі правила роботи Snort. Частина “-c /etc/snort/snort.conf” дозволяє програмі використовувати конфігурацію, яка включає активацію правил виявлення та визначення сканованої мережі. Конфігураційний файл є центральним елементом роботи Snort, тому його ретельно налаштовано перед запуском програми.

Для запису результатів аналізу Snort потрібно вказати директорію для лог-файлів за допомогою параметра -l. Наприклад, опція -l /var/log/snort задає місце, куди Snort записуватиме інформацію про виявлені загрози. Лог-файли можна переглядати за допомогою стандартних інструментів Linux, таких як `tail`, що дозволяє відстежувати вміст файлу в реальному часі.

Важливим етапом є перевірка конфігурації перед запуском Snort у робочому режимі. Це можна зробити за допомогою команди `snort -T -c /etc/snort/snort.conf`, яка тестує конфігураційний файл на наявність помилок. Якщо перевірка проходить успішно, програма повідомить про готовність до аналізу. У разі виявлення помилок вони будуть виведені в термінал із відповідними підказками.

Для запуску Snort у реальному часі використовується команда `snort -A full -i enp0s3 -l /var/log/snort`, яка забезпечує аналіз трафіку з повним записом у лог-файл. Опція -A full дозволяє включити детальний формат виведення, що спрощує обробку даних аналітиками з кібербезпеки. Цей формат містить повну інформацію про виявлені загрози, включаючи час, джерело та призначення трафіку, а також опис спрацювань правил.

Для моніторингу лог-файлів у режимі реального часу використовується команда `tail -f /var/log/snort/alert.log`. Вона дозволяє відображати останні зміни у

файлі, що дуже зручно для оперативного аналізу ситуації. Аналітики можуть миттєво реагувати на підозрілу активність, оскільки всі спрацювання правил з'являються у терміналі у міру їх реєстрації.

Додатково, для роботи Snort у фоновому режимі можна використовувати параметр `-D`. Ця опція дозволяє запускати програму без блокування терміналу, залишаючи її активною у системі. Це корисно для постійного моніторингу, коли Snort працює на сервері чи іншому пристрої, відповідальному за аналіз трафіку. Як ще одну додаткову опцію для розширення функціональності Snort можна додати власні правила у файл `local.rules`. Цей файл дозволяє налаштовувати сигнатури під специфічні потреби вашої мережі. Наприклад, можна створити правило для виявлення специфічного типу атак або моніторингу підозрілих IP-адрес. Коректне налаштування правил дозволяє значно підвищити ефективність аналізу трафіку. В контексті цієї роботи, ця можливість може бути використана для оновлення списку файлу сигнатурами, раніше невідомих атак, які виявлятиме модель машинного навчання.

3.3 Розробка системи машинного навчання для виявлення аномалій

Датасет CIC-IDS-2017 є одним із найбільш поширених наборів даних для аналізу мережевого трафіку та виявлення аномалій. Він складається з набору CSV-файлів, які містять інформацію про різні типи мережевих атак і нормальний трафік. Першим кроком у підготовці моделі є завантаження цих файлів і ознайомлення з їхньою структурою. Кожен файл містить записи, які представляють мережеві сесії, з численними ознаками (features), такими як IP-адреси, порти, тривалість сесії, кількість переданих байтів тощо. Також присутні мітки класів, які вказують, чи є сесія нормальною, чи вона належить до однієї з атак.

Завантаження та попередній аналіз даних виконуються за допомогою бібліотек Python, таких як `pandas`. Основні дії включають перевірку відсутніх значень, оцінку розподілу міток та ознайомлення зі статистичними

характеристиками ознак. Цей етап дозволяє зрозуміти якість даних і підготувати їх до подальшої обробки.

Дані з датасету можуть містити пропуски, дублікати або нерелевантні ознаки, які потрібно усунути перед використанням у моделі. Наприклад, IP-адреси та порти, що є категоріальними даними, часто не несуть корисної інформації для алгоритмів машинного навчання. Тому їх або виключають, або перетворюють у числовий формат за допомогою методів, таких як one-hot encoding.

Крім того, деякі ознаки можуть мати різні масштаби. Наприклад, кількість переданих байтів може варіюватися у широких межах, тоді як часова тривалість сесії може бути значно меншою за значенням. Для забезпечення коректності роботи моделі дані нормалізують або стандартизують. Це можна зробити за допомогою бібліотеки scikit-learn та її функцій, таких як StandardScaler.

У датасеті CIC-IDS-2017 часто спостерігається дисбаланс класів: нормальний трафік переважає, а атаки можуть бути представлені значно меншою кількістю записів. Такий дисбаланс може вплинути на ефективність моделі, оскільки алгоритм може схилитися до переваги домінуючого класу. Для вирішення цієї проблеми використовують методи, такі як:

- Oversampling: штучне збільшення кількості записів менш представленого класу.
- Undersampling: зменшення кількості записів переважаючого класу.
- Метод SMOTE (Synthetic Minority Oversampling Technique): генерація нових зразків менш представленого класу.

Ці методи дозволяють збалансувати дані й забезпечити рівномірний внесок кожного класу у процес навчання моделі.

Для аналізу мережевого трафіку доцільно використовувати алгоритми, які добре працюють із великими та складними наборами даних. Одним із таких алгоритмів є Random Forest. Він базується на ансамблі дерев рішень і добре підходить для задач класифікації. Random Forest автоматично враховує важливість ознак і є стійким до переобчислення, що дозволяє використовувати

його навіть на шумних даних. Його основна перевага полягає у здатності обробляти нерівномірно розподілені ознаки та мінімізувати ризик перенавчання.

Іншим підходом є DBSCAN — алгоритм кластеризації, що працює на основі щільності даних. Він є ефективним для виявлення аномалій, оскільки виділяє точки, які не належать до жодного кластера, як шум. У випадку мережевого трафіку DBSCAN дозволяє ідентифікувати нехарактерну активність, яка може вказувати на атаку, навіть без чітко визначених міток.

Random Forest є ідеальним вибором для задачі класифікації в контексті датасету CIC-IDS-2017, оскільки він:

- Ефективно працює з великим числом ознак, автоматично вибираючи найбільш релевантні.
- Справляється з дисбалансом класів, оскільки може бути налаштований на врахування ваги кожного класу.
- Пояснюваний: дозволяє оцінити важливість кожної ознаки для прийняття рішення. Завдяки цим властивостям Random Forest є потужним інструментом для виявлення атак, які мають чітко визначені патерни.

Алгоритм DBSCAN підходить для виявлення аномалій завдяки своїй здатності працювати із даними, які не потребують попереднього маркування. Його основні переваги:

- Можливість обробки точок, які значно відрізняються від інших, що робить його корисним для аналізу невідомих атак.
- Відсутність необхідності задавати кількість кластерів заздалегідь, що є зручним для аналізу реального мережевого трафіку, де наперед невідомо, скільки груп поведінкових патернів існує.
- Здатність працювати з нелінійними кластерами, які є типовими для мережевого трафіку.

Під час підготовки моделі датасет необхідно розділити на тренувальну та тестову вибірки. Це дозволяє забезпечити об'єктивну оцінку ефективності моделі. Стандартний підхід — розділення у співвідношенні 80/20 або 70/30 за допомогою функції `train_test_split` з бібліотеки `scikit-learn`. Це забезпечує, що

модель навчатиметься на більшості даних, але перевірятиметься на окремій частині, яку вона не бачила раніше.

Ці етапи підготовки є фундаментом для ефективного навчання моделі машинного навчання, яка буде здатна аналізувати мережевий трафік і виявляти потенційні загрози. Дотримання цих принципів забезпечує високу якість моделі та її здатність до виявлення аномалій у реальних умовах.

Як початок розробки моделі машинного навчання імпортуємо потрібні бібліотеки Python, а саме:

- pandas для роботи з даними у форматі таблиць;
- numpy для обробки числових даних;
- RandomForestClassifier для класифікації на основі методу Random Forest;
- train_test_split для розділення даних на тренувальну та тестову вибірки;
- StandardScaler для нормалізації даних;
- GridSearchCV інструмент для перебору комбінацій параметрів з крос-валідацією;
- LabelEncoder для перетворення текстових міток у числовий формат.

