

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ  
ВОЛИНСЬКИЙ НАЦІОНАЛЬНИЙ УНІВЕРСИТЕТ  
ІМЕНІ ЛЕСІ УКРАЇНКИ



**Наталія ЧЕРНЯЩУК**

# ІНДИКАТОРИ КОМПРОМЕТАЦІЇ ЯК ІНСТРУМЕНТ ВИЯВЛЕННЯ КІБЕРАТАК

*Монографія*

Луцьк  
Терен  
2025

*Рекомендовано до друку вченою радою  
Волинського національного університету імені Лесі Українки  
(протокол №11 від 26 вересня 2025 року).*

**Рецензенти:**

**Пастернак Я.М.** – доктор фізико-математичних наук, професор, професор кафедри комп'ютерних наук та кібербезпеки Волинського національного університету імені Лесі Українки;

**Яцків В.В.** – доктор технічних наук, професор, завідувач кафедри кібербезпеки Західноукраїнського національного університету.

**Чернящук Н. Л.**

Індикатори компрометації як інструмент виявлення кібератак: монографія / Н. Л. Чернящук. – Луцьк: Терен, 2025. – 112 с.

ISBN 978-617-8761-11-0

Мета дослідження полягає у вивченні можливостей застосування штучного інтелекту для виявлення вразливостей у мережевій інфраструктурі на основі міток компрометації, з урахуванням особливостей прихованих каналів зв'язку, динаміки атакуючої поведінки та недосконалості традиційних статистичних методів.

У роботі проаналізовано системи виявлення та запобігання вторгненням (IDS та IPS), що дозволило оцінити їхню ефективність у виявленні атак. Досліджено природу стеганоканалів і факторів, що впливають на їх виявлення. Проаналізовані та вивчені індикатори атак, сформовані за допомогою штучного інтелекту на основі аналізу мережевого трафіку.

Розглянуто можливості використання інструменту Splunk Machine для побудови моделей виявлення атак та аналізу аномальної поведінки в мережі. Проведено розробку класифікаторів для побудови системи виявлення вторгнень на основі машинного навчання.

Запропоновано архітектуру системи, здійснено вибір відповідного набору даних для навчання моделі, проведено обробку дисбалансу класів, визначено найбільш значущі ознаки, здійснено їх відбір і скорочення ознакового простору, а також налаштовано модель. Модель пройшла тестування, за результатами якого була проведена оцінка її ефективності.

Для науковців, аспірантів, а також усіх, хто цікавиться проблемами кібербезпеки.

ISBN 978-617-8761-11-0

УДК 004.056.55

© Чернящук Н. Л., 2025

© Чернящук Н. Л., 2025

## ЗМІСТ

|                                                                                                    |    |
|----------------------------------------------------------------------------------------------------|----|
| Вступ .....                                                                                        | 6  |
| 1. Аналіз предметної області виявлення атак на основі індикаторів компрометації.....               | 9  |
| 1.1. Системи виявлення та запобігання вторгненням (IDS/IPS).....                                   | 9  |
| 1.2. Бездротові технології у протидії вторгненням.....                                             | 12 |
| 1.3. Стегано канали як інструмент прихованого перенесення даних у контексті детектування атак..... | 18 |
| 1.4. Атакувальні індикатори, згенеровані за допомогою штучного інтелекту....                       | 20 |
| 2. Детектування атак за допомогою Splunk.....                                                      | 25 |
| 2.1. Застосування Splunk Learning Toolkit у сфері кібербезпеки.....                                | 25 |
| 2.2. Аналіз сучасних рішень з виявлення кібератак.....                                             | 28 |
| 2.3. Побудова класифікаторів мережеских атак .....                                                 | 32 |
| 3. Розробка системи виявлення вторгнень на базі методів машинного навчання.                        | 45 |
| 3.1. Проектування ML-системи виявлення вторгнень.....                                              | 45 |
| 3.2. Підбір та конфігурація моделі.....                                                            | 53 |
| 3.3. Тестування та оцінка ефективності розробленої системи.....                                    | 59 |
| Висновки .....                                                                                     | 64 |
| Список використаних джерел .....                                                                   | 66 |
| Додатки.....                                                                                       | 71 |

## ВСТУП

**Актуальність теми.** У сучасних умовах стрімкого зростання кількості та складності кіберзагроз особливої важливості набуває застосування інтелектуальних технологій для виявлення та прогнозування атак на мережеву інфраструктуру. Сучасні атаки є високодинамічними, гнучко змінюють свою поведінку в реальному часі, маскуються під легітимну активність і часто реалізуються через приховані канали зв'язку, зокрема стеганоканали. Це суттєво ускладнює процес їх виявлення за допомогою традиційних інструментів, які базуються на статичних сигнатурах або простому порівнянні статистичних характеристик трафіку з еталонними зразками.

Більш того, ідеальні моделі виявлення, що ґрунтуються на відхиленні статистики переданого повідомлення від середніх характеристик порожніх контейнерів, не є ефективними в умовах реальних інформаційно-приховуючих систем. Зловмисники можуть використовувати нестационарні джерела, спеціально створені контейнери або зашумлювати канал зв'язку для імітації нормальної активності, що унеможливорює надійне виявлення прихованої передачі інформації звичайними засобами.

У свою чергу, мережеві системи генерують великі обсяги даних, що створює сприятливе середовище для застосування штучного інтелекту. Технології машинного навчання здатні автоматично обробляти трафік, виявляти аномалії, характерні шаблони поведінки атак та ідентифікувати мітки компрометації – артефакти, які залишаються в системі після зловмисної активності.

Таким чином, дослідження методів детектування атак на основі ІоС із використанням інструментів штучного інтелекту є надзвичайно актуальним, оскільки дозволяє ефективно виявляти загрози, що залишаються непоміченими класичними механізмами захисту.

**Мета дослідження** полягає у вивченні можливостей застосування штучного інтелекту для виявлення вразливостей у мережевій інфраструктурі на

основі міток компрометації, з урахуванням особливостей прихованих каналів зв'язку, динаміки атакуючої поведінки та недосконалості традиційних статистичних методів.

**Завдання дослідження:**

- проаналізувати принципи роботи систем виявлення та запобігання вторгненням (IDS/IPS), їх обмеження в контексті прихованих атак;
- дослідити природу стеганоканалів і фактори, що впливають на їх виявлення;
- вивчити індикатори компрометації, які можуть бути сформовані за допомогою ШІ;
- оцінити ефективність використання платформи Splunk Machine для детектування аномалій;
- побудувати та дослідити класифікатори атак на основі міток компрометації;
- обрати відповідний датасет для навчання моделі;
- провести попередню обробку даних: балансування класів, оцінку важливості та відбір ознак;
- здійснити проєктування та оптимізацію моделі машинного навчання;
- протестувати створену модель та оцінити її ефективність.

**Об'єкт дослідження** – процеси навчання моделей штучного інтелекту та виявлення ознак атак у мережевій інфраструктурі.

**Предмет дослідження** – алгоритми машинного навчання для аналізу міток компрометації та аномалій з метою виявлення атак у мережевому трафіку.

**Методи дослідження** базуються на методах машинного навчання (класифікація, відбір ознак, аналіз аномалій), обробка IoC, використання сучасних SIEM-систем (Splunk), а також аналіз поведінкових патернів та моделювання загроз.

**Наукова новизна.** У роботі запропоновано підхід до виявлення атак, що базується на поєднанні міток компрометації та алгоритмів штучного інтелекту,

здатного враховувати змінність, прихованість та складність сучасних атак, зокрема реалізованих через стеганоканали.

**Практична значущість.** Результати дослідження можуть бути використані для розробки програмного рішення або модуля до існуючих систем моніторингу безпеки, здатного виявляти приховані атаки в режимі реального часу, адаптуватися до нових загроз і знижувати ризик успішного проникнення в корпоративну мережу.

# 1 АНАЛІЗ ПРЕДМЕТНОЇ ОБЛАСТІ ВИЯВЛЕННЯ АТАК НА ОСНОВІ ІНДИКАТОРІВ КОМПРОМЕТАЦІЇ

## 1.1 Системи виявлення та запобігання вторгненням (IDS та IPS)

Назва IDS (система виявлення вторгнень) говорить сама за себе – це програмний або апаратний засіб, призначений для виявлення шкідливої активності. Така активність може спостерігатися як у мережевому трафіку, так і в операційній або файловій системі [1].

Основною функцією IDS є виявлення порушень політик безпеки. Однією з характерних рис цих систем є формування звітів про зафіксовані інциденти. IDS, на відміну від брандмауера, не завжди блокує підозрілий трафік. Проте деякі IDS здатні здійснювати активні дії, такі як блокування трафіку або відкидання пакетів – у такому випадку вони наближаються до функціоналу систем запобігання вторгненням (IPS). Незалежно від типу, усі IDS мають одну спільну властивість – здатність повідомляти про виявлені загрози.

Перш ніж детальніше розглянути принципи роботи IDS, наведемо кілька прикладів. Міжмережеві екрани (брандмауери) часто порівнюють із замками на дверях вони пропускають лише авторизований трафік, зупиняючи всі інші спроби доступу. IDS може діяти подібно, не лише фіксуючи загрози, але й зупиняючи їх.

Ще однією важливою функцією багатьох IDS є журналювання мережевих подій. Система може зберігати інформацію про мережеву активність для подальшого аналізу, хоча самотійно вона не припиняє атаку. Так само, як відеоспостереження дозволяє фіксувати інциденти для подальшого перегляду, системні адміністратори аналізують лог-файли IDS для виявлення шкідливих шаблонів поведінки.

Часто IDS працює в режимі пасивного аналізатора – «прослуховувача» (сніффера), який спостерігає за мережевим трафіком і сигналізує про аномальні дії. Це можна порівняти з датчиком сигналізації він не запобігає події, але подає сигнал тривоги при підозрілих діях. Важливо зазначити, що IDS не завжди може

запобігти інциденту. Наприклад, вона може не зупинити зміну файлу адміністратором, але здатна сповістити про цю дію як потенційно небезпечну, що потребує перевірки [2].

Порівняймо IDS із брандмауером. Обидва рішення є складовими систем інформаційної безпеки. Якщо брандмауер переважно виконує функцію контролю доступу на межі між мережами, то IDS має ширший спектр спостереження, включаючи файлові системи та конфігурації. IDS здатна фіксувати спроби вторгнення, навіть якщо вони вже пройшли через брандмауер.

Деякі IDS подібні до брандмауерів у тому, що аналізують мережевий трафік, однак вони також можуть виконувати складніші завдання. IDS можна розгорнути в різних точках інфраструктури, зокрема всередині внутрішньої мережі, де вона функціонує як хост або мережевий сенсор. Важливо усвідомлювати масштаби активності за брандмауером – у сучасних мережах існує багато внутрішніх сегментів, серверів та баз даних, що є цінними цілями для зломисників [2, 3].

У сучасному середовищі, де наявні десятки або сотні клієнтських пристроїв, виникає постійна потреба в захисті від шкідливого програмного забезпечення – наприклад, за допомогою антивірусів. Проте навіть найкращі захисні засоби не гарантують, що зломисник не зможе скомпрометувати хоча б один пристрій, використавши його як точку входу для подальшого розширення атаки в мережі. Наприклад, модифікована прошивка принтера може бути джерелом загрози, навіть якщо сам пристрій розташований за брандмауером.

Смартфони – ще один популярний вектор атак. Це повноцінні комп'ютери, які носять у кишені, часто підключені до тієї ж мережі, що й критично важливі ресурси. Із розвитком концепції Інтернету речей (IoT) з'явилася велика кількість інших типів пристроїв, підключених до мережі, які мають вразливості та можуть бути використані зломисниками. Захист від шкідливої активності, що походить від IoT-пристроїв, стає особливо складним завданням, зокрема через їхню бездротову природу. Якщо атакуючий отримує доступ до бездротової мережі, він фактично обходить периметральний захист і опиняється за брандмауером.

Навіть повне виключення мережевих інтерфейсів не усуває ризиків. Фізичні носії, такі як USB-накопичувачі, залишаються популярним способом поширення шкідливого ПЗ. Антивірус може спрацювати вже після того, як шкідливе ПЗ почне взаємодіяти із серверами чи базами даних.

У таких випадках на перший план виходить система виявлення вторгнень (IDS), яка виконує функції, недоступні звичайному брандмауеру. Один із типових підходів до побудови IDS – це реалізація у вигляді аналізатора трафіку, або сніффера. Його роль можна порівняти з відеокамерою: він фіксує те, що відбувається, але не зупиняє події безпосередньо. Основне призначення такої системи – виявлення підозрілої поведінки й інформування адміністратора про потенційну загрозу.

Однією з ключових функцій IDS є запис подій і трафіку, що дозволяє згодом повернутись до аналізу даних та виявити, що саме сталося. Деякі IDS-системи фіксують увесь трафік у мережі, навіть якщо на момент збору не виявляється жодної аномалії. Після відкриття нових типів атак можна повернутись до вже зібраних даних і перевірити, чи не мали вони місце в минулому.

Таким чином, навіть у своєму пасивному режимі IDS є важливим інструментом безпеки. Вона може просто спостерігати, фіксувати події, інформувати – або, в активнішому варіанті, повністю блокувати шкідливий трафік.

Головна відмінність між системою виявлення вторгнень і превентивними засобами полягає в тому, що останні створюють бар'єр, який зупиняє загрозу ще до її реалізації. IDS також може включати такі функції, як логування й аналіз, але не обов'язково реагує на загрозу негайно.

Одна з найпоширеніших реалізацій IDS – це сигнатурний підхід, де система використовує базу відомих шаблонів шкідливої активності (так званих сигнатур). Принцип роботи такий, IDS аналізує трафік, що проходить через неї, і порівнює його з базою відомих шкідливих шаблонів. Якщо збіг виявлено,

система або блокує трафік (виконуючи функцію IPS), або реєструє інцидент для подальшого аналізу.

Однак сигнатурний метод має обмеження завжди існує ризик хибних спрацювань, коли нормальна активність помилково розцінюється як загроза, або навпаки – нова атака проходить непоміченою, бо її шаблон ще не внесено до бази [4-7].

Передбачливим кроком є завчасне отримання сигнатури, її інтеграція в систему виявлення вторгнень (IDS) і усунення відповідної вразливості. Однак використання сигнатур має свої обмеження. Якщо розглянути життєвий цикл сигнатури, спочатку з'являється новий тип атаки – наприклад, спрямованої на розширення TLS Heartbeat у протоколі SSL. Лише після того, як цю атаку було виявлено і визнано серйозною загрозою, фахівці почали розробляти відповідні сигнатури. Як результат, захисні системи отримали оновлення, і IDS змогли виявляти таку активність та захищати мережу.

## 1.2 Бездротові технології у протидії вторгненням

Через те, що радіохвилі поширюються у просторі без обмежень, не зважаючи на фізичні бар'єри, такі як стіни будівель, бездротові корпоративні мережі постійно наражаються на ризик атак з боку злоумисників. Тому питання безпеки Wi-Fi середовища потребує особливої уваги з боку організацій [8].