Тому перші рядки файлу mlIDS.py вказані в лістингу 3.1. [16]

Лістинг 3.1 – Підключення необхідних Python-бібліотек

```
import pandas as pd
import numpy as np
from sklearn.ensemble import RandomForestClassifier
from sklearn.model_selection import GridSearchCV
from sklearn.model_selection import train_test_split
from sklearn.metrics import classification_report, accuracy_score
from sklearn.preprocessing import StandardScaler, LabelEncoder
```

Тепер розберемо декілька ключових функцій для навчання моделі машинного навчання та збереження навченої моделі у окремий файл.

Лістинг 3.2 – Функція split_data

```
def split_data(data):
```

```

"""
Розділення даних на ознаки та мітки, тренувальну і тестову
вибірки.
"""
X = data.drop(columns=['Label'])
y = data['Label']
X_train, X_test, y_train, y_test = train_test_split(X, y,
test_size=0.3, random_state=42)
print("Дані розділено на тренувальну та тестову вибірки.")
return X_train, X_test, y_train, y_test

```

Ця функція розділяє оброблені дані на ознаки (X) і мітки (y), а потім виконує розділення на тренувальну та тестову вибірки за допомогою функції `train_test_split`. Такий підхід дозволяє оцінити продуктивність моделі на даних, які вона раніше не бачила, забезпечуючи об'єктивність тестування.

Лістинг 3.3 – Функція `train_random_forest`

```

def train_random_forest(X_train, y_train):
    """
    Навчання моделі Random Forest із підбором оптимальних
    параметрів.
    """
    # Набір параметрів для підбору
    param_grid = {
        'n_estimators': [50, 100, 200], # Кількість дерев у лісі
        'max_depth': [None, 10, 20, 30], # Максимальна глибина
        дерева
        'min_samples_split': [2, 5, 10], # Мінімальна кількість
        зразків для поділу вузла
        'min_samples_leaf': [1, 2, 4], # Мінімальна кількість
        зразків у листі
        'bootstrap': [True, False] # Чи використовувати метод
        підстановки
    }

    # Ініціалізація класифікатора Random Forest
    rf = RandomForestClassifier(random_state=42)

    # Налаштування GridSearchCV для пошуку оптимальних параметрів
    grid_search = GridSearchCV(estimator=rf, param_grid=param_grid,
                               cv=5, scoring='accuracy', n_jobs=-1,
    verbose=2)

    # Підбір параметрів на тренувальних даних
    grid_search.fit(X_train, y_train)

    print("Оптимальні параметри для Random Forest:",
    grid_search.best_params_)

```

```

print("Найкраща точність на крос-валідації:",
grid_search.best_score_)

# Отримання моделі з найкращими параметрами
best_rf = grid_search.best_estimator_

print("Модель Random Forest із оптимальними параметрами
навчено.")
return best_rf

```

Функція виконує навчання моделі Random Forest на тренувальних даних. Random Forest створює кілька дерев рішень і використовує їхній ансамбль для підвищення точності класифікації. Навчена модель повертається для подальшого використання.

Лістинг 3.4 – Функція save_model

```

def save_model(model, file_name):
    """
    Збереження моделі у файл.
    """
    joblib.dump(model, file_name)
    print(f"Модель збережено у файл {file_name}.")

```

Ця функція зберігає навчений алгоритм у файл за допомогою бібліотеки joblib. Це дозволяє уникнути повторного навчання моделі, зберігаючи час і ресурси при подальшому використанні.

Лістинг 3.5 – Функція evaluate_model

```

def evaluate_model(model, X_test, y_test):
    """
    Оцінка продуктивності моделі.
    """
    y_pred = model.predict(X_test)
    accuracy = accuracy_score(y_test, y_pred)
    print(f"Точність моделі: {accuracy:.2f}")
    print("Класифікаційний звіт:")
    print(classification_report(y_test, y_pred))

```

Функція оцінює продуктивність моделі Random Forest на тестовій вибірці. Вона обчислює точність класифікації та створює класифікаційний звіт, який містить метрики точності, повноти й F_1 -міри для кожного класу.

Точність визначає, наскільки модель є правильною серед позитивних передбачень. Вона обчислюється як відношення правильно класифікованих позитивних зразків до всіх зразків, передбачених як позитивні. Висока точність вказує на те, що модель рідко помиляється, класифікуючи нормальні дані як аномальні або навпаки.

Повнота визначає, наскільки добре модель знаходить всі позитивні зразки. Вона обчислюється як відношення правильно класифікованих позитивних зразків до всіх фактично позитивних зразків. Висока повнота свідчить про те, що модель рідко пропускає аномалії чи інші важливі події.

F_1 -міра є гармонічним середнім між точністю та повнотою. Вона забезпечує збалансовану оцінку, особливо коли є дисбаланс між позитивними та негативними класами. А функція `classification_report` автоматично обчислює ці метрики для кожного класу та узагальнює їх середнє значення (макро- або мікросереднє). Це дає змогу отримати всебічну оцінку продуктивності моделі, що є важливим для розуміння її здатності класифікувати дані у складних умовах, наприклад, при дисбалансі класів у наборі даних.

```
polishchukv@polishchukv-VirtualBox:~/Desktop/dataset/archive$ python3 ../mLIDS_model.py
Завантаження даних з Friday1.csv...
Завантаження даних з Friday2.csv...
Завантаження даних з Friday3.csv...
Завантаження даних з Monday1.csv...
Завантаження даних з Tuesday1.csv...
Завантаження даних з Thursday1.csv...
Завантаження даних з Thursday2.csv...
Завантаження даних з Wednesday1.csv...
Об'єднаний датасет має 2830743 рядків і 79 колонок.
Мітки перетворено у числовий формат.
Дані нормалізовано.
Дані розділено на тренувальну та тестову вибірки.
Модель Random Forest навчено.
Модель збережено у файл random_forest_model.pkl.
Точність моделі: 0.97
Класифікаційний звіт:
```

	precision	recall	f1-score	support
0	0.98	0.96	0.97	423456
1	0.96	0.98	0.97	358743
accuracy			0.97	782199
macro avg	0.97	0.97	0.97	782199
weighted avg	0.97	0.97	0.97	782199

```
polishchukv@polishchukv-VirtualBox:~
```

Рисунок 3.3 – Вивід програми для навчання моделі

Далі потрібно розробити програму (лістинг якої викладено у додатку В), що забезпечує аналіз мережевого трафіку в реальному часі, використовуючи навчений алгоритм машинного навчання. Основною функцією є ідентифікація аномалій у трафіку та їх запис у лог-файл. Використовується навчена модель Random Forest, яка завантажується з файлу `random_forest_model.pkl`, щоб оперативно виконувати аналіз без потреби в повторному навчанні.

Перший етап — завантаження моделі. Функція `load_model` використовує бібліотеку `joblib` для завантаження навченої моделі. Це дозволяє працювати з раніше створеною та збереженою моделлю, що суттєво економить час на етапі ініціалізації. Завантажена модель використовується для аналізу характеристик кожного мережевого пакета.

Для підготовки кожного пакета до аналізу створена функція `preprocess_packet`. Вона витягує основні параметри пакета, такі як довжина пакета, порти джерела та призначення, TTL (час життя пакета) і протокол. Ці параметри об'єднуються в таблицю `pandas` для подальшого передавання в модель.

Функція `analyze_packet` виконує безпосередній аналіз пакета (лістинг 3.6). Після обробки параметри пакета нормалізуються за допомогою `StandardScaler`. Це гарантує, що значення ознак перебувають у прийнятному діапазоні, що відповідає тому, який використовувався під час навчання моделі. Потім модель передбачає, чи є цей пакет аномальним. Якщо результатом є клас 1, це означає, що пакет визнаний аномалією.

Лістинг 3.6 – Функція `analyze_packet`

```
def analyze_packet(packet, model, scaler):  
    """  
    Аналіз мережевого пакета за допомогою моделі.  
    """  
    if IP not in packet:  
        return # Ігнорувати пакети, які не містять IP  
  
    features_df = preprocess_packet(packet)  
    scaled_features = scaler.transform(features_df)  
  
    prediction = model.predict(scaled_features)[0]
```



```

if prediction == 1: # 1 вказує на аномалію
    print("Аномалія виявлена!")
    log_anomaly(packet, prediction)

```

Важливою частиною функціоналу є логування аномалій. Функція `log_anomaly` записує всі виявлені аномалії у лог-файл `real_time_anomalies.log`. Лог містить часову мітку, короткий опис пакета та клас передбачення. Це дозволяє аналітикам безпеки мати детальний журнал для подальшого аналізу мережевих інцидентів.

Реалізація аналізу трафіку в реальному часі здійснюється за допомогою бібліотеки `scapy`. Функція `sniff` перехоплює мережеві пакети, застосовуючи до кожного з них функцію `analyze_packet`. Фільтр IP гарантує, що аналізуються тільки пакети, які містять IP-протокол.

Програма також включає механізм масштабування ознак за допомогою `StandardScaler`. Цей процес забезпечує сумісність нових даних із даними, використаними під час тренування моделі. Нормалізація є критично важливою для коректної роботи моделі, оскільки без цього передбачення можуть бути некоректними.

Загалом програма надає ефективний інструмент для виявлення аномалій у мережевому трафіку. Вона об'єднує можливості машинного навчання та аналізу в режимі реального часу, що робить її цінним рішенням для моніторингу безпеки мережі. Усі виявлення автоматично записуються до лог-файлу, що дозволяє аналізувати мережеві інциденти та оперативно реагувати на загрози.

3.4 Тестування системи шляхом імітації атак

Для перевірки ефективності виявлення вторгнень обрано декілька типів сканування порта 22 (SSH), зокрема `TCP Scan`, `XMAS Scan` та `FIN Scan`. Ці типи сканування є поширеними методами мережевої розвідки, що використовуються зловмисниками для виявлення активних портів і потенційно вразливих сервісів.

Програма із залученням машинного навчання працює у реальному часі, аналізуючи вхідний мережевий трафік за допомогою попередньо навченої

моделі Random Forest. Завдяки витягуванню ключових параметрів з мережесих пакетів (наприклад, довжина пакета, порти джерела та призначення, TTL і протокол), модель може ідентифікувати пакети, які відповідають шаблонам аномалій. У випадку виявлення аномального пакета, інформація про нього записується в лог-файл разом із часовою міткою і деталями про сам пакет (рис. 3.4).

```
polishchukv@polishchukv-VirtualBox:~/Desktop/dataset/archive$ python3 ../anomaly_monitoring.py
Завантаження моделі...
Модель завантажено з random_forest_model.pkl.
Підготовка нормалізатора...
Початок аналізу мережевого трафіку...
Аномалія виявлена!
[2024-12-26 04:36:17] Аномалія виявлена!
Деталі пакету: IP / TCP 192.168.0.101:36720 > 192.168.0.103:22
Клас передбачення: 1

Аномалія виявлена!
[2024-12-26 04:36:40] Аномалія виявлена!
Деталі пакету: IP / TCP 192.168.0.101:35174 > 192.168.0.103:22
Клас передбачення: 1

Аномалія виявлена!
[2024-12-26 04:36:50] Аномалія виявлена!
Деталі пакету: IP / TCP 192.168.0.101:51311 > 192.168.0.103:22
Клас передбачення: 1

Аномалія виявлена!
[2024-12-26 04:36:57] Аномалія виявлена!
Деталі пакету: IP / TCP 192.168.0.101:40285 > 192.168.0.103:22
Клас передбачення: 1
```

Рисунок 3.4 – Виявлення аномалій за допомогою машинного навчання

Ці самі події також були виявлені системою Snort, що показано на рисунку 3.5. Snort є потужним інструментом для виявлення вторгнень (IDS), який здатний фіксувати схожі типи мережесих атак. У логах Snort можна побачити класифікацію цих подій, таких як TCP Scan і XMAS Scan, а також їхній рівень пріоритету. Це підтверджує ефективність обох систем у виявленні спроб мережевої розвідки.

```

Preprocessor Object: SF_SSH Version 1.1 <Build 3>
Preprocessor Object: SF_MODBUS Version 1.1 <Build 1>
Preprocessor Object: SF_SDF Version 1.1 <Build 1>
Preprocessor Object: SF_DNS Version 1.1 <Build 4>
Commencing packet processing (pid=4556)
12/26-04:36:17.246384  [**] [1:10000005:2] NMAP TCP Scan [**] [Priority: 0] {TCP} 192.168.0.101:36720 -> 192.168.0.103:22
12/26-04:36:17.246547  [**] [1:10000005:2] NMAP TCP Scan [**] [Priority: 0] {TCP} 192.168.0.101:36720 -> 192.168.0.103:22
12/26-04:36:17.247025  [**] [1:10000005:2] NMAP TCP Scan [**] [Priority: 0] {TCP} 192.168.0.101:36720 -> 192.168.0.103:22
12/26-04:36:40.109578  [**] [1:10000005:2] NMAP TCP Scan [**] [Priority: 0] {TCP} 192.168.0.101:35172 -> 192.168.0.103:22
12/26-04:36:40.109578  [**] [1:1228:7] SCAN nmap XMAS [**] [Classification: Attempted Information Leak] [Priority: 2] {TCP} 192.168.0.101:35172 -> 192.168.0.103:22
12/26-04:36:40.209580  [**] [1:10000005:2] NMAP TCP Scan [**] [Priority: 0] {TCP} 192.168.0.101:35174 -> 192.168.0.103:22
12/26-04:36:40.209580  [**] [1:1228:7] SCAN nmap XMAS [**] [Classification: Attempted Information Leak] [Priority: 2] {TCP} 192.168.0.101:35174 -> 192.168.0.103:22
12/26-04:36:50.256709  [**] [1:10000005:2] NMAP TCP Scan [**] [Priority: 0] {TCP} 192.168.0.101:51311 -> 192.168.0.103:22
12/26-04:36:50.256709  [**] [1:621:7] SCAN FIN [**] [Classification: Attempted Information Leak] [Priority: 2] {TCP} 192.168.0.101:51311 -> 192.168.0.103:22
12/26-04:36:50.357049  [**] [1:10000005:2] NMAP TCP Scan [**] [Priority: 0] {TCP} 192.168.0.101:51313 -> 192.168.0.103:22
12/26-04:36:50.357049  [**] [1:621:7] SCAN FIN [**] [Classification: Attempted Information Leak] [Priority: 2] {TCP} 192.168.0.101:51313 -> 192.168.0.103:22
12/26-04:36:57.137344  [**] [1:10000005:2] NMAP TCP Scan [**] [Priority: 0] {TCP} 192.168.0.101:40285 -> 192.168.0.103:22
12/26-04:36:57.238007  [**] [1:10000005:2] NMAP TCP Scan [**] [Priority: 0] {TCP} 192.168.0.101:40287 -> 192.168.0.103:22

```

Рисунок 3.5 – Реєстрування атаки за допомогою Snort

Фіксація подій, таких як XMAS і FIN-сканування, є особливо важливою, оскільки ці методи є частиною розширеного арсеналу зловмисників. Вони часто використовуються для обходу стандартних правил брандмауера або IDS. Наприклад, XMAS-сканування надсилає пакети зі специфічними прапорцями, щоб спробувати отримати інформацію про те, які порти активні. У свою чергу, FIN-сканування може допомогти приховати активність зловмисника в мережі.

Однак є відмінності між Snort і системою на базі машинного навчання. Snort покладається на сигнатури — заздалегідь відомі шаблони атак, тоді як модель машинного навчання здатна виявляти нові, раніше невідомі аномалії, якщо вони істотно відрізняються від нормального трафіку. Це робить модель корисною для виявлення атак, які ще не додані в бази сигнатур.

Результати тестування системи виявлення вторгнень, яка поєднує сигнатурний підхід (Snort) і поведінковий аналіз на основі машинного навчання, підтвердили ефективність такого комбінованого підходу. Snort забезпечує точне виявлення відомих атак, використовуючи бази сигнатур, і дозволяє миттєво реагувати на загрози, які вже мають зафіксовані шаблони. З іншого боку, модель машинного навчання доповнює Snort, виявляючи аномалії в мережевому трафіку, які не відповідають жодному відомому шаблону, таким чином, розширюючи можливості системи для роботи з новими або модифікованими типами атак.