З метою підвищення рівня захищеності бездротових локальних мереж, багато компаній вже впровадили або планують впровадження бездротових систем запобігання вторгненням (WIPS – Wireless Intrusion Prevention System). Такі системи виконують моніторинг бездротової активності та виявляють або блокують як внутрішні, так і зовнішні загрози. Завдяки аналізу інформації на фізичному та каналному рівнях моделі OSI, WIPS здатні ефективно виявляти несанкціоновані точки доступу, атаки на бездротові мережі та спроби порушити їхню роботу за допомогою атак типу «відмова в обслуговуванні» (DoS) [9].

З розширенням використання бездротових мереж у корпоративному середовищі, хакерські атаки стають дедалі складнішими і більш цілеспрямованими. У відповідь на ці виклики організації дедалі частіше впроваджують WIPS-рішення для дотримання політики заборони бездротового доступу, виявлення проникнень і блокування атак ще до того, як вони завдадуть шкоди. Такі системи забезпечують розширені можливості спостереження та звітності, що дозволяє оперативно реагувати на загрози та підтримувати безпеку WLAN-інфраструктури.

На рисунку 1.1 наведено схему, яка ілюструє основні загрози, що можуть становити небезпеку для корпоративної мережі.



Рисунок 1.1 – Бездротова мережа та її загрози

Подібно до дротових мереж, у бездротових системах присутні елементи, які безпосередньо або опосередковано виконують функції забезпечення безпеки. Контролер бездротової мережі здійснює базове керування точками доступу й може додатково виконувати функції автентифікації користувачів [9].

Сканери мережевої безпеки, які зовні нагадують точки доступу, призначені виключно для моніторингу трафіку та передачі зібраної інформації до контролера або системи запобігання вторгненням (IPS). Зазвичай при розгортанні WIPS-систем співвідношення між кількістю сканерів і точок доступу становить приблизно 1 до 4 або 1 до 5.

Точки доступу забезпечують підключення пристроїв до мережі, хоча в деяких конфігураціях WIPS вони можуть одночасно виступати в ролі сканерів [10].

Wireless IPS аналізує дані, отримані від точок доступу та сканерів, і, за наявності загроз, надсилає інструкції контролеру для вжиття відповідних заходів. Такий IPS може функціонувати як окремий модуль або бути інтегрованим до бездротового контролера (див. рисунок 1.2).



Рисунок 1.2 – WIPS-мережа, її структура

Організації мають кілька варіантів впровадження систем WIPS, а саме оверлейна (накладена) модель; інтегрована (вбудована) модель; гібридна модель.

Оверлейна архітектура передбачає створення додаткової WIPS-мережі поверх наявної WLAN за допомогою спеціалізованих сенсорів і централізованої системи управління. У цьому підході організація доповнює свою чинну інфраструктуру бездротової мережі шляхом встановлення бездротових сенсорів, відомих як Air Monitors (AM). Такі AM-пристрої інтегруються в мережу аналогічно звичайним точкам доступу, кріпляться на стінах або стелі та живляться через PoE [11].

На відміну від типових Access Points (AP), Air Monitors, як правило, є пасивними компонентами, призначеними для безперервного спостереження за мережевим середовищем з метою виявлення атак чи іншої підозрілої активності (див. рисунок 1.3).



Рисунок 1.3 – Оверлейна WIPS-мережа

Інтегрована модель передбачає застосування єдиної консолі управління для одночасного контролю WLAN і WIPS, що робить можливим розгортання WIPS навіть у мережах із заборonoю використання Wi-Fi. У цьому випадку

організація модернізує існуючу WLAN-інфраструктуру за рахунок використання комбінованих пристроїв AP/AM, які забезпечують як підключення користувачів до мережі, так і моніторинг бездротового трафіку з метою виявлення атак і небажаної активності [13]. Такий підхід зазвичай є більш економічно доцільним, ніж оверлейна модель, оскільки дозволяє поєднати функції доступу й захисту в одному апаратному рішенні, без потреби в додаткових сенсорах (див. рисунок 1.4).



Рисунок 1.4 – Інтегрована WIPS-мережа

Гібридна модель моніторингу поєднує переваги інтегрованої та оверлейної архітектур, використовуючи їх сильні сторони для підвищення ефективності виявлення та запобігання бездротовим загрозам. У межах цього підходу організації можуть комбінувати точки доступу (AP) зі спеціалізованими сенсорами (AM) для посилення безпеки або ж розгорнути окрему інфраструктуру моніторингу, сформовану виключно з AM-пристроїв [14]. Обробка зібраних даних виконується централізованим контролером, що характерно для оверлейної

моделі, на відміну від інтегрованої, де аналіз базується на даних, отриманих від стандартних точок доступу (див. рисунок 1.5).



Рисунок 1.5 – Гібридна мережа WIPS

Порівняльний аналіз варіантів впровадження демонструє, що інтегрована модель вирізняється більшою гнучкістю, кращим радіочастотним охопленням і нижчими сукупними витратами порівняно з оверлейною. Вона також відзначається ефективними механізмами аналізу та потужними засобами протидії вторгненням [15].

У свою чергу, гібридна модель має низку переваг над інтегрованою та оверлейною, зокрема вищу гнучкість при впровадженні, точніший аналіз подій, ширший спектр засобів для виявлення атак і дієві механізми їхнього блокування. Існує безліч підходів до реалізації WIDS/WIPS-рішень, кожен з яких має свої плюси й мінуси. Загалом, гібридна модель нині демонструє помітну перевагу над іншими архітектурами. Водночас, для повноцінного використання потенціалу інтегрованої моделі, важливо ретельно підходити до вибору постачальника, переконуючись у відповідності його рішень функціональним вимогам і здатності ефективно протидіяти як внутрішнім, так і зовнішнім загрозам [16].

### 1.3. Стегано канали як інструмент прихованого перенесення даних у контексті детектування атак

У сучасному кіберпросторі стегано канали дедалі частіше використовуються зловмисниками як засіб прихованої передачі даних, команд або індикаторів компрометації (IoC), що ускладнює виявлення атак традиційними методами. Такі канали дають змогу передавати шкідливу інформацію поза увагою систем моніторингу, використовуючи звичайний, на перший погляд, мережевий або файловий трафік. Це суттєво ускладнює завдання виявлення та аналізу вторгнень, особливо у випадках використання складних багатоетапних атак (APT – Advanced Persistent Threats).

Стегано канали використовуються як частина механізмів прихованої комунікації між елементами зловмисної інфраструктури (наприклад, між зараженим хостом і сервером управління – C2 server), або для непомітного ексфільтрування конфіденційної інформації. У таких випадках мітки компрометації, згенеровані на основі поведінкових ознак (наприклад, нетипова активність у певних мережевих портах, зміна шаблонів запитів тощо), можуть бути недостатні для своєчасного виявлення загрози без глибокого аналізу прихованих каналів.

У роботі були проаналізовані та змодельовані типові реалізації стегано каналів.

DNS-стеганографія – використання запитів до неіснуючих піддоменів, де корисне навантаження кодується у вигляді символів у доменних іменах. Такі запити часто виглядають легітимно, проте їхній надлишок або нестандартна структура може виступати міткою компрометації.

ICMP-тунелювання – приховане передавання даних у полях службових повідомлень протоколу ICMP (наприклад, ping). Звичайні IDS/IPS-системи часто ігнорують цей трафік, вважаючи його діагностичним.

HTTP/HTTPS-обфускація – мітки компрометації можуть бути вбудовані у POST/GET-запити або в заголовки запитів, при цьому дані передаються з

використанням шифрування TLS, що ускладнює глибоку перевірку без розшифрування.

Мережеві таймінгові канали – дані кодуються через часові інтервали між послідовними пакетами. Це один із найбільш складних для детектування типів стеганоканалів.

Хоча самі стеганоканали спрямовані на уникнення виявлення, їх використання залишає певні артефакти в системі. Такі артефакти можуть стати індикаторами компрометації за умови правильного моделювання поведінки. Приклади таких індикаторів:

- аномалії у шаблонах DNS-запитів (частота, довжина, рідкісні домени);
- нестандартна періодичність пакетів або часові затримки (у таймінгових каналах);
- повторювані підозрілі HTTP-запити;
- висока ентропія трафіку або нетиповий розподіл розмірів пакетів.

Інтеграція таких поведінкових ІоС у систему виявлення атак, особливо з використанням методів машинного навчання, значно підвищує ймовірність виявлення прихованих загроз.

Існують різні підходи до виявлення стеганоканалів як елементів інфраструктури атак:

- підписні методи (signature-based) здебільшого малоефективні, оскільки стеганоканали динамічно змінюються;
- аномалійно-орієнтовані методи (anomaly-based) мають більший потенціал, особливо у поєднанні з індикаторами компрометації та профілями поведінки;
- машинне навчання дозволяє побудувати моделі виявлення на основі історичних даних та шаблонів поведінки, що відрізняються від звичайного фону трафіку.

Однак проблемами залишаються нестача якісно анотованих наборів даних, обробка великої кількості помилково позитивних спрацювань, а також

необхідність глибокої інспекції зашифрованого трафіку (що часто неможливо через обмеження безпеки та приватності) [17].

Використання стеганоканалів є одним із ключових викликів у сфері кібербезпеки та детектування атак. Виявлення таких каналів потребує комплексного підходу – поєднання глибокого аналізу трафіку, поведенкових характеристик та індикаторів компрометації. Запропонована методологія доповнює традиційні IDS/IPS-рішення та підвищує точність виявлення складних цілеспрямованих атак.

#### 1.4 Атакувальні індикатори, згенеровані за допомогою штучного інтелекту

Індикатори атак, згенеровані за допомогою штучного інтелекту, підсилюють наявну систему безпеки (див. рисунок 1.6) завдяки використанню хмарних технологій машинного навчання та можливості виявляти загрози в режимі реального часу.

Зібрані дані слугують основою для аналізу активності під час виконання програм і дозволяють оперативно передавати відповідні індикатори на сенсорні пристрої.

Клієнт CrowdStrike зіставляє ці AI-згенеровані індикатори, що відображають поведінкові події, із локальними журналами та файловою інформацією з метою оцінки ризиків, пов'язаних з операціями введення-виведення.

Робота сенсорів та штучного інтелекту відбувається асинхронно, із урахуванням результатів машинного навчання на основі сенсорних даних і вже наявних індикаторів загроз.

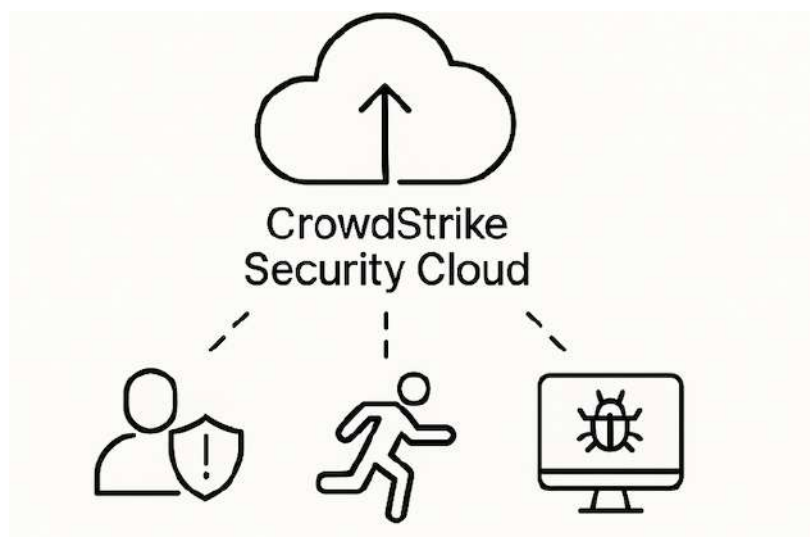


Рисунок 1.6 – Хмарні моделі машинного навчання, використовуючи обширні телеметричні дані CrowdStrike Security Cloud, генерують індикатори атак (IOA) за допомогою штучного інтелекту

Індикатори атак (IOA), згенеровані на основі штучного інтелекту, мають низку ключових переваг. Передусім, вони дозволяють вдвічі швидше виявляти нові типи загроз, що дає змогу випереджати дії зловмисників, передбачати зміни в їхній тактиці та забезпечувати активний локальний захист, який ефективно доповнює існуючі заходи безпеки.

Крім того, завдяки точному розпізнаванню загроз хмарні AI-моделі в режимі реального часу обмінюються індикаторами з сенсорами CrowdStrike Falcon, що дозволяє припиняти атаки незалежно від конкретного шкідливого ПЗ або інструментів, які при цьому використовуються.

Іншою важливою перевагою є зниження кількості помилкових спрацьовувань. Завдяки знанням команди фахівців CrowdStrike із загроз і масштабованим можливостям хмарної інфраструктури, індикатори атак дають змогу ефективніше аналізувати поведінкові дані, що суттєво підвищує продуктивність роботи аналітиків навіть у великих корпоративних середовищах.

ІОА являють собою логічну послідовність поведінкових ознак, які вказують на активну або поточну спробу компрометації системи. Це дає змогу

фахівцям з безпеки відстежити способи проникнення, цілі та мотиви зловмисників, а також сформувані повне уявлення про хід атаки — незалежно від використаного шкідливого програмного забезпечення чи технічних засобів.

Також ІОА використовує поглиблений поведінковий аналіз для моделювання й передбачення можливих дій атакуючих, що істотно підвищує ефективність захисту від майбутніх загроз.

У порівнянні з традиційними індикаторами компрометації (ІОС), ІОА забезпечує виявлення атак ще на ранніх етапах або в процесі їхнього розвитку, дозволяючи реалізувати проактивний підхід до кібербезпеки. Натомість ІОС орієнтовані переважно на реагування після завершення атаки, що робить їх менш ефективними в запобіганні загрозам.



Рисунок 1.7 – Дані про компрометацію та поведінкові сигнали атак, зібрані за допомогою рішень CrowdStrike

Зосереджуючись на мотиваціях зловмисників, а не на конкретних шкідливих програмах чи інструментах, ІОА дозволяють клієнтам оперативно адаптуватися до нових типів атак та змін у тактиках зловмисників. Це включає, зокрема, атаки без застосування шкідливого ПЗ або безфайлові атаки, які протягом останнього року становили 62% усіх зафіксованих інцидентів.

Завдяки своїй універсальності IOA можна аналізувати паралельно, що забезпечує високу обчислювальну ефективність і масштабованість. Крім того, вони потребують менш частого оновлення порівняно з сигнатурними методами, які використовуються для традиційних IOC [18].

Раніше створення IOA значною мірою спиралося на практичний досвід провідних експертів із пошуку загроз, що дозволяло формувати складні й точні індикатори. Однак для ефективного виявлення майбутніх загроз необхідно прискорити процеси ідентифікації та класифікації активних атак, зберігаючи при цьому високу точність, властиву експертно створеним IOA. Для цього CrowdStrike об'єднала людський досвід із машинним навчанням, що дозволяє масштабувати створення IOA й підвищувати їхню якість без втрати точності.

Завдяки використанню обчислювальних потужностей CrowdStrike Security Cloud для навчання моделей на платформі Falcon, машинне навчання здатне швидко й точно аналізувати великі обсяги даних про загрози. Впровадження хмарного машинного навчання у процес створення IOA забезпечує клієнтам проактивні, високоякісні індикатори атак з високою швидкістю та масштабованістю.

Від моменту запуску хмарні моделі машинного навчання ідентифікували понад 20 нових шаблонів індикаторів, які були підтверджені експертами й інтегровані в платформу Falcon для автоматичного виявлення та блокування атак. Нижче наведено два приклади тактик зловмисників, що стали основою нових IOA, пов'язаних із корисним навантаженням після експлуатації та атаками з використанням PowerShell.

Корисне навантаження після експлуатації — це код, який зловмисник завантажує на хост після отримання контролю над ним. AI-орієнтовані IOA виявляють такі навантаження, комбінуючи дані зі статичної AI-моделі сенсора Windows, інформацію про виконуваний файл і унікальні дані, отримані через CrowdStrike Security Cloud. Цей підхід враховує походження процесу та спосіб його запуску, що забезпечує високу точність виявлення і дозволяє формувати

детальні індикатори, значно перевершуючи ефективність традиційних статичних або поведінкових методів.

Атаки з використанням PowerShell часто застосовуються зловмисниками для доставки шкідливого коду або виконання зловмисних дій, коли класичні ІОС неефективні. Такі атаки важко виявити, а традиційні сигнатурні методи легко обходяться. Використання моделей глибокого навчання для автоматичного виявлення ключових фрагментів PowerShell-скриптів дозволяє ідентифікувати безфайлові загрози, керовані PowerShell, і захищати від них системи [19].

Машинне навчання залишається важливим інструментом для виявлення нових закономірностей у даних і проведення глибокого аналізу поведінки зловмисників, що допомагає зрозуміти їхні наміри та цілі. CrowdStrike, як лідер у сфері хмарного захисту кінцевих точок, хмарних робочих навантажень і аналітики даних, планує й надалі поєднувати штучний інтелект та хмарні технології для підвищення ефективності захисту, протидії методам кіберзловмисників і забезпечення клієнтів можливістю швидко блокувати атаки.

## 2 ДЕТЕКТУВАННЯ АТАК ЗА ДОПОМОГОЮ SPLUNK

### 2.1 Застосування Splunk Learning Toolkit у сфері кібербезпеки

В умовах стрімкої цифрової трансформації в різних секторах економіки, активної цифровізації державного управління, охорони здоров'я, освіти та наукової сфери, а також з огляду на зростання популярності інтернет-сервісів і мобільних пристроїв, питання захисту мобільних мереж стає все більш актуальним.

Сучасне суспільство дедалі більше покладається на безперебійну роботу мобільного зв'язку, що забезпечує доступ до критичних сервісів, банківських операцій, обміну конфіденційною інформацією та підтримки життєво важливої інфраструктури.

У зв'язку з цим зловмисники дедалі частіше обирають саме мобільні мережі як об'єкт своїх атак, використовуючи як загальновідомі вразливості, так і нові, ще недостатньо вивчені вектори проникнення.

Із поширенням джерел та ускладненням методів здійснення кібератак зростає складність своєчасного виявлення різноманітних кіберзагроз. Кібератаки стають дедалі більш таргетованими, багаторівневими та обфускованими, що ускладнює їхнє виявлення традиційними методами [20].

Наприклад, зловмисники можуть комбінувати соціальну інженерію, шкідливе програмне забезпечення, експлойти в мережевих протоколах і навіть атаки з використанням штучного інтелекту. Такі методи дозволяють обходити системи захисту, залишаючись непоміченими протягом тривалого часу.

Традиційні підходи до виявлення атак у мережі, які переважно базуються на статичних правилах, зокрема сигнатурному аналізі, використанні чорних списків або регулярних виразів, виявляються недостатньо гнучкими та малоефективними для ранньої діагностики аномалій і оперативного реагування на інциденти безпеки.

Такі підходи часто не здатні адаптуватися до нових або модифікованих форм атак, які ще не внесені до бази сигнатур. Крім того, вони генерують велику

кількість хибнопозитивних або хибнонегативних спрацювань, що перевантажує аналітиків з кібербезпеки та знижує загальну ефективність захисту.

У зв'язку з цим все більше уваги приділяється розвитку інтелектуальних систем виявлення атак, які базуються на машинному навчанні, поведінковому аналізі, індикаторах компрометації (IoC) та індикаторах атак (IoA). Застосування таких підходів дозволяє створювати більш адаптивні та ефективні системи виявлення, здатні оперативно реагувати на нові, ще невідомі загрози, виявляти аномальну активність на ранніх етапах та знижувати ризики компрометації критичних ресурсів. Для подолання цих обмежень доцільним є використання алгоритмів машинного навчання, що відкривають нові можливості для точнішого виявлення шкідливої активності в інформаційних мережах. У межах цього дослідження було застосовано платформу аналізу даних Splunk у поєднанні з Splunk, що дозволило створити, навчити, протестувати та оцінити класифікатори для виявлення мережових атак [21].

Оцінювання ефективності моделі здійснювалося із застосуванням чотирьох популярних алгоритмів машинного навчання дерева рішень, методу опорних векторів, випадкового лісу та подвійного випадкового лісу. Результати експериментів продемонстрували, що всі зазначені алгоритми здатні ефективно виявляти мережові атаки, а найвищу точність у розпізнаванні атак типу «відмова в обслуговуванні» (DoS) показав саме метод подвійного випадкового лісу.

Таким чином, у сучасних умовах розширення цифрового простору та загострення кіберзагроз, оператори мобільного зв'язку стикаються з новими викликами безпеки, що потребують впровадження інтелектуальних аналітичних рішень на основі машинного навчання. Системи стільникового зв'язку постійно удосконалюються, що проявляється у розвитку мережевої архітектури, модернізації інтерфейсів та протоколів, а також у збільшенні пропускної здатності каналів передачі даних. Водночас із технічним прогресом виникають нові вразливості, які можуть бути використані зловмисниками для здійснення атак як на рівні мережі доступу, так і в межах базової мережі.

Такі атаки можуть спричинити перебої в обслуговуванні або активувати небажані функції, що загрожує стабільності роботи систем. У зв'язку з цим особливої актуальності набуває розробка надійних засобів для виявлення мережових атак і оперативного реагування на кіберінциденти. Одним із рішень у цій сфері є системи управління інформацією та подіями безпеки (SIEM), які відіграють ключову роль у сучасній архітектурі кіберзахисту разом із міжмережевими екранами, системами виявлення та запобігання вторгненням (IDS/IPS), а також інструментами для аналізу шкідливого програмного забезпечення – як статичного, так і динамічного. Раніше нами було реалізовано програмний комплекс для оперативного центру інформаційної безпеки підприємства (ОЦІБ), у якому вбудована SIEM-система для моніторингу та дослідження інцидентів кібербезпеки. Система підтримує інтерактивну візуалізацію за допомогою аналітичних панелей (дашбордів), які автоматично формуються на основі даних, отриманих із різноманітних апаратно-програмних компонентів ОЦІБ. Серед доступних SIEM-рішень – як з відкритим кодом, так і комерційних – було обрано Splunk, побудовану на платформі Splunk [23].

Ця платформа є провідним рішенням для операційної аналітики, яке дозволяє ефективно працювати з машинними даними, що сприяє підвищенню продуктивності, швидшому виявленню несправностей і посиленню інформаційної безпеки організацій. У межах цієї роботи представлено результати експериментального дослідження, спрямованого на реалізацію класифікатора мережових атак на основі алгоритмів машинного навчання. Було використано інструментарій Splunk Machine Learning Toolkit, еталонний набір даних UNSW-NB15, мову програмування Python, а також мову обробки даних SPL (Search Processing Language). Оцінювання точності моделі здійснювалося за допомогою чотирьох алгоритмів дерева рішень, методу опорних векторів, випадкового лісу та подвійного випадкового лісу.

Результати експериментів підтвердили ефективність використання машинного навчання як для розпізнавання нормальної та аномальної активності в мережі, так і для точного виявлення різних типів мережових атак.

## 2.2 Аналіз сучасних рішень з виявлення кібератак

За останні роки значно збільшилася кількість досліджень, спрямованих на удосконалення методів виявлення мережових атак, зокрема тих, що застосовують методи машинного навчання та штучного інтелекту для ідентифікації аномалій у трафіку та виявлення нових, раніше невідомих загроз. Такий підхід дозволяє не лише автоматизувати процеси аналізу даних, а й забезпечити високу точність виявлення складних атак у динамічних мережових середовищах.

Розвиток технологій Big Data, зростання обсягів телеметрії, що збирається з кінцевих пристроїв, серверів, хмарних сервісів і мережових вузлів, створює нові можливості для створення більш інтелектуальних та адаптивних систем захисту. В роботі [8] наведено широкий спектр методів виявлення вторгнень, які можна умовно класифікувати на дві основні групи системи, що базуються на сигнатурах (signature-based IDS), та системи, що орієнтовані на виявлення аномалій (anomaly-based IDS). Сигнатурні системи є традиційним підходом у сфері інформаційної безпеки та працюють шляхом аналізу мережових пакетів або поведінки системи, зіставляючи їх з відомими шаблонами атак, які зберігаються в попередньо сформованих базах даних. Їх основна перевага полягає в високій точності при виявленні відомих загроз із мінімальною кількістю хибнопозитивних спрацювань. Зокрема, вони не здатні ефективно виявляти атаки нульового дня (zero-day attacks), цілеспрямовані атаки (targeted attacks), а також поліморфні та метаморфні варіанти шкідливого програмного забезпечення (ВПЗ), які динамічно змінюють свій код або поведінку з метою обходу сигнатурного аналізу. Унаслідок цього сигнатурні системи потребують постійного оновлення баз даних, що не завжди можливо в реальному часі, особливо в умовах активної еволюції загроз [25].

Натомість системи виявлення аномалій ґрунтуються на побудові моделей нормальної поведінки мережі або користувачів. Відхилення від цих моделей можуть свідчити про потенційні загрози або підозрілу активність. У поєднанні з

алгоритмами машинного навчання, ці системи здатні виявляти невідомі або модифіковані форми атак, адаптуватися до нових шаблонів поведінки та підвищувати ефективність детектування в реальному часі. Однак вони також стикаються з проблемами – зокрема, високим рівнем хибнопозитивних спрацювань, складністю побудови точних моделей поведінки в різномірних середовищах і потребою в якісному наборі даних для навчання. Таким чином, сучасні наукові дослідження зосереджуються на пошуку балансу між точністю, адаптивністю та швидкістю систем виявлення загроз, із використанням гібридних підходів, глибокого навчання, обробки потоків даних у реальному часі, а також впровадження концепцій індикаторів компрометації (IoC) та атак (IoA), що значно розширює можливості виявлення як відомих, так і нових форм кіберзагроз. Натомість аналіз мережевих пакетів із застосуванням підходів на основі аномалій виявляється більш ефективним для розпізнавання складних кібератак, оскільки виявлення аномальної поведінки користувачів не залежить від сигнатурної бази [28]. У цій же роботі методи виявлення кібератак на основі аномалій поділяють на три основні категорії статистичні, знаннєві та засновані на машинному навчанні. Ці групи разом із прикладами їхніх підкласів ілюструються на рисунку 2.1.

Статистичні методи передбачають збір і аналіз кожного запису даних із побудовою статистичної моделі нормальної поведінки користувача. Знаннєві підходи зосереджуються на визначенні нормальних і аномальних специфікацій протоколів та аналізі екземплярів мережевого трафіку. Методи на основі машинного навчання створюють шаблони нормального трафіку та аномальних ситуацій, навчають і тестують класифікатори на цих даних. Автори детально розглянули кілька робіт із застосуванням машинного навчання для виявлення атак нульового дня і вказали на існуючі проблеми, пов'язані з оновленням інформації про нові атаки та високою кількістю помилкових спрацювань або низькою точністю. Це зумовлено тим, що більшість моделей машинного навчання тренуються на застарілих наборах даних DARPA/KDD99, зібраних у 1999 році, які не відображають сучасні типи шкідливих дій. Таким чином,

виникає необхідність у нових і більш повних наборах даних, що охоплюють широкий спектр актуальних загроз.

Нижче наведено короткий огляд досліджень, у яких застосовували різні методи машинного навчання для виявлення аномалій або певних типів мережових атак. Так, у роботі [20] для виявлення атак типу «відмова в обслуговуванні» (Denial of Service, DoS) запропоновано модель, що базується на зборі даних із файлів логів платформи Spark та класифікації атак методом випадкового лісу (Random Forest) з точністю 99,2%. Автори довели, що ця модель ефективно обробляє великі потоки DNS-запитів, що робить її придатною для практичного застосування. Запропоноване рішення складається з двох основних компонентів модуля збору мережового трафіку в реальному часі та модуля виявлення атак. Модуль збору відповідає за збір і вилучення характеристик мережового трафіку з використанням моделі мікропакетної обробки. Точність роботи алгоритму була перевірена та порівняна на наборах даних NSL-KDD і UNSW-NB15. Типологія систем виявлення кібератак, що ґрунтуються на виявленні аномалій [22] наведено на рисунку 2.1.

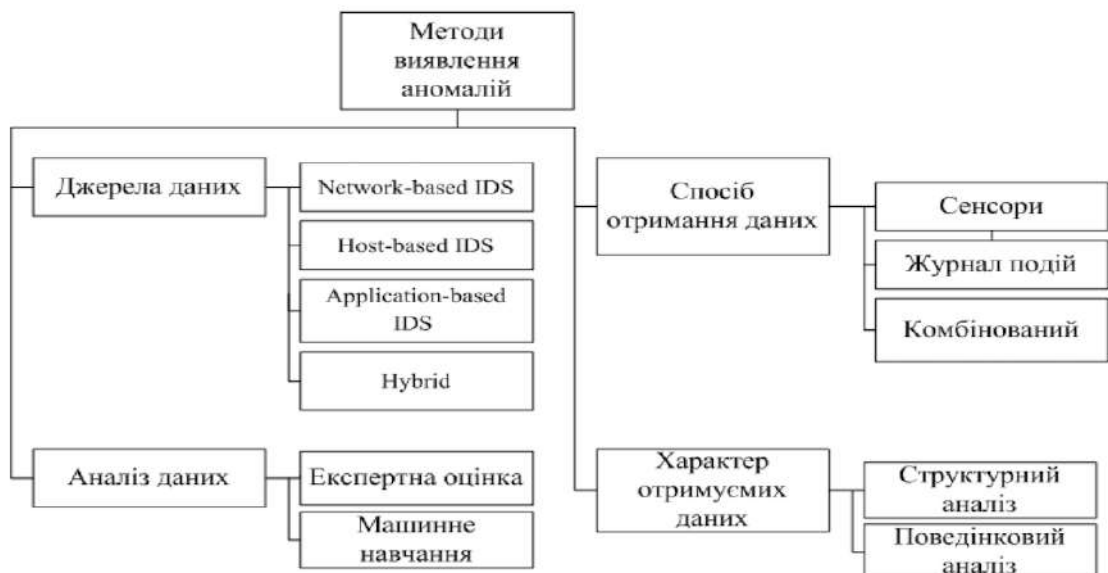


Рисунок 2.1 – Типологія систем виявлення кібератак, що ґрунтуються на виявленні аномалій

У дослідженні запропоновано моделі глибокого навчання для виявлення та зменшення ризику DDoS-атак, спрямованих на централізований контролер у програмно-визначуваних мережах, з використанням рекурентної нейронної мережі з довгою короткостроковою пам'яттю (LSTM) та згорткової нейронної мережі (CNN) [23].