Комбіноване використання цих підходів дозволяє ефективно підтверджувати підозрілий трафік. Наприклад, коли Snort фіксує подію на основі сигнатури, модель машинного навчання може додатково перевірити, чи є цей трафік типовим для нормальної поведінки мережі. Це значно зменшує ймовірність правдиво негативних (True Negative) або хибно позитивних (False Positive) спрацьовувань, забезпечуючи кращу точність у виявленні загроз.

Система продемонструвала, що поєднання двох підходів мінімізує шанс правдиво негативних (True Negative) спрацьовувань, оскільки сигнатурний аналіз фокусується на відомих загрозах, а поведінковий аналіз забезпечує виявлення нетипового трафіку. Це особливо корисно для сучасних мереж, де атаки стають дедалі складнішими та частіше включають невідомі техніки, які можуть пройти непоміченими у звичайних сигнатурних системах.

Підсумовуючи, тестування підтвердило, що така інтегрована система виявлення вторгнень забезпечує комплексний захист мережі, поєднуючи швидкість і точність сигнатурного аналізу з гнучкістю і здатністю до навчання поведінкових моделей. Це робить її ефективним інструментом для виявлення як традиційних, так і нових атак, а також забезпечує високу надійність і мінімізує ризик пропуску загроз у реальному часі.

РОЗДІЛ 4 ОХОРОНА ПРАЦІ ТА БЕЗПЕКА В НАДЗВИЧАЙНИХ СИТУАЦІЯХ

4.1 Охорона праці

У сучасних умовах повномасштабної війни, розв'язаної росією проти України, питання кібербезпеки набула особливої гостроти. Українські державні установи, критична інфраструктура та комерційні організації щоденно зазнають масованих кібератак з боку ворожих хакерських угруповань. Ворог використовує методи, спрямовані на дестабілізацію систем управління, викрадення даних, блокування роботи важливих інформаційних ресурсів та завдання шкоди інфраструктурі. У цих умовах впровадження ефективних систем виявлення вторгнень у мережевому трафіку стало не лише технічною необхідністю, а й питанням національної безпеки.

Системи виявлення вторгнень (IDS) відіграють критичну роль у виявленні аномальної активності та попередженні атак на інформаційні системи. Вони дозволяють своєчасно ідентифікувати підозрілу поведінку у мережевому трафіку, що може свідчити про спроби вторгнення чи зловмисну активність. Особливо важливою є здатність таких систем адаптуватися до нових загроз, які постійно еволюціонують, що вимагає використання сучасних методів аналізу аномалій, зокрема машинного навчання. У контексті війни проти України, де кібератаки часто координуються з фізичними діями на фронті, ефективність IDS є одним із ключових факторів для забезпечення стабільності та працездатності цифрових систем держави та бізнесу.

Важливою складовою впровадження систем виявлення вторгнень є дотримання норм охорони праці під час експлуатації апаратури та програмного забезпечення. Згідно з ДСанПІН 3.3.2.007-98, для працівників, які займаються впровадженням та підтримкою IDS, повинні бути розроблені інструкції з охорони праці з урахуванням специфіки робіт. Інструкції мають містити чіткі вимоги щодо безпеки під час роботи з електронним обладнанням, програмними інструментами та мережевими ресурсами. До них належить обов'язкове

ознайомлення з правилами безпеки, періодичне проведення інструктажів, а також регулярний моніторинг робочих умов для своєчасного виявлення порушень.[20]

Особливої уваги потребує техніка безпеки при роботі з комп'ютерним обладнанням і мережею. Усі робочі місця повинні відповідати вимогам безпечного розміщення технічних засобів, включаючи захист від можливих електричних уражень та перегріву обладнання. Зокрема, слід забезпечити використання захисного заземлення, пристроїв автоматичного вимикання в разі короткого замикання чи перевантаження системи, а також регулярну перевірку електричної мережі на відповідність нормам безпеки. Згідно з нормами, передбаченими для електрообладнання ДСН 3.3.6.042-99, всі кабелі повинні мати відповідну ізоляцію, а працівники повинні користуватися засобами індивідуального захисту при роботі з електронними пристроями.[21]

Важливим аспектом є протипожежна безпека приміщень, у яких розміщується апаратура систем виявлення вторгнень. Згідно з правилами пожежної безпеки для підприємств та організацій, необхідно забезпечити наявність первинних засобів пожежогасіння (вогнегасників вуглекислотного або порошкового типу) поруч із серверними кімнатами та робочими місцями. Крім того, повинна бути встановлена система автоматичного пожежогасіння, що використовує газові склади або порошкові склади для мінімізації шкоди обладнанню у разі пожежі. Обов'язковим є проведення навчань з евакуації персоналу та регулярні перевірки протипожежного інвентарю.[22]

Організація вентиляції та кондиціонування повітря є важливим фактором для підтримання стабільної роботи комп'ютерних систем. Наявність сучасних систем охолодження запобігає перегріванню апаратури, що може призвести до її виходу з ладу або короткого замикання. Відповідно до чинних норм [22], у приміщеннях серверних кімнат необхідно підтримувати температуру в діапазоні 18-25°C та забезпечити безперебійну циркуляцію повітря. Також рекомендовано використовувати автономні системи живлення для кондиціонування з метою підтримання необхідних умов навіть у разі аварійного відключення електроенергії.

Працівники, які обслуговують системи виявлення вторгнень, повинні бути ознайомлені з особливостями використання електрообладнання та ризиками, пов'язаними з його експлуатацією. Згідно з вимогами Про затвердження Порядку проведення медичних оглядів працівників певних категорій, важливо забезпечити регулярні інструктажі та навчання з охорони праці для всіх фахівців, які працюють у сфері кібербезпеки. Особливу увагу слід приділити правилам дій у разі аварійних ситуацій, як-от коротке замикання, відключення системи живлення або виявлення неполадок у роботі апаратури.[22]

Додаткову увагу слід приділити ергономічному облаштуванню робочих місць. Робоче місце оператора IDS має відповідати нормам безпеки праці, включаючи зручне розташування моніторів, клавіатур та інших пристроїв. Згідно з чинними вимогами охорони праці, слід передбачити регулярні перерви для працівників, які тривалий час працюють за комп'ютером, з метою зниження втоми та напруги зору. Також необхідно забезпечити наявність спеціальних підставок для моніторів, які дозволяють регулювати висоту та кут нахилу екрана.[21]

Особливе значення має забезпечення захисту від статичної електрики, яка може накопичуватися у приміщеннях з комп'ютерним обладнанням. Відповідно до встановлених правил ДСН 3.3.6.042-99, необхідно використовувати антистатичні покриття для підлоги, заземлені стійки та додаткові пристрої нейтралізації статичних зарядів. Окрім цього, слід регулярно контролювати рівень статичної електрики у приміщенні та проводити технічне обслуговування засобів захисту для запобігання виходу з ладу обладнання.[21]

Таким чином, дотримання вимог охорони праці, техніки безпеки та протипожежної безпеки є важливою умовою успішного впровадження систем виявлення вторгнень. Виконання зазначених заходів сприятиме підвищенню ефективності функціонування IDS, забезпеченню безпеки працівників і стабільній роботі критичних інформаційних систем держави.

4.2 Безпека в надзвичайних ситуаціях

В умовах військових дій об'єкти приладобудівної галузі стають потенційними цілями для фізичних атак, кіберзагроз і деструктивного впливу. Першочерговим стратегічним завданням для забезпечення економічної, технологічної та оборонної спроможності України є підвищення стійкості таких об'єктів. Приладобудування є високотехнологічною та найбільш наукомісткою галуззю машинобудування, що спеціалізується на виробництві засобів виміру, аналізу, обробки і представлення інформації, автоматичних та автоматизованих систем управління, а також пристроїв регулювання [25]. Здатність підприємств цієї галузі стабільно функціонувати під час війни визначає рівень забезпечення критичних потреб держави, включаючи виробництво оборонної техніки, підтримання роботи стратегічної інфраструктури та сприяння економічної стабільності в країні.