Точність класифікації DDoS-атак за допомогою CNN була відносно низькою – 66%, тоді як модель LSTM продемонструвала кращий результат – 89,63%, перевищуючи показники класичних методів машинного навчання, таких як метод опорних векторів (SVM) і наївний байєсівський класифікатор, які досягали 86,85% та 82,61% відповідно.

Застосовуючи методи випадкового лісу та багат шарового персептрону на платформі обробки великих даних Spark ML, вдалося досягти точності виявлення DoS-атак на рівні 99,5% у реальному часі за кілька мілісекунд [24].

Підсумовуючи огляд існуючих рішень, можна припустити, що інтеграція сучасних методів машинного навчання з технологіями збору машинних даних у SIEM-системах на базі операційної аналітики є перспективним і ефективним підходом для створення систем раннього виявлення аномалій та класифікації мережевих атак у контексті кібербезпеки.

Для перевірки методики проведено експерименти з детектування аномалій на штучно згенерованих DNS-запитах та найбільш поширених категоріях кібератак.

Для реалізації досліджень була використана платформа Splunk Enterprise з додатковим інструментарієм Machine Learning Toolkit [25].

Ця бібліотека містить понад 30 найпоширеніших алгоритмів машинного навчання з відкритим вихідним кодом на мові Python.

Підготовка даних здійснювалася на сервері Splunk із використанням вбудованої мови пошуку Search Processing Language (SPL).

Для виконання SPL-скриптів над наборами даних застосовується інтерпретатор мови Python.

Для проведення експерименту було обрано такі алгоритми машинного навчання, які часто використовуються в аналогічних дослідженнях, що дозволяє порівнювати отримані результати з результатами інших науковців: дерево рішень (Decision Tree, DT); випадковий ліс (Random Forest, RF); метод опорних векторів (Support Vector Machine, SVM).

Ці методи є класичними і їх детальний опис наведено у численних підручниках та наукових публікаціях, зокрема в роботах [25, 26].

Також в експерименті застосовано метод подвійного випадкового лісу (Double Random Forest, DRF), який є модифікацією класичного випадкового лісу.

Автори роботи [26] продемонстрували, що точність класифікації методу випадкового лісу можна покращити, використовуючи ансамбль різноманітних дерев замість дерев з мінімальним розміром вузлів.

Запропоновано новий підхід, у якому навчальна вибірка завантажується безпосередньо на кожному вузлі під час побудови дерева, на відміну від традиційного методу, де дані завантажуються лише в кореневому вузлі.

Експериментальні дослідження з розпізнавання рукописних цифр від 0 до 9 показали, що метод DRF значно перевершує класичні ансамблеві алгоритми класифікації [27]. Саме тому було вирішено застосувати алгоритм DRF з урахуванням специфіки задачі класифікації мережевих атак.

### 2.3 Побудова класифікаторів мережевих атак

Для навчання та тестування використана база даних UNSW-NB15, створена в лабораторії Cyber Range ACCS за допомогою генератора аномального мережевого трафіку, який імітує 9 різних типів мережевих атак (див. таблиця 2.1).

Як зазначено авторами роботи [20], для збору мережевих пакетів загальним обсягом 100 ГБ застосовувався програмний інструмент tcpdump, який також використовувався для розбиття трафіку на фрагменти розміром по 1000 МБ.

Надалі з pcap-файлів за допомогою програмного забезпечення Argus було сформовано 4 CSV-файли, що містять ключові характеристики – надійні ознаки (атрибути) для класифікації кількох типів атак.

Ці атрибути включають параметри мережевого трафіку, такі як IP-адреси, порти, протоколи, тривалість сесій та обсяг переданих даних, що дозволяє ефективно ідентифікувати та розрізняти різні види кібератак у процесі машинного навчання. Крім того, база даних містить виведені ознаки, отримані в результаті обробки первинних даних, зокрема рівень ентропії, коефіцієнт фрагментації, кількість сесій на одиницю часу тощо.

Було виокремлено 49 основних характеристик (атрибутів для класифікації атак), які поділяються на кілька категорій базові (Basic), контентні (Content) та часові (Time), що наведено у таблиці 2.1.

Таблиця 2.1 – Категоризація записів у датасеті UNSW-NB15 за видами атак

| Типи мережевих атак                                                   | Кількість записів | Description                                                                                        |
|-----------------------------------------------------------------------|-------------------|----------------------------------------------------------------------------------------------------|
| Normal (нормальний трафік)                                            | 2 218 761         | природні дані нормального трафіку                                                                  |
| Fuzzer (трафік, створюваний програмою фаззером)                       | 24 246            | спроба призупинити роботу мережі шляхом передачі випадково згенерованих даних                      |
| Analysis (трафік, створюваний програмою аналізатором)                 | 2 677             | трафік містить різні атаки сканування портів, html-файлів та проникнення спам                      |
| Backdoors (бекдори)                                                   | 2 329             | техніка, при якій механізм безпеки системи потай обходить для доступу до комп'ютера або його даних |
| DoS (атаки типу «відмова в обслуговування», в тому числі розподілені) | 16 353            | зловмисна спроба зробити сервер або мережевий ресурс недоступний для користувачів                  |

Продовження таблиці 2.1

|                                                 |         |                                                                                                                                        |
|-------------------------------------------------|---------|----------------------------------------------------------------------------------------------------------------------------------------|
| Exploits (атаки з використанням експлойтів)     | 44 525  | зловмисник знає про проблему безпеки в операційній системі, використовує ці знання, застосовуючи вразливість                           |
| Generic (атака проти шифрів)                    | 215 481 | загальний метод працює проти всіх блокових шифрів (із заданим розміром блоку та ключа), без обліку структури блочного шифру            |
| Reconnaissance (пасивні атаки з метою розвідки) | 13 987  | містить усі спроби проникнення в мережу, які можуть імітувати атаки зі збором інформації                                               |
| Shellcode (двійковий шкідливий код)             | 1 511   | невеликий фрагмент коду, який використовується в якості корисного навантаження при експлуатації вразливості у програмному забезпеченні |
| Worms (шкідливий код типу «хробак»)             | 174     | шкідливий код, який копіює себе, щоб поширитись на інші комп'ютери                                                                     |

Зразки базових характеристик наведено в таблиці 2.2. Дані з файлу **UNSW-NB15\_1.csv** були використані як навчальна вибірка, а дані з **UNSW-NB15\_2.csv** – як тестова вибірка.

Схематичне представлення експерименту наведено на рисунку 2.2. Вибірка необхідної категорії мережових атак виконувалась за допомогою команди **»search»** мови запитів Splunk.

*search attack\_cat=Normal OR attack\_cat=Dos,*

Останній параметр визначає тип атаки (див. таблицю 2.2); у наведеному прикладі використано параметр Dos, що відповідає атаці типу «відмова в обслуговуванні» (Denial of Service, DoS).

На етапі попередньої обробки текстові та символні поля перетворюються у числові за допомогою методу one-hot кодування, оскільки алгоритми машинного навчання працюють лише з числовими значеннями.

Після цього проводиться нормалізація даних, у результаті чого кожен запис мережевого трафіку подається у вигляді числової матриці атрибутів, яка надалі використовується для класифікації одним із чотирьох обраних методів машинного навчання.

Таблиця 2.2 – Типові характеристики потоків у мережі

| № атрибуту | Позначення | Тип даних                  | Опис                                                                                                                          |
|------------|------------|----------------------------|-------------------------------------------------------------------------------------------------------------------------------|
| 1          | srcip      | N (номінальний)            | IP Адреса джерела (Source IP address)                                                                                         |
| 2          | sport      | I (ціле число)             | Номер порту джерела (Source port number)                                                                                      |
| 3          | dstip      | N                          | IP Адреса одержувача (Destination IP address)                                                                                 |
| 4          | dsport     | I                          | Номер порту одержувача (Destination port number)                                                                              |
| 5          | proto      | N                          | Протокол транзакції (Transaction protocol)                                                                                    |
| 29         | stime      | T (мітка часу)             | Час початку запису (record start time)                                                                                        |
| 34         | synack     | F(число з плаваючою комою) | Година між пакетами SYN та SYN_ACK у TCP The time between the SYN and the SYN_ACK packets of the TCP                          |
| 48         | attack_cat | N                          | Назва категорії атаки (Fuzzers, Analysis, Backdoors, DoS, Exploits, Generic, Reconnaissance, Shellcode, Worms) див. Таблицю 1 |
| 49         | Label      | B (двійковий)              | 0 для запису нормального трафіку та 1 для запису трафіку при мережевій атаці                                                  |

Приклад нижче демонструє, як формується запит для пошуку даних з індексованого набору даних та подальшого застосування до нього обраного алгоритму машинного навчання:

```
index=nb15_test|search attack_cat=Normal OR attack_cat=Dos|fields label,
ackdat, ct_dst_ltm, ct_dst_sport_ltm, ct_dst_src_ltm, ct_flw_http_mthd, ct_ft_
ct_src_dport_ltm, ct_src_ltm, ct_srv_dst, ct_srv_src, ct_state_ttl, dbytes,
dinpkt, djit,
dload, dloss, dmean, dpkts, dtcpb, dttl, dur, dwin, is_ftp_login,
is_sm_ips_ports, proto, rate, response_body_len, sbytes, service, sinpkt, sjit, sload,
sloss, smean, spkts, state, synack, tcprrt, trans_depth |apply DRF_DOS
```

Задача виявлення аномалій або розпізнавання конкретного типу мережових атак є прикладом бінарної класифікації, ефективність якої оцінюється за допомогою таких метрик: матриця неточностей (confusion matrix, CM), точність класифікації (accuracy), точність передбачення (precision), повнота (recall) та F-міра (F-score).

У таблиці 2.3 представлено приклад матриці невідповідностей для задачі класифікації мережових атак. У ній стовпці відповідають прогнозованим класам, а рядки – фактичним класам.

Таблиця 2.3 – Матриця невідповідностей для задачі класифікації мережових атак

| Actual Class | Predicted Class     |                     |
|--------------|---------------------|---------------------|
| Class        | Normal              | Attack              |
| Normal       | True negative (TN)  | False Positive (FP) |
| Attack       | False Negative (FN) | True positive (TP)  |

Результативність роботи алгоритму класифікації зазвичай оцінюється за допомогою чотирьох основних метрик, формули яких подано нижче з урахуванням позначень, наведених у матриці помилок (див. таблицю 2.3).

Метрика *Accuracy* (точність, або правильність класифікації) визначає частку коректно класифікованих випадків серед загальної кількості прогнозів і демонструє загальну ефективність класифікатора (рисунк 2.2.).

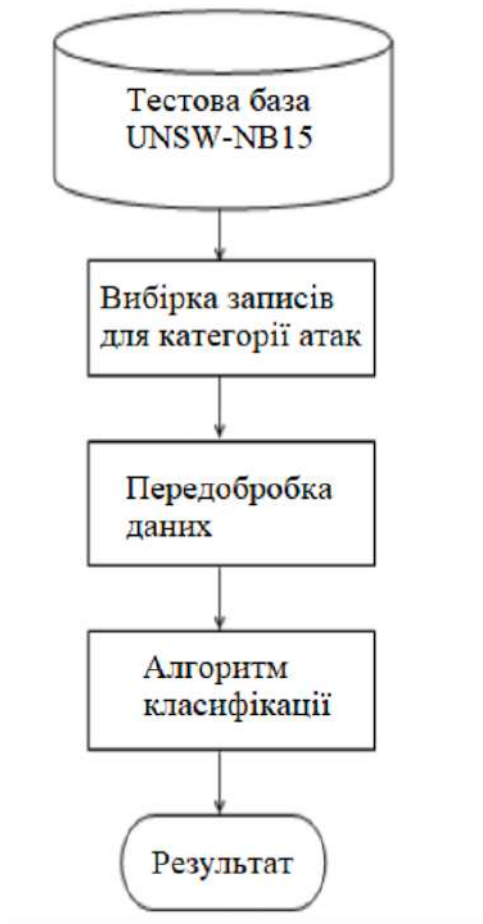


Рисунок 2.2 – Структура експериментального дослідження з класифікації атак за допомогою алгоритмів машинного навчання

$$\text{Accuracy} = \frac{TP + TN}{TP + FP + FN + TN} \quad (2.1)$$

Метрика *Precision* (прецизійність, або точність класифікації) відображає частку випадків, коли класифікатор правильно ідентифікував мережеву атаку серед усіх випадків, які він класифікував як атаки, тобто скільки з передбачених атак дійсно є справжніми атаками:

$$\text{Precision} = \frac{TP}{TP + FP} \quad (2.2)$$

Метрика *Recall* (чутливість, повнота) або *True Positive Rate* (*TPR*) визначається як частка правильно виявлених атак серед усіх реальних випадків атак, тобто це відношення кількості вірно класифікованих атак до загальної кількості фактичних атак.

$$\text{Recall} = \frac{TP}{TP + FN} \quad (2.3)$$

У разі, якщо всі мережеві атаки були правильно ідентифіковані, значення показника *TPR* (чутливість) дорівнює 1. Проте така ситуація є надзвичайно рідкісною для реальних систем виявлення вторгнень. Як правило, значення *TPR* лише наближається до одиниці.

Крім того, для узагальненої оцінки ефективності алгоритму часто використовується агрегована метрика *F1-score* (F-міра), яка обчислюється як гармонійне середнє між точністю (*Precision*) та чутливістю (*Recall*).

$$F1 - score = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (2.4)$$

F-міра досягає максимального значення лише тоді, коли як точність, так і повнота дорівнюють одиниці. Якщо хоча б один з цих показників наближається до нуля, F-міра також зменшується до значення, близького до нуля.

Усі формули для розрахунку зазначених метрик ефективності алгоритмів машинного навчання реалізовані в бібліотеці *scikit-learn*, яка використовується в рамках *Machine Learning Toolkit*. Цей інструмент також підтримує реалізацію трьох основних алгоритмів машинного навчання: дерево рішень (DT), випадковий ліс (RF) та метод опорних векторів (SVM).

Для впровадження нового підходу – *подвійного випадкового лісу (DRF)* – було розроблено окремий плагін мовою Python, який отримав назву *SPL-DRF*. З метою підвищення продуктивності під час реалізації цього модуля була використана *паралельна побудова дерев і паралельне обчислення прогнозів* за допомогою параметра `n_jobs`. Якщо вказано `n_jobs=k`, то обчислення розподіляються на  $k$  завдань, які виконуються одночасно на  $k$ -процесорних ядрах (рисунок 2.3).

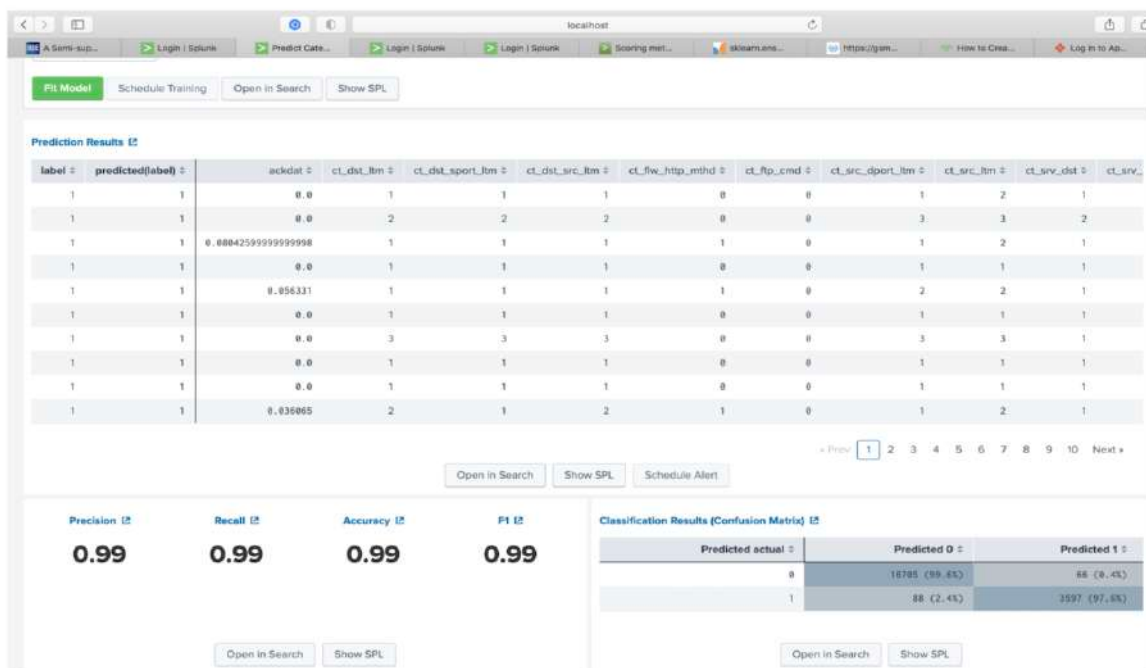


Рисунок 2.3 – Результати класифікації мережевого трафіку для виявлення DoS-атак із застосуванням алгоритму випадкового лісу (RF)

У випадку, коли параметр `n_jobs` встановлено на значення -1, задіюються всі доступні процесорні ядра, що дозволяє суттєво прискорити процес навчання

моделі. Це особливо ефективно при побудові великої кількості дерев або при роботі з об'ємними наборами даних, де створення одного дерева є ресурсоемним.

На рисунках 2.4-2.6 наведено скріншоти з результатами тестування класифікаторів, що розрізняють нормальний та аномальний мережевий трафік із використанням бази даних UNSW-NB15.

Для класифікації було застосовано три алгоритми: дерево рішень (DT), випадковий ліс (RF) та метод опорних векторів (SVM). У всіх прикладах розглядалася атака типу «відмова в обслуговуванні» (DoS).

Як видно зі згаданих рисунків, у лівій нижній частині інтерфейсу відображено значення метрик ефективності класифікації, округлені до двох знаків після коми. Через це складно точно оцінити різницю між метриками різних алгоритмів. У правій нижній частині показані відповідні матриці помилок, які дозволяють вручну перерахувати метрики за формулами (2.1)–(2.4) для більш точного аналізу.

У таблиці 2.4 зібрано результати класифікації, отриманих за допомогою різних методів машинного навчання.

Таблиця 2.4 – Зіставлення результатів класифікації, отриманих за допомогою різних методів машинного навчання

| Алгоритм машинного навчання    | Accuracy | Precision | Recall | F1-score |
|--------------------------------|----------|-----------|--------|----------|
| Випадковий ліс (RF)            | 0,9925   | 0,9948    | 0,9961 | 0,9954   |
| Дерево рішень (DT)             | 0,9885   | 0,9932    | 0,9927 | 0,9930   |
| Машина опорних векторів (SVM)  | 0,6985   | 0,8370    | 0,6945 | 0,7329   |
| Подвійний випадковий ліс (DRF) | 0,9984   | 0,9987    | 0,9993 | 0,9990   |

На рисунку 2.4 показано оцінку ефективності класифікації мережевого трафіку для виявлення атак із використанням алгоритму дерева рішень (DT).

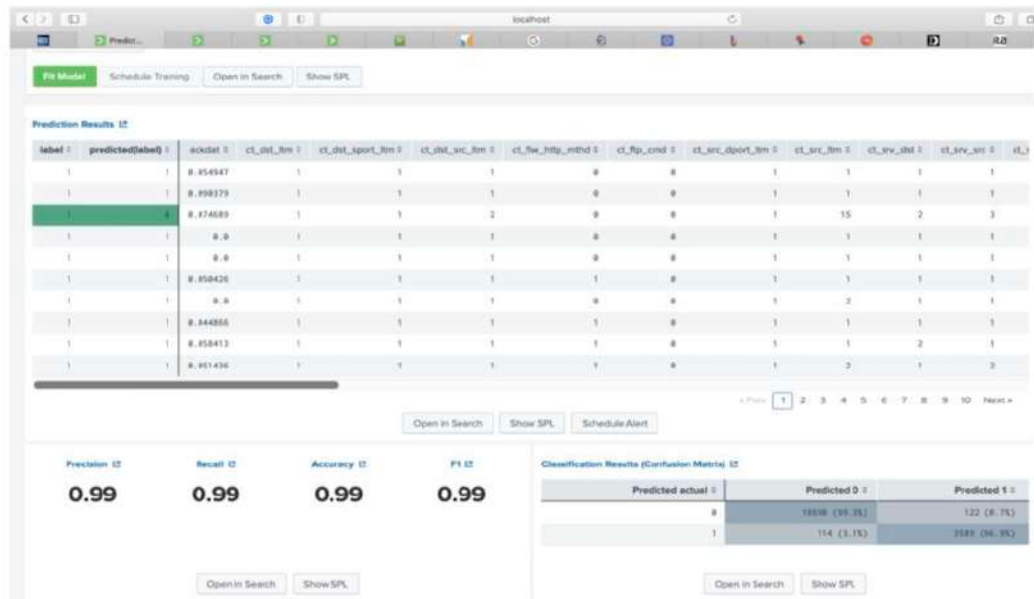


Рисунок 2.4 – Оцінка ефективності класифікації трафіку з метою виявлення DoS-атак із застосуванням алгоритму DT

Результати застосування SVM-алгоритму для ідентифікації DoS-атак у мережевому трафіку наведені на рисунку 2.5.

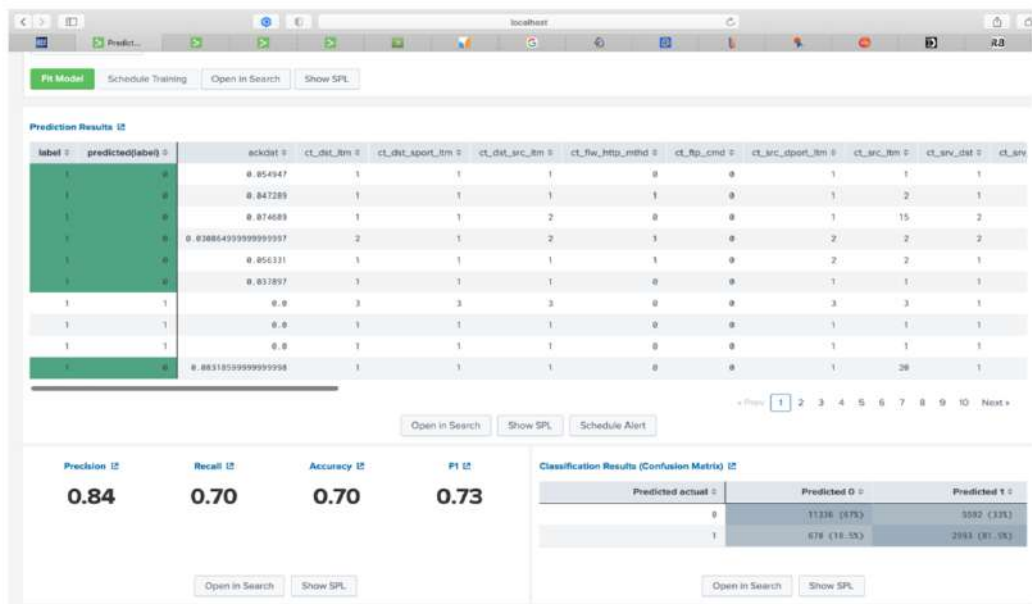


Рисунок 2.5 – Результати застосування SVM-алгоритму для ідентифікації DoS-атак у мережевому трафіку

Проведемо аналіз точності класифікації атак за використанням запропонованих алгоритмів обробки даних.

Таблиця 2.5 – Метрики точності класифікації DoS-атак

| Алгоритм машинного навчання    | Accuracy | Precision | Recall | F1-score |
|--------------------------------|----------|-----------|--------|----------|
| Випадковий ліс (RF)            | 0,9925   | 0,9948    | 0,9961 | 0,9954   |
| Дерево рішень (DT)             | 0,9885   | 0,9932    | 0,9927 | 0,9930   |
| Машина опорних векторів (SVM)  | 0,6985   | 0,8370    | 0,6945 | 0,7329   |
| Подвійний випадковий ліс (DRF) | 0,9984   | 0,9987    | 0,9993 | 0,9990   |

На рисунку 2.6 наведено підсумки класифікації трафіку за допомогою DRF з метою ідентифікації DoS-атак.

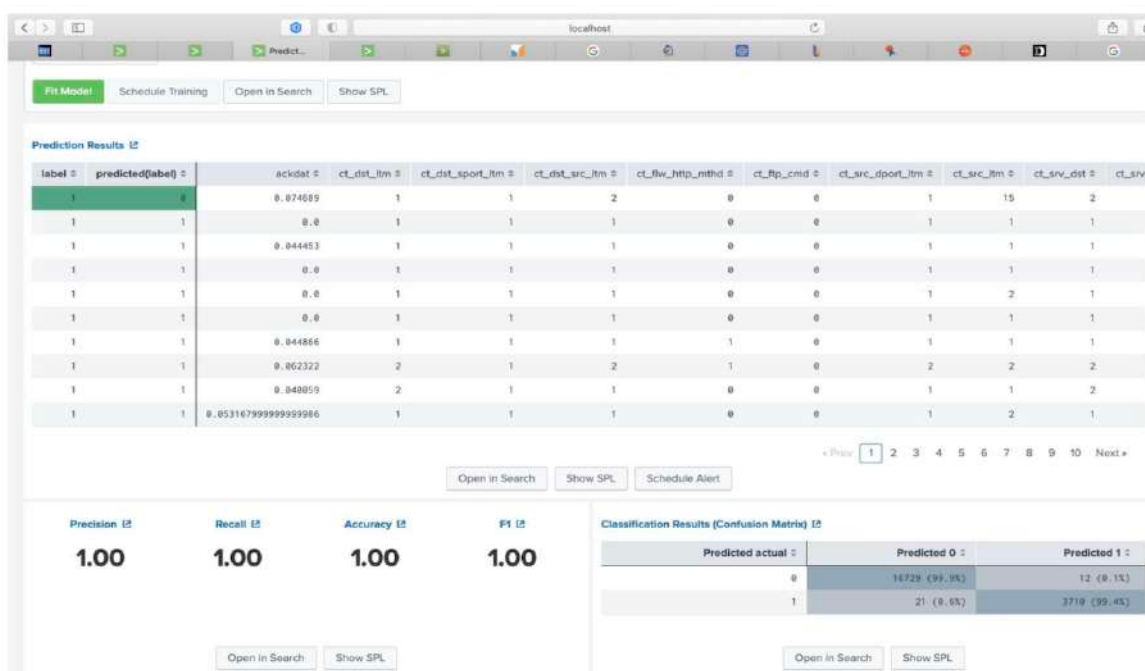


Рисунок 2.6 – Підсумки класифікації трафіку за допомогою DRF з метою ідентифікації DoS-атак

Експериментальні дані, отримані під час тестування навченої моделі на базі даних UNSW-NB15B, демонструють, що алгоритм «випадковий ліс»

забезпечує найкращі показники точності класифікації DoS-атак та мінімальний рівень помилкових спрацьовувань у порівнянні з іншими класичними методами машинного навчання. Результати методу RF також узгоджуються з оцінками точності 99,2% та 98,75%, наведені в роботах [18] та [20] відповідно. Водночас реалізація подвійного випадкового лісу (double random forest, DRF) із базою UNSW-NB15B показала покращення точності класифікації у порівнянні з класичним випадковим лісом.

Представлені результати досліджень демонструють застосування сучасних методів машинного навчання у поєднанні з технологіями збору машинних даних на платформі операційної аналітики Splunk Enterprise Security. Показано, що запропонований підхід до створення систем детектування аномалій та класифікації мережевих атак є ефективним для завдань раннього виявлення та оперативного реагування на інциденти кібербезпеки.

Отримано оцінки продуктивності моделей детектування аномалій на основі класичних алгоритмів машинного навчання – дерева рішень, методу опорних векторів та випадкового лісу – за допомогою відповідних метрик класифікації та розпізнавання. Експериментальні дані підтверджують, що всі розглянуті алгоритми успішно справляються із задачами розпізнавання нормальної та аномальної мережевої активності, а також класифікації мережевих атак.

Розроблений алгоритм побудови дерев рішень, що використовує всі дані навчальної вибірки на кожному проміжному вузлі, включно з кореневим. Такий підхід дозволяє створювати більш розгалужені ансамблі порівняно з класичним методом випадкового лісу, який формує окремі дерева, завантажуючи початкові дані лише у кореневий вузол і використовуючи випадкові підвибірки навчальної бази.

Функціональність Splunk Machine Learning Toolkit була розширена за допомогою додаткового плагіна, розробленого мовою Python, для реалізації методу подвійного випадкового лісу. Для забезпечення високої продуктивності

обробки великих ансамблів дерев використано можливості мови SPL, що дозволяють паралельно будувати дерева та паралельно обчислювати прогнози.

Результати експериментів з детектування аномалій і класифікації мережових атак свідчать, що модифікований алгоритм випадкового лісу забезпечує найвищу точність виявлення DoS-атак у порівнянні з іншими класичними алгоритмами машинного навчання.

## 3 РОЗРОБКА СИСТЕМИ ВИЯВЛЕННЯ ВТОРГНЕНЬ НА БАЗІ МЕТОДІВ МАШИННОГО НАВЧАННЯ

### 3.1 Проектування ML-системи виявлення вторгнень

Сьогодні майже кожен чув про машинне навчання. Розумні колонки та онлайн-кінотеатри ніби відчують ваш настрій і пропонують у рекомендаціях майже ідеальний наступний трек чи фільм. Коли ви дзвоните на «гарячу лінію» банку, часто важко зрозуміти, чи спілкуєтеся з роботом, чи з живою людиною. Безпілотні автомобілі вже їздять звичайними дорогами.

Важко уявити собі проекти та додатки в інформаційній безпеці без використання технологій штучного інтелекту та машинного навчання. У звіті Стенфордського університету «AI Index 2019 Report» на 220 сторінках докладно описано сучасний стан розвитку галузі штучного інтелекту.