Фізичні загрози включають авіаудари, ракетні обстріли, диверсії, спрямовані на руйнування виробничої інфраструктури, складів сировини, обладнання та логістичних ланцюгів. Пошкодження таких об'єктів зупиняє виробничі процеси, а також призводить до значних економічних втрат і технологічної деградації. Ракетні удари по об'єктах критичної інфраструктури під час війни в Україні призвели до тимчасового паралічу виробничих потужностей, зокрема у машинобудівній та енергетичній галузях. За даними ДСНС у 2023 році в Україні відбулося 109 надзвичайних ситуацій (для порівняння 66 у 2022 році), з яких – 60 природного та 48 техногенного характеру, зокрема за масштабами: 4 державного рівня, 5 – регіонального рівня, 54 місцевого та 46 об'єктового рівня. Серед надзвичайних ситуацій державного рівня воєнного характеру власне саме руйнування та пошкодження критичної інфраструктури: з початку 2022 року пошкоджено, зруйновано понад 214 тисяч її об'єктів (8 тис. 642 – це об'єкти життєзабезпечення, 1 тис. 592 об'єктів транспортної інфраструктури, 3 тис. 679 освітніх закладів та 1 тис. 569 медичних закладів тощо) [23].

Крім фізичних атак, об'єкти приладобудівної галузі зазнають масштабних кібератак, які стають невід'ємною складовою сучасних військових конфліктів. Використання шкідливого програмного забезпечення, спрямованого на порушення роботи автоматизованих систем управління, викрадення конфіденційної інформації чи шифрування важливих даних значно впливає на виробничі процеси. Відомі приклади таких атак, як поширення вірусу NotPetya, продемонстрували масштабність наслідків від застосування кіберзброї проти промислових об'єктів, що створює додаткові виклики для забезпечення стійкості галузі.

Ураховуючи специфіку роботи приладобудівної галузі, важливим аспектом є ризики, пов'язані з порушенням роботи систем постачання. Транспортна інфраструктура під час війни стає критично вразливою, що ускладнює доставку необхідних матеріалів і компонентів. Це впливає на цикли виробництва, особливо якщо ланцюги постачання не мають резервних механізмів або альтернативних маршрутів.

Зважаючи на важливість приладобудівної галузі для стратегічної стійкості України, впровадження науково обґрунтованих підходів до забезпечення безпеки є вкрай актуальним. Інтеграція фізичних, кібернетичних і інформаційних заходів захисту дозволить мінімізувати ризики зупинки роботи об'єктів та забезпечити їхню здатність до оперативного відновлення функціонування навіть у критичних умовах. Слід акцентувати на діях, спрямованих на поліпшення структури ресурсного потенціалу машинобудівних підприємств, зокрема приладобудівних. Насамперед це стосується його інноваційної складової, що може надати значну допомогу в підвищенні конкурентоспроможності підприємств і їх продукції, інвестиційній привабливості [24]. Вивчення міжнародного досвіду захисту критичної інфраструктури та адаптація передових технологій до українських реалій є важливим елементом у розробці ефективної стратегії підвищення стійкості приладобудівної галузі.

Фізичний захист об'єктів є однією з найважливіших складових забезпечення їхньої стійкості. Інфраструктура таких підприємств повинна бути

укріпленою, щоб протистояти прямим загрозам, зокрема ракетним ударам чи артилерійським обстрілам. Для цього необхідно будувати бомбосховища та захищені приміщення, як серверні кімнати, які здатні витримувати екстремальні впливи. Захищені сховища даних і приміщення для збереження ключового обладнання допомагають мінімізувати втрати навіть у разі пошкодження інших частин об'єкта. Наприклад, заводи, які забезпечують оборонну промисловість, можуть бути оснащені мобільними укриттями для персоналу та техніки, що вже довело свою ефективність у країнах, які стикалися з подібними викликами.

Одним із важливих заходів підвищення стійкості є впровадження систем відеоспостереження та датчиків вторгнення. Вони дозволяють своєчасно виявляти небезпеку та координувати дії для запобігання вторгненню. Такі системи мають працювати в інтеграції з аварійними сигналізаціями, що дозволяє миттєво інформувати відповідні служби та персонал. У військових зонах Ізраїлю впровадження систем раннього виявлення вторгнень, які включають мережу камер відеоспостереження, тепловізорів і датчиків руху, значно зменшило втрати від диверсій та атак. Зокрема, на кордоні з Сектором Газа та іншими потенційно небезпечними регіонами ці системи дозволяють ізраїльській армії та силам безпеки оперативно виявляти рух підозрілих осіб чи об'єктів, оцінювати рівень загрози в режимі реального часу та координувати дії відповідних підрозділів. Інтеграція таких систем із центрами управління та безпілотними літальними апаратами (БПЛА) забезпечує ще більш точне реагування та мінімізацію ризиків для населення і критичної інфраструктури. Цей підхід продемонстрував свою ефективність завдяки поєднанню високої технологічності обладнання та ретельно налагоджених процедур реагування. Застосування ізраїльського досвіду в Україні може суттєво посилити рівень безпеки об'єктів приладобудівної галузі в умовах війни.

До того ж системи моніторингу стану об'єктів у реальному часі дозволяють відслідковувати аномалії в роботі обладнання та своєчасно вживати заходів. Інтеграція штучного інтелекту забезпечує проактивний захист від загроз, адже алгоритми можуть аналізувати великі обсяги даних і виявляти потенційні небезпеки ще до їх реалізації.

Автономне електроживлення також відіграє критичну роль у забезпеченні безперервності роботи об'єктів. У випадках пошкодження енергомереж генератори, сонячні батареї чи інші альтернативні джерела енергії дозволяють підтримувати роботу виробничого обладнання та систем моніторингу. Підприємства в Україні вже активно впроваджують генераторні установки для резервного живлення, що дозволяє уникнути повної зупинки виробництва навіть під час перебоїв із постачанням електроенергії.

Для підвищення стійкості роботи об'єктів приладобудівної галузі важливо розробити чіткі кризові плани дій. Такі плани повинні включати протоколи реагування на надзвичайні ситуації, зокрема організацію резервних каналів зв'язку, які забезпечують безперебійне управління навіть у разі пошкодження основних мереж. Евакуація співробітників також має бути передбачена для збереження життя та здоров'я людей, що працюють на підприємствах. Використання мобільних додатків для координації евакуаційних дій може значно підвищити ефективність таких заходів.

Навчання персоналу є ще одним ключовим аспектом. Працівники повинні бути обізнані щодо алгоритмів дій у разі загрози. Регулярне проведення симуляцій і тестування відмовостійкості інформаційних і фізичних систем дозволяє виявити слабкі місця в системах захисту. Впровадження тестових сценаріїв на підприємствах оборонної галузі в Європі дозволило скоротити час реагування на загрози.

Не менш важливим є використання блокчейн-рішень для забезпечення цілісності даних. Ця технологія забезпечує неможливість змінювання або втрати даних, що особливо важливо для критичної інфраструктури. Наприклад, впровадження блокчейну для управління виробничими процесами може гарантувати безперервність роботи навіть у разі кібератак.

Досвід інших країн, які стикалися зі схожими викликами, є важливим для адаптації в Україні. Захист критичної інфраструктури в Ізраїлі чи системи мобільних укриттів демонструють ефективність цих рішень у контексті воєнних загроз. Впровадження таких практик в українських реаліях сприятиме зниженню втрат і підвищенню загальної стійкості об'єктів приладобудівної галузі.

Отже, підвищення стійкості роботи об'єктів приладобудівної галузі у воєнний час вимагає комплексного підходу, що охоплює фізичний захист, резервування ресурсів, розробку кризових планів і впровадження сучасних технологій. Це забезпечить стабільну роботу підприємств, а також підвищить загальну обороноздатність країни в умовах надзвичайних ситуацій.