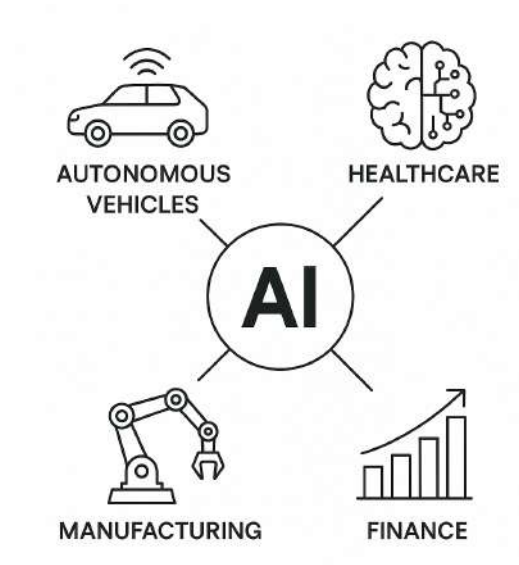


Рисунок 3.1 – Інтеграція ШІ у різні галузі діяльності та прикладні проекти

Розроблювана система виявлення атак покликана не замінити, а доповнити сигнатурний аналізатор з метою підвищення загальної ефективності, особливо у виявленні раніше невідомих атак.

Основні етапи проектування системи виявлення вторгнень на основі машинного навчання включають дибір набору даних для навчання системи детектування комп'ютерних атак; попередню обробку даних; семплювання для боротьби з дисбалансом класів; оцінку важливості та відбір ознак; зменшення розмірності простору ознак; вибір моделі, налаштування параметрів і навчання моделі; тестування та валідація.

Нижче наведено базову схему супервізованого навчання, що ілюструється на рисунку 3.2.

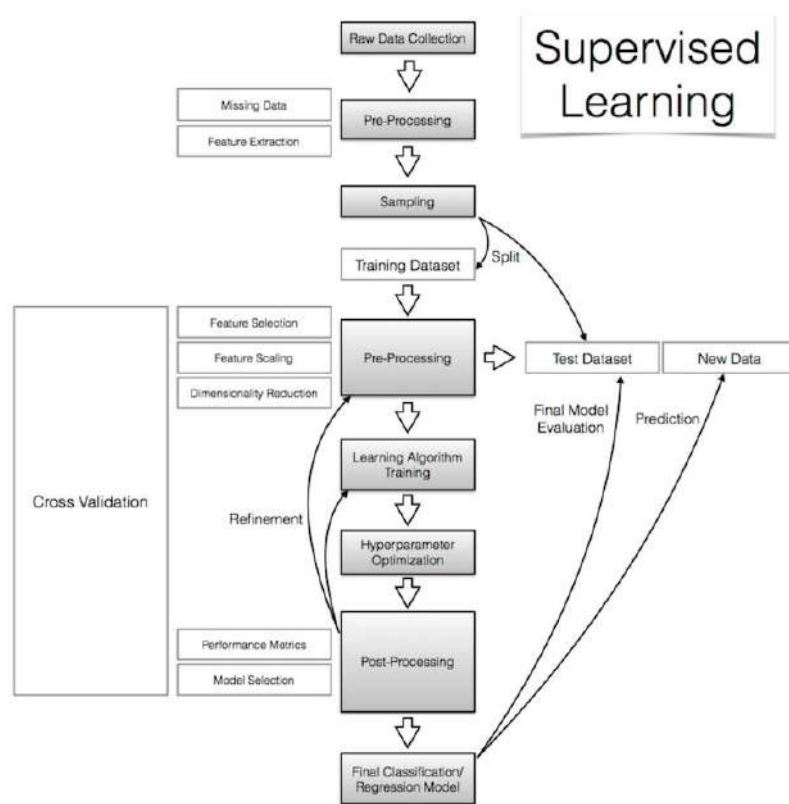


Рисунок 3.2 – Базова схема супервізованого навчання

Навчання системи доступних публічних датасетів (таких як DARPA1998, KDD1999, ISCX2012, ADFA2013 та інші) обрали один із найактуальніших на момент початку дослідження – «Intrusion Detection Evaluation Dataset» CICIDS2017. Розробником цього набору даних є Канадський інститут кібербезпеки (Canadian Institute for Cybersecurity).

CICIDS2017 сформований на основі аналізу мережевого трафіку в ізолюваному середовищі, де моделювалися дії 25 легітимних користувачів, а також різні типи шкідливих атак порушників.

Набір містить понад 50 Гб сирих даних у форматі PCAP і включає 8 попередньо оброблених CSV-файлів з розміченими сесіями та виділеними ознаками, зібраними у різні дні спостережень.

Короткий опис цих файлів і кількісний склад набору наведено у таблицях 3.1 та 3.2.

Таблиця 3.1 – Інформація про структуру та вміст файлів датасету CICIDS2017

| № | Назва файлу                                                | Різновиди атак                                                                 |
|---|------------------------------------------------------------|--------------------------------------------------------------------------------|
| 1 | Monday-WorkingHours.pcap_ISCX.csv                          | Benign (обычный трафик)                                                        |
| 2 | Tuesday-WorkingHours.pcap_ISCX.csv                         | Benign, FTP-Patator, SSH-Patator                                               |
| 3 | Wednesday-workingHours.pcap_ISCX.csv                       | Benign, DoS GoldenEye, DoS Hulk, DoS Slowhttptest, DoS slowloris, Heartbleed   |
| 4 | Thursday-WorkingHours-Morning-WebAttacks.pcap_ISCX.csv     | Benign, Web Attack - Brute Force, Web Attack - Sql Injection, Web Attack - XSS |
| 5 | Thursday-WorkingHours-Afternoon-Infiltration.pcap_ISCX.csv | Benign, Infiltration                                                           |
| 6 | Friday-WorkingHours-Morning.pcap_ISCX.csv                  | Benign, Bot                                                                    |
| 7 | Friday-WorkingHours-Afternoon-PortScan.pcap_ISCX.csv       | Benign, PortScan                                                               |
| 8 | Friday-WorkingHours-Afternoon-DDos.pcap_ISCX.csv           | Benign, DDoS                                                                   |

Нижче наведено обсяг і розподіл даних у наборі CICIDS2017 у таблиці 3.2.

Таблиця 3.2 – Обсяг і розподіл даних у наборі CICIDS2017

| №  | Тип запису                 | Кількість записів |
|----|----------------------------|-------------------|
| 1  | BENIGN                     | 2359087           |
| 2  | DoS Hulk                   | 231072            |
| 3  | PortScan                   | 158930            |
| 4  | DDoS                       | 41835             |
| 5  | DoS GoldenEye              | 10293             |
| 6  | FTP-Patator                | 7938              |
| 7  | SSH-Patator                | 5897              |
| 8  | DoS slowloris              | 5796              |
| 9  | DoS Slowhttptest           | 5499              |
| 10 | Bot                        | 1966              |
| 11 | Infiltration               | 36                |
| 12 | Heartbleed                 | 11                |
| 13 | Web Attack - Brute Force   | 1507              |
| 14 | Web Attack - XSS           | 652               |
| 15 | Web Attack - SQL Injection | 21                |

Один приклад із набору даних CICIDS2017 складається з кількох записів. Кожен із них відповідає мережевій сесії та описується 85 ознаками, як показано на рисунку 3.3.

```
Flow ID, Source IP, Source Port, Destination IP, Destination Port, Protocol, Timestamp,
Flow Duration, Total Fwd Packets, Total Backward Packets, Total Length of Fwd Packets,
Total Length of Bwd Packets, Fwd Packet Length Max, Fwd Packet Length Min, Fwd Packet
Length Mean, Fwd Packet Length St. Flow IAT Max, Flow IAT Min, Fwd IAT Total, Fwd IAT
Mean, Fwd IAT Std, Fwd IAT Max, Fwd IAT Min, Bwd IAT Total, Bwd IAT Mean, Bwd IAT Std,
Bwd IAT Max, Bwd IAT Min, Fwd PSH Flags, Bwd PSH Flags, Fwd URG Flags, Bwd URG Flags,
Fwd Header Length, Bwd Header Length, Fwd Packets/s, Bwd Packets/s, Min Packet Length,
Max Packet Length, Packet Length Mean, Packet Length Std, Packet Variance, FIN Flag Count,
SYN Flag Count, RST Flag Count, PSH Flag Count, ACK Flag Count, URG Flag Count, CWE Flag
Count, ECE Flag Count, Down/Up Ratio, Average Packet Size, Avg Fwd Segment Size, Avg Bwd
Segment Size, Fwd Header Length, Fwd Avg Bytes/Bulk, Fwd Avg Packets/Bulk, Fwd Avg Bulk
Rate, Bwd Avg Bytes/Bulk, Bwd Avg Packets/Bulk, Bwd Avg Bulk Rate, Subflow Fwd Packets,
Subflow Fwd Bytes Bwd Packets, Subflow Bwd Bytes, Init Win bytes forward, Init Win bytes
backward, act data pkt fwd, min seg size forward, Active Mean, Active Std, Active Max,
Active Min, Idle Mean, Idle Std, Idle Max, Idle Min, Label
```

Рисунок 3.3 – Фрагмент даних, що демонструє окремий запис із набору  
CICIDS2017

Далі представлено приклад одного із записів, який буде використаний для  
аналізу.

```
192.168.10.14-65.55.44.109-59135-443-6, 65.55.44.109, 443, 192.168.10.14, 59135, 6,
6/7/2017 9:0 , 6, 0, 6, 6, 6, 0, 250000, 41666.66667, 48, 0, 48, 48, 0, 0, 0, 0, 0,
0, 0, 0, 0, 0, 0, 0, 20, 20, 20833.33333, 20833.33333, 6, 6, 6, 0, 0, 0, 0, 0,
0, 1, 1, 0, 0, 1, 9, 6, 6, 20, 0, 0, 0, 0, 0, 0, 0, 1, 6, 1, 6, 513, 253, 0, 20, 0, 0,
0, 0, 0, 0, 0, 0, BENIGN
```

Рисунок 3.4 – Типовий заповнений запис у таблиці вибірки

Якісні дані є необхідною умовою для створення ефективного  
класифікатора.

У рецензіях на набір даних CICIDS2017 (Intrusion2017, Panigrahi2018,  
Sharafaldin2018) відмічалися проблеми із дисбалансом класів, складною  
структурою файлів та відсутністю деяких значень. Однак ці недоліки  
вважаються не критичними.

Під час аналізу виникли питання щодо точності розмітки даних у наборі  
CICIDS2017, тому розглянемо весь процес обробки – від збору даних за  
допомогою сніфера та передобробки мережевих сесій до побудови моделі  
машинного навчання та тестування в реальному мережевому середовищі.

Точкою для проведення експериментів із набором даних CICIDS2017 стало  
дослідження Kahraman Kostas «Anomaly Detection in Networks Using Machine

Learning». Під час спроби відтворити результати цього дослідження були виявлені розбіжності, а також помилки в коді автора.

Щоб скоротити час обчислень, у навчальній вибірці залишили лише один клас атак – веб-атаки (Brute Force, XSS, SQL Injection). Для цього була сформована підвибірка «WebAttacks» на основі обробки файлу Thursday-WorkingHours-Morning-WebAttacks.pcap\_ISCX.csv із набору CICIDS2017. Набір WebAttacks містить 458 968 записів, з яких 2 180 відповідають веб-атакам, а решта – нормальному трафіку.

Етапи попередньої обробки даних CICIDS2017 та підготовки підвибірки WebAttacks включали:

1. Видалення ознаки «Fwd Header Length.1» (вона дублює ознаку «Fwd Header Length»).
2. Видалення записів із null-значеннями у ідентифікаторі сесії Flow ID (після цього з 458 968 записів залишилось 170 366).
3. Заміну нечислових значень ознак Flow Bytes/s та Flow Packets/s на -1.
4. Заміну невизначених (NaN) та нескінченних значень також на -1.
5. Перетворення рядкових значень ознак Flow ID, Source IP, Destination IP, Timestamp у числові за допомогою label encoding.
6. Кодування міток у навчальній вибірці за правилом: 0 – «немає атаки», 1 – «атака».

Jupyter-блокноти з кодом збережені у репозиторії ml-cybersecurity на Github, а посилання на них надано у Google Colaboratory, що дозволяє запускати код безпосередньо у браузері. Під час етапу попередньої обробки даних були виявлені неточності в ознаковому просторі, зокрема, підозра викликала наявність двох різних ознак із ідентичними значеннями. Для перевірки було прийнято рішення взяти «сирий» мережевий трафік CICIDS2017, виділити з нього мережеві сесії та сформувати власний датасет. Отриманий набір даних мав співпадати з оригінальним датасетом CICIDS2017.

Для цього ми обробили pcap-файл із записаним трафіком за допомогою власного сніффера, витягнули сесії та ознаки, а потім порівняли їх із датасетом

Thursday-WorkingHours-Morning WebAttacks.pcap\_ISCX.csv, намагаючись виявити та виправити розбіжності. Опис процедури знаходження помилок показано на рисунку 3.5.

| Flow ID | Source IP     | Source Port   | Destination | Destination Port | Protocol | Timestamp       | Flow Duration   | Total Fwd Pkts | Total Backward Pkts | Total Length of Fwd Pkts | Total Length of Bwd Pkts |       |
|---------|---------------|---------------|-------------|------------------|----------|-----------------|-----------------|----------------|---------------------|--------------------------|--------------------------|-------|
| 1       | 192.168.10.3  | 192.168.10.50 | 389         | 192.168.10.3     | 389      | 6               | 06.07.2017 8:59 | 113095485      | 48                  | 24                       | 9668                     | 10012 |
| 2       | 192.168.10.3  | 192.168.10.50 | 389         | 192.168.10.3     | 389      | 6               | 06.07.2017 8:59 | 113473706      | 68                  | 40                       | 11364                    | 12718 |
| 3       | 192.168.10.14 | 65.55.44.109  | 59135       | 443              | 6        | 06.07.2017 8:59 | 60261528        | 9              | 7                   | 2130                     | 4221                     |       |
| 4       | 192.168.10.14 | 65.55.44.109  | 59135       | 443              | 6        | 06.07.2017 9:00 |                 | 48             | 1                   | 6                        | 8                        |       |
| 5       | 192.168.10.3  | 192.168.10.19 | 32791       | 192.168.10.3     | 389      | 6               | 06.07.2017 9:00 | 487            | 11                  | 4                        | 172                      |       |
| 6       | 192.168.10.3  | 192.168.10.19 | 41567       | 192.168.10.3     | 88       | 6               | 06.07.2017 9:00 | 1206           | 10                  | 7                        | 3150                     |       |
| 7       | 192.168.10.3  | 192.168.10.19 | 32792       | 192.168.10.3     | 389      | 6               | 06.07.2017 9:00 | 27779          | 17                  | 11                       | 3450                     |       |
| 8       | 192.168.10.3  | 192.168.10.19 | 41569       | 192.168.10.3     | 88       | 6               | 06.07.2017 9:00 | 1133           | 9                   | 6                        | 3150                     |       |
| 9       | 192.168.10.3  | 192.168.10.19 | 32794       | 192.168.10.3     | 389      | 6               | 06.07.2017 9:00 | 118034439      | 106                 | 61                       | 17896                    |       |
| 10      | 192.168.10.3  | 192.168.10.19 | 41567       | 192.168.10.3     | 88       | 6               | 06.07.2017 9:00 | 62             | 1                   | 3                        | 0                        |       |
| 11      | 192.168.10.3  | 192.168.10.19 | 41569       | 192.168.10.3     | 88       | 6               | 06.07.2017 9:00 | 75             | 1                   | 4                        | 12                       |       |
| 12      | 192.168.10.3  | 192.168.10.19 | 32792       | 192.168.10.3     | 389      | 6               | 06.07.2017 9:00 | 75             | 1                   | 4                        | 0                        |       |
| 13      | 192.168.10.3  | 192.168.10.19 | 32791       | 192.168.10.3     | 389      | 6               | 06.07.2017 9:00 | 79             | 1                   | 4                        | 0                        |       |
| 14      | 192.168.10.3  | 192.168.10.19 | 32798       | 192.168.10.3     | 389      | 6               | 06.07.2017 9:00 | 190            | 10                  | 4                        | 160                      |       |
| 15      | 192.168.10.3  | 192.168.10.19 | 32798       | 192.168.10.3     | 389      | 6               | 06.07.2017 9:00 | 2              | 2                   | 0                        | 7                        |       |
| 16      | 192.168.10.3  | 192.168.10.19 | 41574       | 192.168.10.3     | 88       | 6               | 06.07.2017 9:00 | 1086           | 9                   | 6                        | 3150                     |       |
| 17      | 192.168.10.3  | 192.168.10.19 | 44766       | 192.168.10.3     | 3268     | 6               | 06.07.2017 9:00 | 116446503      | 72                  | 47                       | 13380                    |       |
| 18      | 192.168.10.3  | 192.168.10.19 | 389         | 192.168.10.19    | 32798    | 6               | 06.07.2017 9:00 | 50             | 4                   | 0                        | 12                       |       |

Рисунок 3.5 – Опис процедури знаходження помилок

Для аналізу обрана конкретна сесія з відбором пакетів: Flow ID = «192.168.10.14-65.55.44.109-59135-443-6», Source IP = «65.55.44.109». У роботі канадських дослідників два останні пакети було виділено в окрему сесію. Ретельно перевірили вихідні дані сніффера та знайдемо підтвердження в методі addPacket.

Важливі моменти для врахування у нашому сніффері:

1. При отриманні пакета з прапором FIN у напрямку forward потрібно завершувати поточну сесію та починати нову. Щоб два останні пакети (FIN ACK і ACK) опинилися в одній другій сесії, необхідно доповнити умову розриву сесії: сесія має містити більше одного пакета.
2. Завершення сесії за таймером встановлено на 120 секунд, хоча у readme зазначено 600 секунд.

Для цієї ж сесії у прямому напрямку зафіксовано один пакет («Total Fwd Packets» = 1), при цьому загальна довжина переданих пакетів у цьому напрямку («Total Length of Fwd Packets») дорівнює 6. За даними Wireshark, довжина пакета становить 0.

Виявлено різницю у 6 байт, які спускаючись від TCP до рівня Ethernet виявляються доповненням (padding) кадру Ethernet. Існує дискусія, чи варто

враховувати ці 6 байт у довжині ТСР-пакета. Проведена перевірка показана на рисунку 3.6.

| Time  | Source | Source Port   | Destination   | Destination Port | Protocol | Length | TCP Segment Len | Info                                                      |
|-------|--------|---------------|---------------|------------------|----------|--------|-----------------|-----------------------------------------------------------|
| 0.000 | 59135  | 65.55.44.109  | 192.168.10.14 | 443              | TCP      | 66     | 0               | 59135 → 443 [SYN] Seq=0 Win=0 Len=0 MSS=1460              |
| 0.000 | 443    | 192.168.10.14 | 59135         | 59135            | TCP      | 66     | 0               | 443 → 59135 [SYN, ACK] Seq=0 Ack=1 Win=0 Len=0 MSS=1460   |
| 0.000 | 59135  | 65.55.44.109  | 192.168.10.14 | 443              | TCP      | 66     | 0               | 59135 → 443 [ACK] Seq=1 Ack=1 Win=0 Len=0                 |
| 0.000 | 59135  | 65.55.44.109  | 192.168.10.14 | 443              | TLSv1.2  | 264    | 238             | Client Hello                                              |
| 0.000 | 443    | 192.168.10.14 | 59135         | 59135            | TCP      | 2514   | 1460            | 443 → 59135 [ACK] Seq=1 Ack=211 Win=131328 Len=1460       |
| 0.000 | 443    | 192.168.10.14 | 59135         | 59135            | TCP      | 2514   | 1460            | 443 → 59135 [ACK] Seq=1461 Ack=211 Win=131328 Len=1460    |
| 0.000 | 443    | 192.168.10.14 | 59135         | 59135            | TLSv1.2  | 901    | 847             | Server Hello, Certificate, Server Key Exchange, Server    |
| 0.000 | 59135  | 65.55.44.109  | 192.168.10.14 | 443              | TCP      | 66     | 0               | 59135 → 443 [ACK] Seq=211 Ack=2921 Win=66048 Len=0        |
| 0.000 | 59135  | 65.55.44.109  | 192.168.10.14 | 443              | TLSv1.2  | 236    | 182             | Client Key Exchange, Change Cipher Spec, Encrypted Hand   |
| 0.000 | 443    | 192.168.10.14 | 59135         | 59135            | TLSv1.2  | 161    | 107             | Change Cipher Spec, Encrypted Handshake Message           |
| 0.000 | 59135  | 65.55.44.109  | 192.168.10.14 | 443              | TLSv1.2  | 875    | 821             | Application Data                                          |
| 0.000 | 59135  | 65.55.44.109  | 192.168.10.14 | 443              | TLSv1.2  | 1147   | 1093            | Application Data                                          |
| 0.000 | 443    | 192.168.10.14 | 59135         | 59135            | TCP      | 66     | 0               | 443 → 59135 [ACK] Seq=3875 Ack=2387 Win=131328 Len=0      |
| 0.000 | 443    | 192.168.10.14 | 59135         | 59135            | TLSv1.2  | 395    | 341             | Application Data                                          |
| 0.000 | 59135  | 65.55.44.109  | 192.168.10.14 | 443              | TCP      | 66     | 0               | 59135 → 443 [ACK] Seq=2387 Ack=4216 Win=64768 Len=0       |
| 0.000 | 59135  | 65.55.44.109  | 192.168.10.14 | 443              | TCP      | 66     | 0               | 59135 → 443 [FIN, ACK] Seq=2387 Ack=4216 Win=64768 Len=0  |
| 0.000 | 443    | 192.168.10.14 | 59135         | 59135            | TCP      | 66     | 0               | 443 → 59135 [FIN, ACK] Seq=4216 Ack=2388 Win=131328 Len=0 |
| 0.000 | 59135  | 65.55.44.109  | 192.168.10.14 | 443              | TCP      | 66     | 0               | 59135 → 443 [ACK] Seq=2388 Ack=4217 Win=64768 Len=0       |

Frame 120: 60 bytes on wire (480 bits), 60 bytes captured (480 bits) on interface unknown, id 0

Ethernet II, Src: Cisco\_14:eb:31 (00:c1:b1:14:eb:31), Dst: Dell\_36:07:ee (b8:ac:6f:36:07:ee)

Destinations: Dell\_36:07:ee (b8:ac:6f:36:07:ee)

Source: Cisco\_14:eb:31 (00:c1:b1:14:eb:31)

Type: IPv4 (0x0008)

Padding: 000000000000

Internet Protocol Version 4, Src: 65.55.44.109, Dst: 192.168.10.14

Transmission Control Protocol, Src Port: 443, Dst Port: 59135, Seq: 4216, Ack: 2388, Len: 0

0000 b8 ac 6f 36 07 ee c0 c1 b1 14 eb 31 00 00 45 00 ...E...1...A7...  
0010 00 20 41 57 40 00 6c 06 95 1e 41 37 2c 6d c0 a8 ...f...Z...  
0020 0a 0e 01 bb c4 ff 66 ed 7f 5a 24 80 66 54 50 11 ...f...Z...  
0030 02 01 1c 18 00 00 00 00 00 00 00 00 00 00 00 ...f...Z...

Рисунок 3.6 – Оцінка довжини переданого пакета через інтерфейс Wireshark

Перевіряємо інші ознаки для цієї сесії всі значення співпадають, крім Average Packet Size, яке дорівнює 9.

Залишається незрозумілим, як при двох пакетах по 6 байт виходить саме це значення. Водночас Packet Length Mean дорівнює 6, що відповідає очікуванням.

Продовжуючи перевірку інших сесій, виявляємо, що часто спостерігаються невеликі розбіжності у таких ознаках, як Packet Length Mean, Packet Length Std, Packet Length Variance, Average Packet Size, Average Fwd Segment Size, Average Bwd Segment Size.

Аналіз перехоплених пакетів (відновлення вихідних значень на основі середніх показників) свідчить, що при спрацьовуванні тайм-ауту сесії довжина пакета, що відкидається, помилково включається до статистичних розрахунків (див. рисунок 3.7).

| Time     | Source        | Source Port | Destination   | Destination Port | Protocol | Length | TCP Segment Len | Info                              |
|----------|---------------|-------------|---------------|------------------|----------|--------|-----------------|-----------------------------------|
| 0.000000 | 192.168.10.14 | 59135       | 65.55.44.109  | 443              | TCP      | 60     | 0               | 59135 → 443 [SYN] Seq=0 Win=8192  |
| 0.000000 | 65.55.44.109  | 443         | 192.168.10.14 | 59135            | TCP      | 60     | 0               | 443 → 59135 [SYN, ACK] Seq=0 Ack= |
| 0.000000 | 192.168.10.14 | 59135       | 65.55.44.109  | 443              | TCP      | 60     | 0               | 59135 → 443 [ACK] Seq=1 Ack=1 Wi  |
| 0.000000 | 192.168.10.14 | 59135       | 65.55.44.109  | 443              | TLSv1.2  | 264    | 210             | Client Hello                      |
| 0.000000 | 65.55.44.109  | 443         | 192.168.10.14 | 59135            | TCP      | 1514   | 1460            | 443 → 59135 [ACK] Seq=1 Ack=211   |
| 0.000000 | 65.55.44.109  | 443         | 192.168.10.14 | 59135            | TCP      | 1514   | 1460            | 443 → 59135 [ACK] Seq=1461 Ack=2  |
| 0.000000 | 65.55.44.109  | 443         | 192.168.10.14 | 59135            | TLSv1.2  | 901    | 847             | Server Hello, Certificate, Serve  |
| 0.000000 | 192.168.10.14 | 59135       | 65.55.44.109  | 443              | TCP      | 60     | 0               | 59135 → 443 [ACK] Seq=211 Ack=29  |
| 0.000000 | 192.168.10.14 | 59135       | 65.55.44.109  | 443              | TLSv1.2  | 236    | 182             | Client Key Exchange, Change Ciph  |
| 0.000000 | 192.168.10.14 | 59135       | 65.55.44.109  | 443              | TLSv1.2  | 161    | 107             | Change Cipher Spec, Encrypted Ha  |
| 0.000000 | 192.168.10.14 | 59135       | 65.55.44.109  | 443              | TLSv1.2  | 875    | 821             | Application Data                  |
| 0.000000 | 192.168.10.14 | 59135       | 65.55.44.109  | 443              | TLSv1.2  | 1147   | 1093            | Application Data                  |
| 0.000000 | 65.55.44.109  | 443         | 192.168.10.14 | 59135            | TCP      | 60     | 0               | 443 → 59135 [ACK] Seq=3875 Ack=2  |
| 0.000000 | 65.55.44.109  | 443         | 192.168.10.14 | 59135            | TLSv1.2  | 395    | 341             | Application Data                  |
| 0.000000 | 65.55.44.109  | 443         | 192.168.10.14 | 59135            | TCP      | 60     | 0               | 59135 → 443 [ACK] Seq=2307 Ack=4  |
| 0.000000 | 192.168.10.14 | 59135       | 65.55.44.109  | 443              | TCP      | 60     | 0               | 59135 → 443 [FIN, ACK] Seq=2307   |
| 0.000000 | 65.55.44.109  | 443         | 192.168.10.14 | 59135            | TCP      | 60     | 0               | 443 → 59135 [FIN, ACK] Seq=4216   |
| 0.000000 | 192.168.10.14 | 59135       | 65.55.44.109  | 443              | TCP      | 60     | 0               | 59135 → 443 [ACK] Seq=2308 Ack=4  |

> Frame 120: 60 bytes on wire (480 bits), 60 bytes captured (480 bits) on interface unknown, id 0  
 > Ethernet II, Src: Cisco\_14:eb:31 (00:c1:b1:14:eb:31), Dst: Dell\_36:07:ee (b8:ac:6f:36:07:ee)  
 > Destination: Dell\_36:07:ee (b8:ac:6f:36:07:ee)  
 > Source: Cisco\_14:eb:31 (00:c1:b1:14:eb:31)  
 > Type: IPv4 (0x0800)  
 > Padding: 000000000000  
 > Internet Protocol Version 4, Src: 65.55.44.109, Dst: 192.168.10.14  
 > Transmission Control Protocol, Src Port: 443, Dst Port: 59135, Seq: 4216, Ack: 2308, Len: 0

0000 b8 ac 6f 36 07 ee 00 c1 b1 14 eb 31 08 00 45 00 ...e6....1..E  
 0010 00 28 41 57 40 00 6c 05 95 1e 41 37 2c 6d c8 a8 ...A7,m--  
 0020 0a 0e 01 bb e6 ff 65 c4 7f 5a 2d 00 66 54 50 11 .....f..ZS-FTP  
 0030 02 01 1c 18 00 00 00 00 00 00 00 00 00 00 00 .....

Рисунок 3.7 – Обробка трафіку, перехопленого у Wireshark

Після усунення більшості розбіжностей між тестовим датасетом і CICIDS2017 стало зрозуміло, як обчислюються значення ознак на основі записаного трафіку.

### 3.2 Підбір та конфігурація моделі

Формована підвибірка «WebAttacks» є збалансованою щодо класів: із загальної кількості 170 366 записів на клас «без атаки» припадає 168 186 екземплярів, тоді як клас «атака» представлений лише 2 180 записами.

Щоб усунути дисбаланс між класами, було використано метод випадкового зменшення вибірки (undersampling), що передбачає випадкове видалення частини прикладів із переважаючого класу «без атаки».

У результаті було досягнуто цільового співвідношення між класами — 70% для «немає атаки» та 30% для «є атака».

З ознакового простору заздалегідь виключено такі ознаки, як «Flow ID», «Source IP», «Source Port», «Destination IP», «Destination Port», «Protocol»,

«Timestamp», виходячи з припущення, що статистичні ознаки мережевого трафіку мають більшу значущість у загальному випадку.

Крім того, зазначені ознаки адресації легко піддаються підробці злоумисниками, тому їх не враховували при навчанні.