ВИСНОВКИ

У сучасному цифровому середовищі, де мережі є ключовим елементом комунікації та функціонування інформаційних систем, питання безпеки стає одним із найважливіших викликів. Проведене дослідження повністю підтвердило актуальність теми забезпечення захисту від мережових атак, таких як DDoS, MITM, фішинг та експлуатація програмних вразливостей. Зростання кількості атак і їхня складність зумовлюють необхідність впровадження новітніх рішень, які поєднують традиційні та інноваційні підходи до виявлення вторгнень.

Мета роботи, яка полягала у розробці ефективної системи виявлення вторгнень, що базується на аналізі аномалій у мережевому трафіку з використанням методів машинного навчання, була досягнута. У процесі дослідження виконано поставлені завдання. Зокрема, було проведено аналіз статистичних даних про мережеві атаки, що дозволило визначити ключові тенденції їх розвитку. Це дало змогу ідентифікувати типи загроз, які є найактуальнішими для сучасного інформаційного середовища.

Також досліджено сучасні інструменти та технології для протидії мережевим атакам. Аналіз переваг і недоліків таких рішень, як сигнатурні системи виявлення вторгнень (IDS) і поведінкові моделі, підтвердив доцільність інтеграції методів машинного навчання для підвищення ефективності захисту. У результаті роботи визначено виклики, які постають перед сучасними системами безпеки, зокрема недостатня адаптивність до нових загроз, складність обробки великих обсягів даних і потреба в швидкому реагуванні на інциденти.

Наукова новизна роботи підтверджена результатами тестування системи виявлення вторгнень скануванням вразливих мережових портів. Ефективність виявлення системи машинного навчання виявилась не меншою за сигнатурний підхід, який пропонує Snort. Таким чином змодельована система, в якій модель машинного навчання може виступати як самостійна IDS, а також як альтернативний спосіб виявлення, який може бути використаний аналітиками для підтвердження правдивості спрацювання.

У процесі розробки системи були створені механізми, що дозволяють інтегрувати її з корпоративними та публічними мережами. Запропонована модель успішно протестована, продемонструвавши високу точність і адаптивність. Вона забезпечує баланс між ефективністю, швидкістю реагування та доступністю, що робить її особливо цінною для малих і середніх підприємств, які часто стають мішенню атак через обмежені ресурси.

Результати роботи підкреслюють важливість інтеграції сигнатурних і поведінкових підходів. Така комбінація мінімізує ймовірність хибних спрацьовувань і забезпечує виявлення загроз як відомого, так і невідомого характеру. Завдяки цьому запропонована система дозволяє не лише знижувати ризики успішних атак, але й оптимізувати процеси моніторингу й реагування.

Таким чином, виконана робота робить значний внесок у розвиток сучасних систем виявлення вторгнень, демонструючи переваги поєднання методів машинного навчання та аналізу аномалій. Подальші дослідження можуть бути спрямовані на розширення функціоналу системи, оптимізацію її продуктивності та інтеграцію з хмарними платформами, що забезпечить масштабованість і глобальний захист мереж. Результати роботи мають практичне значення для вдосконалення систем кібербезпеки як у корпоративних, так і у державних структурах.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Cost of a data breach 2024 | IBM. IBM - United States. URL: <https://www.ibm.com/reports/data-breach> (дата звернення: 26.10.2024).
2. What is an intrusion detection system (IDS)? + best IDS tools | upguard. Third-Party Risk and Attack Surface Management Software | UpGuard. URL: <https://www.upguard.com/blog/intrusion-detection-system> (дата звернення: 26.10.2024).
3. Tymoshchuk, V., Vorona, M., Dolinskyi, A., Shymanska, V., & Tymoshchuk, D. (2024). SECURITY ONION PLATFORM AS A TOOL FOR DETECTING AND ANALYSING CYBER THREATS. Collection of scientific papers «ΛΟΓΟΣ», (December 13, 2024; Zurich, Switzerland), 232-237.
4. Top 20 most common types of cyber attacks | fortinet. Fortinet. URL: <https://www.fortinet.com/resources/cyberglossary/types-of-cyber-attacks> (дата звернення: 29.10.2024).
5. Tymoshchuk, D., & Yatskiv, V. (2024). Slowloris ddos detection and prevention in real-time. Collection of scientific papers «ΛΟΓΟΣ», (August 16, 2024; Oxford, UK), 171-176.
6. 15 cybersecurity tools for small and medium businesses (smbs) | fortinet. Fortinet. URL: <https://www.fortinet.com/resources/cyberglossary/smb-cybersecurity-tools> (дата звернення: 15.11.2024).
7. Tymoshchuk, V., Mykhailovskyi, O., Dolinskyi, A., Orlovska, A., & Tymoshchuk, D. (2024). OPTIMISING IPS RULES FOR EFFECTIVE DETECTION OF MULTI-VECTOR DDOS ATTACKS. Матеріали конференцій МЦНД, (22.11. 2024; Біла Церква, Україна), 295-300.
8. IBM. What is an intrusion detection system (IDS)? | IBM. IBM - United States. URL: <https://www.ibm.com/think/topics/intrusion-detection-system> (дата звернення: 10.11.2024).

9. Tymoshchuk, V., Vantsa, V., Karnaukhov, A., Orlovska, A., & Tymoshchuk, D. (2024). COMPARATIVE ANALYSIS OF INTRUSION DETECTION APPROACHES BASED ON SIGNATURES AND ANOMALIES. Матеріали конференцій МЦНД, (29.11. 2024; Житомир, Україна), 328-332.
10. Network traffic anomaly detection with machine learning. Eyer - headless AIOps. URL: <https://eyer.ai/blog/network-traffic-anomaly-detection-with-machine-learning/> (дата звернення: 15.11.2024).
11. Quiroz-Vázquez C. Anomaly detection in machine learning: examples, applications & use cases | IBM. IBM - United States. URL: <https://www.ibm.com/think/topics/machine-learning-for-anomaly-detection> (дата звернення: 15.11.2024).
12. Tymoshchuk, V., Pakhoda, V., Dolinskyi, A., Karnaukhov, A., & Tymoshchuk, D. (2024). MODELLING CYBER THREATS AND EVALUATING THE PERFORMANCE OF INTRUSION DETECTION SYSTEMS. Grail of Science, (46), 636–641. <https://doi.org/10.36074/grail-of-science.29.11.2024.081>
13. What is generative AI in cybersecurity?. Palo Alto Networks. URL: <https://www.paloaltonetworks.com/cyberpedia/generative-ai-in-cybersecurity> (дата звернення: 22.11.2024).
14. Tymoshchuk, D., Yasniy, O., Mytnyk, M., Zagorodna, N., Tymoshchuk, V., (2024). Detection and classification of DDoS flooding attacks by machine learning methods. CEUR Workshop Proceedings, 3842, pp. 184 - 195.
15. Rigatti S. J. Random Forest. Journal of Insurance Medicine. 2017. Т. 47, № 1. С. 31–39. URL: <https://doi.org/10.17849/insm-47-01-31-39.1> (дата звернення: 26.11.2024).
16. Lupa, B., Horyn, I., Zagorodna, N., Tymoshchuk, D., Lechachenko T., (2024). Comparison of feature extraction tools for network traffic data. CEUR Workshop Proceedings, 3896, pp. 1-11.
17. Xie Y., Yu S. A deep learning approach to network intrusion detection // IEEE Transactions on Emerging Topics in Computing. 2015.

18. ТИМОЩУК, Д., & ЯЦКІВ, В. (2024). USING HYPERVISORS TO CREATE A CYBER POLYGON. MEASURING AND COMPUTING DEVICES IN TECHNOLOGICAL PROCESSES, (3), 52-56.

19. Sommer R., Paxson V. Outside the closed world: On using machine learning for network intrusion detection // IEEE Symposium on Security and Privacy. 2010.

20. ТИМОЩУК, Д., ЯЦКІВ, В., ТИМОЩУК, В., & ЯЦКІВ, Н. (2024). INTERACTIVE CYBERSECURITY TRAINING SYSTEM BASED ON SIMULATION ENVIRONMENTS. MEASURING AND COMPUTING DEVICES IN TECHNOLOGICAL PROCESSES, (4), 215-220.

21. Про затвердження державних санітарних правил і норм роботи з візуальними дисплейними терміналами електронно-обчислювальних машин : Наказ Міністерства охорони здоров'я України від 10.12.1998 № 7. URL: <https://zakon.rada.gov.ua/rada/show/v0007282-98#Text> (дата звернення: 17.12.2024).

22. Про затвердження порядку проведення медичних оглядів працівників певних категорій : Наказ Міністерства охорони здоров'я України від 21.05.2007 № 246. URL: <https://zakon.rada.gov.ua/laws/show/z0846-07#Text> (дата звернення: 17.12.2024).

23. Інформаційно-аналітична довідка про надзвичайні ситуації в Україні у 2023 році. Вебсайт ДСНС України. URL: <https://dsns.gov.ua/upload/2/0/2/2/3/2/1/2023-rik.pdf>

24. Матвійчук І.О. Сучасний стан та перспективи розвитку приладобудування в Україні / І. О. Матвійчук // Глобальні та національні проблеми економіки. - 2015. - Вип. 3. - С. 360-365

25. Приладобудування // Вікіпедія – вільна енциклопедія [Електронний ресурс]. – URL: <http://uk.wikipedia.org/wiki/Приладобудування>.

ДОДАТОК А ПУБЛІКАЦІЯ

**МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
ТЕРНОПІЛЬСЬКИЙ НАЦІОНАЛЬНИЙ ТЕХНІЧНИЙ УНІВЕРСИТЕТ
імені ІВАНА ПУЛЮЯ
НАУКОВЕ ТОВАРИСТВО ім. ШЕВЧЕНКА**

М А Т Е Р І А Л И

ХІІ НАУКОВО-ТЕХНІЧНОЇ КОНФЕРЕНЦІЇ

**«ІНФОРМАЦІЙНІ МОДЕЛІ,
СИСТЕМИ ТА ТЕХНОЛОГІЇ»**



18-19 грудня 2024 року

**ТЕРНОПІЛЬ
2024**

ПРОГРАМНИЙ КОМІТЕТ

Голова: Приймак Микола – професор кафедри комп’ютерних систем та мереж, д.т.н., професор.

Співголови: Марущак Павло – проректор з наукової роботи, докт. техн. наук, професор.

Баран Ігор – канд. техн. наук, доцент, декан факультету ФІС.

Науковий секретар: Надія Крива – старший викладач.

Члени: Василь Кривень - завідувач кафедри математичних методів в інженерії д.ф.-м.н., професор; Галина Осухівська – завідувач кафедри комп’ютерних систем та мереж, к.т.н., доцент; Микола Карпінський - професор кафедри кібербезпеки, д.т.н., професор; Жанна Баб’як - завідувач кафедри української та іноземних мов, к.пед. н., доцент; Ярослав Литвиненко – професор кафедри комп’ютерних наук, д.т.н., професор; Михайло Петрик - завідувач кафедри програмної інженерії, д.ф.-м.н., професор; Наталія Загородна – завідувач кафедри кібербезпеки, к.т.н., доцент.

ОРГАНІЗАЦІЙНИЙ КОМІТЕТ

Голова: Скоренький Юрій Любомирович – канд. техн. наук, доцент кафедри фізики.

Члени: доцент кафедри комп’ютерних наук, к.т.н. В. Никитюк; доцент кафедри програмної інженерії, к.т.н. Д. Михалик; доцент кафедри кібербезпеки, к.т.н. М. Стадник; доцент кафедри комп’ютерних систем та мереж, к.т.н. Є. Тиш; ст. викладач Л. Джиджора.

Матеріали XII науково-технічної конфіції «Інформаційні моделі, системи та технології» Тернопільського національного технічного університету імені Івана Пулюя, (Тернопіль, 18–19 грудня 2024 р.). – Тернопіль : Тернопільський національний технічний університет імені Івана Пулюя, 2024.

Адреса оргкомітету: ТНТУ ім. І. Пулюя, м. Тернопіль, вул. Руська, 56, 46001, тел. (0352) 52-41-33, факс (0352) 25-49-83.

E-mail: confis2024@gmail.com

Редагування, оформлення та верстка: Надія Крива

СЕКЦІЇ КОНФЕРЕНЦІЇ, ЯКІ ПРЕДСТВЛЕНІ В ЗБІРНИКУ

- Математичне моделювання
- Інформаційні системи та технології, кібербезпека
- Комп’ютерні системи та мережі
- Програмна інженерія та моделювання складних розподілених систем
- Новітні фізико-технічні та освітні технології

В збірнику надруковано тези доповідей XII науково-технічної конференції «Інформаційні моделі, системи та технології» (Тернопіль, 18–19 грудня 2024 р.) за такими науковими напрямками: математичне моделювання; інформаційні системи та технології, кібербезпека; комп’ютерні системи та мережі; програмна інженерія та моделювання складних розподілених систем; новітні фізико-технічні та освітні технології.

Розрахований на науковців, викладачів та студентів вузів.

За зміст тез та дотримання норм академічної доброчесності відповідальність несе автор.

РОЗРОБКА СИСТЕМИ ВИЯВЛЕННЯ ВТОРГНЕНЬ НА ОСНОВІ АНАЛІЗУ АНОМАЛІЙ У МЕРЕЖЕВому ТРАФІКУ З ВИКОРИСТАННЯМ МАШИННОГО НАВЧАННЯ

Vladyslav Polishchuk, student of gr. СБм-62

DEVELOPMENT OF AN ANOMALY-BASED INTRUSION DETECTION SYSTEM USING MACHINE LEARNING FOR NETWORK TRAFFIC ANALYSIS

Сучасні інформаційні системи стикаються з низкою проблем, серед яких зростання кількості кібератак, еволюція методів атак, а також недостатня адаптивність традиційних систем виявлення вторгнень (IDS). Традиційні IDS здебільшого базуються на сигнатурному або евристичному аналізі, що дозволяє ефективно виявляти лише добре відомі атаки. Їх основним недоліком є нездатність розпізнавати нові, невідомі загрози і частота хибнопозитивних (false positive) спрацьовувань, що ускладнює роботу аналітиків безпеки. Тому розробка системи виявлення вторгнень на основі аналізу аномалій у мережевому трафіку з використанням машинного навчання є надзвичайно актуальною. Такий підхід дозволяє забезпечити адаптацію до нових загроз та ефективне виявлення невідомих атак [1].

Використання машинного навчання для обробки мережевого трафіку дозволяє побудувати моделі, які навчаються розпізнавати відхилення від нормальної поведінки. Алгоритми, такі як кластеризація (наприклад, DBSCAN) чи метод опорних векторів (SVM), успішно застосовуються для аналізу великих обсягів мережевих даних [2].

Система виявлення вторгнень була реалізована з використанням TensorFlow, Snort та Python-бібліотеки Scikit-learn для обробки даних. Тестування системи виконано на основі набору даних CICIDS2017, який містить як звичайні сесії, так і сесії з атаками, які відображає відповідний мережевий трафік. Розроблена система дозволила ефективно виявляти аномалії у мережевому трафіку, зокрема атаки типу Brute Force, DoS та DDoS, з високою точністю, забезпечуючи адаптивний підхід до нових загроз.

Загалом, розвиток технологій обробки даних і штучного інтелекту відкриває нові можливості для галузі кібербезпеки та вирішення її основних завдань щодо захисту. Зокрема, автоматизація аналізу мережевого трафіку дозволяє значно скоротити час реакції на загрози, підвищуючи ефективність роботи ІТ-фахівців. Такі системи можуть працювати як у локальних мережах, так і в хмарних середовищах, що робить їх універсальними для різних типів організацій [3].

Завдяки використанню таких систем у корпоративних мережах різного масштабу можна підвищити загальний рівень безпеки, зменшити ризики витоку даних і забезпечити стабільну роботу інформаційних ресурсів. Інтеграція алгоритмів машинного навчання дозволяє адаптувати системи до нових загроз, виявляючи аномалії навіть у разі невідомих атак. Реалізація подібних рішень сприяє підвищенню ефективності моніторингу мереж і забезпечує довготривалий захист від новітніх кіберзагроз.

Література:

1. Xie Y., Yu S. A deep learning approach to network intrusion detection // IEEE Transactions on Emerging Topics in Computing. 2015.
2. Kim D. Y., Kang H. Anomaly detection in network traffic using unsupervised deep learning methods // IEEE Access. 2020.
3. Sommer R., Paxson V. Outside the closed world: On using machine learning for network intrusion detection // IEEE Symposium on Security and Privacy. 2010.

ДОДАТОК Б ЛІСТИНГ ФАЙЛУ mlIDS_teach.py

```
import pandas as pd
import numpy as np
from sklearn.ensemble import RandomForestClassifier
from sklearn.model_selection import train_test_split
from sklearn.metrics import classification_report, accuracy_score
from sklearn.preprocessing import StandardScaler, LabelEncoder
import joblib

# Завантаження даних з кількох файлів

def load_multiple_csv(file_paths):
    """
    Завантажує декілька CSV-файлів і об'єднує їх у загальний
    DataFrame.
    """
    dataframes = []
    for file_path in file_paths:
        print(f"Завантаження даних з {file_path}...")
        data = pd.read_csv(file_path)
        dataframes.append(data)
    combined_data = pd.concat(dataframes, ignore_index=True)
    print(f"Об'єднаний датасет має {combined_data.shape[0]} рядків
i {combined_data.shape[1]} колонок.")
    return combined_data

# Попередня обробка даних

def preprocess_data(data):
    """
    Попередня обробка даних: очищення, нормалізація, перетворення
    міток.
    """
    irrelevant_columns = ['Flow ID', 'Source IP', 'Destination IP',
'Timestamp']
    data = data.drop(columns=irrelevant_columns, errors='ignore')
    data = data.fillna(0)

    if 'Label' in data.columns:
        label_encoder = LabelEncoder()
        data['Label'] = label_encoder.fit_transform(data['Label'])
        print("Мітки перетворено у числовий формат.")

    scaler = StandardScaler()
    numeric_columns = data.select_dtypes(include=['float64',
'int64']).columns
    data[numeric_columns] = scaler.fit_transform(data[numeric_columns])
    print("Дані нормалізовано.")
```

```

return data

# Розділення даних на тренувальну та тестову вибірки
def split_data(data):
    """
    Розділення даних на ознаки та мітки, тренувальну і тестову
    вибірки.
    """
    X = data.drop(columns=['Label'])
    y = data['Label']
    X_train, X_test, y_train, y_test = train_test_split(X, y,
test_size=0.3, random_state=42)
    print("Дані розділено на тренувальну та тестову вибірки.")
    return X_train, X_test, y_train, y_test

# Навчання моделі Random Forest
def train_random_forest(X_train, y_train):
    """
    Навчання моделі Random Forest із підбором оптимальних
    параметрів.
    """
    # Набір параметрів для підбору
    param_grid = {
        'n_estimators': [50, 100, 200], # Кількість дерев у лісі
        'max_depth': [None, 10, 20, 30], # Максимальна глибина
дерева
        'min_samples_split': [2, 5, 10], # Мінімальна кількість
зразків для поділу вузла
        'min_samples_leaf': [1, 2, 4], # Мінімальна кількість
зразків у листі
        'bootstrap': [True, False] # Чи використовувати метод
підстановки
    }

    # Ініціалізація класифікатора Random Forest
    rf = RandomForestClassifier(random_state=42)

    # Налаштування GridSearchCV для пошуку оптимальних параметрів
    grid_search = GridSearchCV(estimator=rf, param_grid=param_grid,
cv=5, scoring='accuracy', n_jobs=-1,
verbose=2)

    # Підбір параметрів на тренувальних даних
    grid_search.fit(X_train, y_train)

    print("Оптимальні параметри для Random Forest:",
grid_search.best_params_)
    print("Найкраща точність на крос-валідації:",
grid_search.best_score_)

```

ДОДАТОК Б ЛІСТИНГ ФАЙЛУ mlIDS_model.py

Продовження лістингу додатку Б

```
# Отримання моделі з найкращими параметрами
best_rf = grid_search.best_estimator_

print("Модель Random Forest із оптимальними параметрами
навчено.")
return best_rf

# Збереження моделі
def save_model(model, file_name):
    """
    Збереження моделі у файл.
    """
    joblib.dump(model, file_name)
    print(f"Модель збережено у файл {file_name}.")

# Оцінка моделі
def evaluate_model(model, X_test, y_test):
    """
    Оцінка продуктивності моделі.
    """
    y_pred = model.predict(X_test)
    accuracy = accuracy_score(y_test, y_pred)
    print(f"Точність моделі: {accuracy:.2f}")
    print("Класифікаційний звіт:")
    print(classification_report(y_test, y_pred))

# Основний скрипт
if __name__ == "__main__":

    # Шляхи до CSV-файлів
    file_paths = [
        'path_to_file1.csv',
        'path_to_file2.csv',
        'path_to_file3.csv',
        'path_to_file4.csv',
        'path_to_file5.csv',
        'path_to_file6.csv',
        'path_to_file7.csv'
    ]

    # Завантаження даних
    combined_data = load_multiple_csv(file_paths)

    # Попередня обробка
    processed_data = preprocess_data(combined_data)

    # Розділення даних
    X_train, X_test, y_train, y_test = split_data(processed_data)
```

ДОДАТОК Б ЛІСТИНГ ФАЙЛУ mlIDS_model.py

Продовження лістингу додатку Б

```
# Навчання Random Forest
rf_model = train_random_forest(X_train, y_train)

# Збереження моделі
save_model(rf_model, "random_forest_model.pkl")

# Оцінка моделі
evaluate_model(rf_model, X_test, y_test)
```


ДОДАТОК В ЛІСТИНГ ФАЙЛУ anomaly_monitoring.py

```
import joblib
import pandas as pd
import numpy as np
from scapy.all import sniff, IP
from sklearn.preprocessing import StandardScaler
import time

# Завантаження моделі
MODEL_PATH = "random_forest_model.pkl"
LOG_FILE = "real_time_anomalies.log"

def load_model(model_path):
    """
    Завантаження навченої моделі з файлу.
    """
    model = joblib.load(model_path)
    print(f"Модель завантажено з {model_path}.")
    return model

def preprocess_packet(packet):
    """
    Обробка пакета для підготовки до аналізу.
    """
    features = {
        "packet_length": len(packet),
        "src_port": packet[IP].sport if IP in packet else 0,
        "dst_port": packet[IP].dport if IP in packet else 0,
        "ttl": packet[IP].ttl if IP in packet else 0,
        "protocol": packet[IP].proto if IP in packet else 0
    }
    return pd.DataFrame([features])

def log_anomaly(packet, prediction):
    """
    Запис аномалії в лог-файл.
    """
    with open(LOG_FILE, "a") as log_file:
        log_entry = (
            f"[{time.strftime('%Y-%m-%d %H:%M:%S')}] Аномалія  

            виявлена!\n"
            f"Деталі пакету: {packet.summary()}\n"
            f"Клас передбачення: {prediction}\n\n"
        )
        log_file.write(log_entry)
    print("Аномалія записана у лог-файл.")

def analyze_packet(packet, model, scaler):
    """
    Аналіз мережевого пакета за допомогою моделі.
    """
    if IP not in packet:
```

```
        return # Ігнорувати пакети, які не містять IP

    features_df = preprocess_packet(packet)
    scaled_features = scaler.transform(features_df)

    prediction = model.predict(scaled_features)[0]

    if prediction == 1: # 1 вказує на аномалію
        print("Аномалія виявлена!")
        log_anomaly(packet, prediction)

def main():
    """
    Основна функція для аналізу мережевого трафіку в реальному часі.
    """
    print("Завантаження моделі...")
    model = load_model(MODEL_PATH)

    print("Підготовка нормалізатора...")
    scaler = StandardScaler()
    scaler.fit(pd.DataFrame([{"packet_length": 0,
                              "src_port": 0,
                              "dst_port": 0,
                              "ttl": 0,
                              "protocol": 0
                              }])) # Масштабування буде застосовуватися до нових даних

    print("Початок аналізу мережевого трафіку...")
    sniff(filter="ip", prn=lambda packet: analyze_packet(packet,
model, scaler), store=0)

if name == "main":
    main()
```