Оцінку значущості ознак проведено за допомогою вбудованого механізму методу `sklearn.ensemble.RandomForestClassifier` (атрибут `feature_importances_`). Перші результати показали сильний зв'язок ознак `Init_Win_bytes_backward` та `Init_Win_bytes_forward` з мітками класів у навчальній вибірці, що свідчить про можливі похибки при формуванні датасету. Ці ознаки було виключено з подальшого аналізу.

Підсумкові результати оцінки впливовості окремих ознак на рисунку 3.8, де наведено топ-20 найбільш важливих ознак.

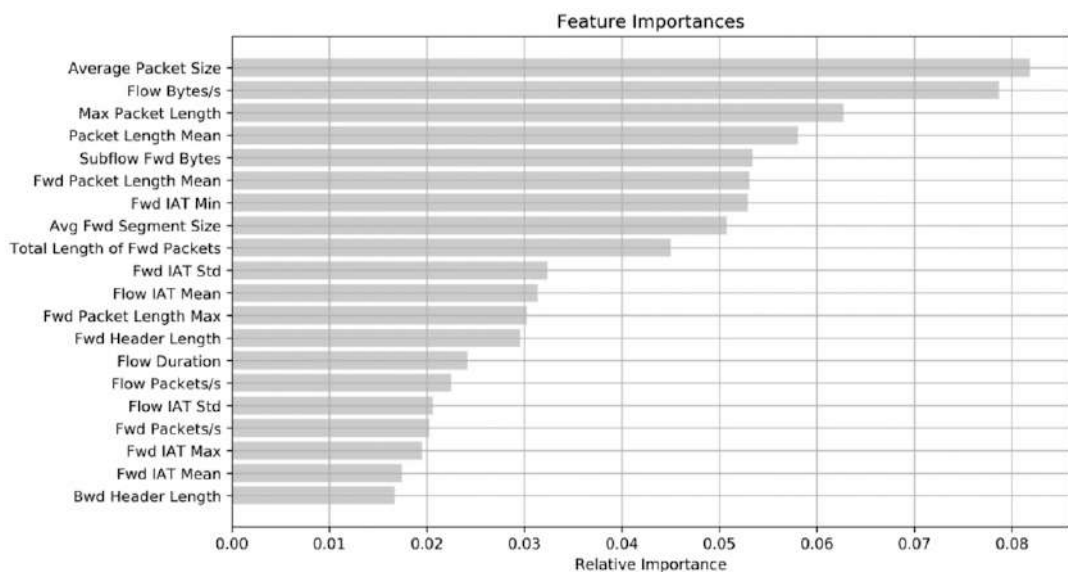


Рисунок 3.8 – Оцінка впливовості окремих ознак

Результати додаткових експериментів засвідчили, що досить високу точність класифікації можна досягти навіть за умови використання лише однієї ознаки – `Init_Win_bytes_backward` або `Init_Win_bytes_forward`.

На наступному рисунку представлено кореляційну матрицю, яка відображає лінійні коефіцієнти кореляції (коефіцієнти Пірсона), розраховані для всіх пар із двадцяти найбільш значущих ознак.

Насиченість кольору клітинок пропорційна величині коефіцієнта кореляції (рисунок 3.9).

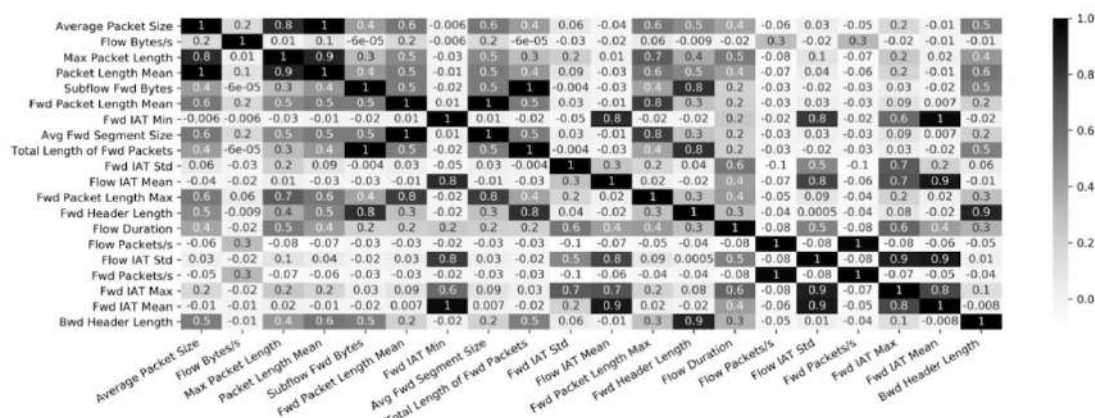


Рисунок 3.9 – Аналіз кореляцій між двадцятьма найважливішими характеристиками

Кореляційний аналіз виявив сильний зв'язок між наступними парами ознак «Average Packet Size» та «Packet Length Mean»; Subflow Fwd Bytes і Total Length of Fwd Packets; «Fwd Packet Length Mean» та «Avg Fwd Segment Size»; «Flow Duration» та «Fwd IAT Total»; «Flow Packets/s» та «Fwd Packets/s»; «Flow IAT Max» та «Fwd IAT Max».

На основі результатів кореляційного аналізу з ознакового простору були вилучені такі ознаки: Packet Length Mean, Subflow Fwd Bytes, Avg Fwd Segment Size, Fwd IAT Total, Fwd Packets/s, Fwd IAT Max.

Після виключення ознак з найменшою значущістю, ознаковий простір було звужено до набору з 10 ознак Average Packet Size – середня довжина поля даних TCP/IP пакета (далі – довжина пакета); Flow Bytes/s – швидкість передачі потоку даних; Max Packet Length – максимальна довжина пакета; Fwd Packet Length Mean – середня довжина пакетів, переданих у прямому напрямку; Fwd IAT Min – мінімальний міжпакетний інтервал (IAT, inter-arrival time) у прямому напрямку; Total Length of Fwd Packets – загальна довжина пакетів, переданих у прямому напрямку; Fwd IAT Std – стандартне відхилення міжпакетного інтервалу в прямому напрямку; Flow IAT Mean – середнє значення міжпакетного інтервалу; Fwd Packet Length Max – максимальна довжина пакета, переданого у

прямому напрямку; Fwd Header Length – сумарна довжина заголовків пакетів, переданих у прямому напрямку.

На етапі вибору моделі було розглянуто 10 найпопулярніших алгоритмів машинного навчання, які оцінювалися за якістю класифікації на підвибірці WebAttacks.

Для порівняння були обрані такі моделі (алгоритми) машинного навчання, з позначеннями та відповідними реалізаціями у пакеті scikit-learn метод k найближчих сусідів (KNN, `sklearn.neighbors.KNeighborsClassifier`); метод опорних векторів (SVM, `sklearn.svm.SVC`); дерево рішень (CART, алгоритм CART, `sklearn.tree.DecisionTreeClassifier`); випадковий ліс (RF, `sklearn.ensemble.RandomForestClassifier`); адаптивний бустинг на основі дерева рішень (AdaBoost, `sklearn.ensemble.AdaBoostClassifier`); логістична регресія (LR, `sklearn.linear_model.LogisticRegression`); Байєсівський класифікатор (NB, `sklearn.naive_bayes.GaussianNB`); лінійний дискримінантний аналіз (LDA, `sklearn.discriminant_analysis.LinearDiscriminantAnalysis`); квадратичний дискримінантний аналіз (QDA, `sklearn.discriminant_analysis.QuadraticDiscriminantAnalysis`); багатошаровий перцептрон (MLP, `sklearn.neural_network.MLPClassifier`).

Якість роботи класифікаторів оцінювалася за допомогою таких метрик:

- точність (accuracy);
- precision (точність, яка показує, наскільки можна довіряти класифікатору);
- recall (повнота, яка характеризує, скільки об'єктів класу «є атака» класифікатор правильно виявляє);
- F1-міра (гармонійне середнє між точністю та повнотою).

Оцінка ефективності моделей виконувалася на збалансованій та попередньо опрацьованій підвибірці WebAttacks із набору даних CICIDS2017 (співвідношення нормального та аномального трафіку – 70% до 30%, з урахуванням 20 найбільш значущих ознак).

У таблиці нижче наведено усереднені за 5-кратною крос-валідацією результати значень цих метрик.

Таблиця 3.3 – Аналіз результатів роботи десяти різних класифікаторів

| Модель (алгоритм) | Acuracy | Precision | Recall | F1    | Час виконання, с |
|-------------------|---------|-----------|--------|-------|------------------|
| KNN               | 0,971   | 0,942     | 0,961  | 0,969 | 4,57             |
| SVM               | 0,705   | 0,669     | 0,036  | 0,602 | 176,04           |
| CART              | 0,975   | 0,973     | 0,946  | 0,969 | 1,53             |
| RF                | 0,971   | 0,978     | 0,943  | 0,970 | 1,14             |
| AdaBoost          | 0,978   | 0,962     | 0,965  | 0,973 | 23,40            |
| LR                | 0,955   | 0,939     | 0,914  | 0,963 | 15,80            |
| Naive Bayes       | 0,722   | 0,520     | 0,956  | 0,754 | 0,47             |
| LDA               | 0,939   | 0,921     | 0,872  | 0,941 | 2,23             |
| QDA               | 0,872   | 0,978     | 0,597  | 0,949 | 1,28             |
| MLP               | 0,904   | 0,921     | 0,912  | 0,776 | 93,83            |

Найкращі результати, як і очікувалося, продемонстрували моделі (алгоритми) KNN, CART, RF, AdaBoost та LR. Враховуючи при цьому мінімальний час виконання, застосування моделі «випадковий ліс» (RF) для розв’язання поставленої задачі є цілком виправданим вибором.

Отже, за основу була взята модель типу «випадковий ліс», реалізована в бібліотеці scikit-learn як RandomForestClassifier.

Серед налаштовуваних гіперпараметрів моделі обрали такі: кількість дерев у лісі (n\_estimators), мінімальна кількість об’єктів у листі дерева (min\_samples\_leaf), максимальна глибина дерева (max\_depth) та максимальна кількість ознак, які використовуються для побудови одного дерева (max\_features).

Ступінь оптимальності параметрів оцінювався за значенням F1-міри. Для уточнення результатів експертного аналізу застосували вбудований метод пошуку гіперпараметрів GridSearchCV із бібліотеки scikit-learn.

В результаті було отримано наступні підсумкові параметри моделі RandomForestClassifier:

```

RandomForestClassifier(
    bootstrap=True, class_weight=None, criterion='gini',
    max_depth=17, max_features=10, max_leaf_nodes=None,
    min_impurity_decrease=0.0, min_impurity_split=None,
    min_samples_leaf=3, min_samples_split=2,
    min_weight_fraction_leaf=0.0, n_estimators=50,
    n_jobs=None, oob_score=False, random_state=1, verbose=0,
    warm_start=False
)

```

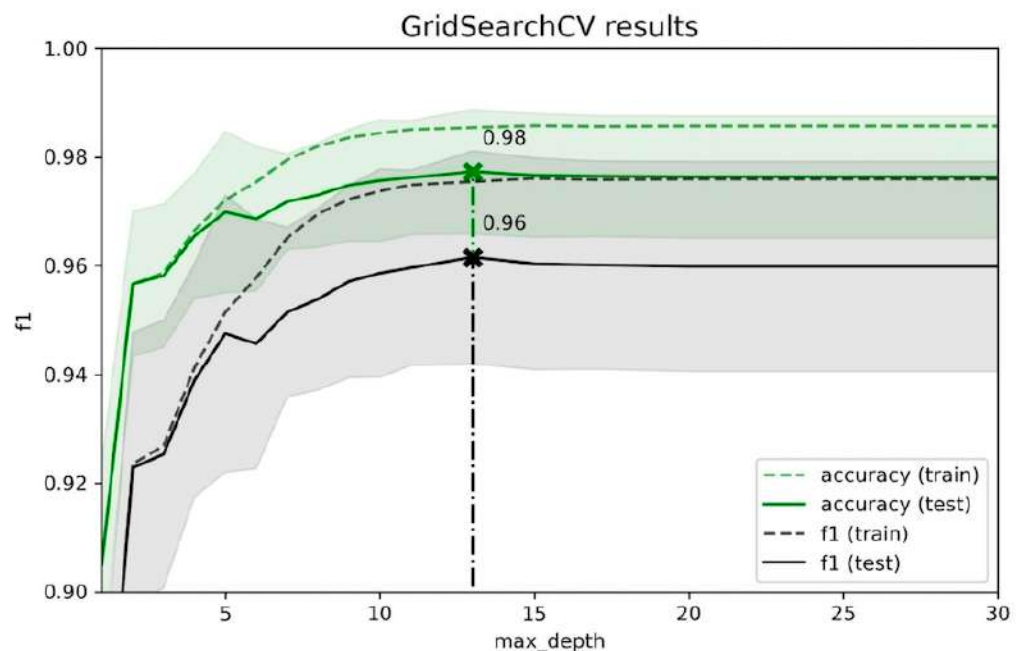


Рисунок 3.10 – Демонстрація процесу налаштування  
RandomForestClassifier

Нижче наведено приклад підбору одного з гіперпараметрів – `max_depth` – при фіксованих значеннях інших параметрів (`n_estimators`, `min_samples_leaf`, `max_features`). На рисунку показано залежність метрики якості (F1-міри) від різних значень параметра `max_depth`.

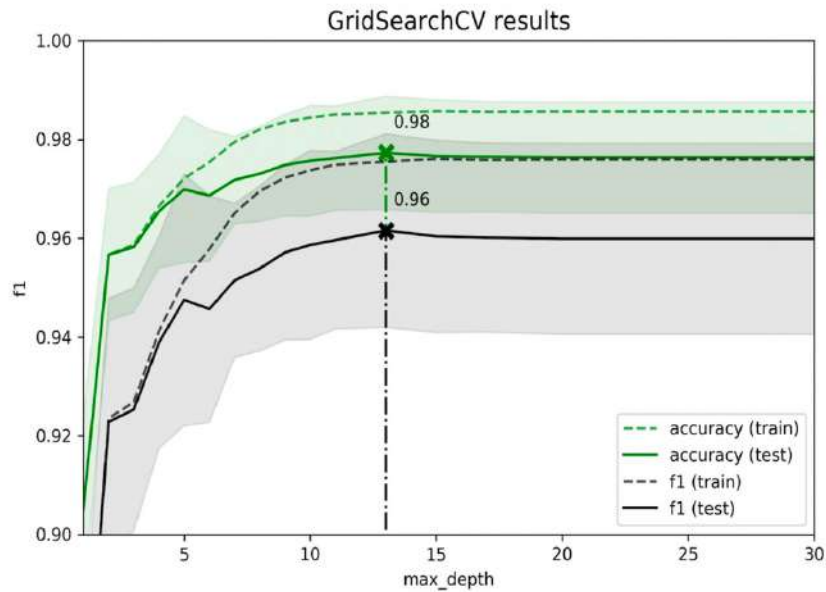


Рисунок 3.11 – Варіація F1-показника моделі при зміні параметра max\_depth

### 3.3 Тестування та оцінка ефективності розробленої системи

Навчена і налаштована модель RandomForestClassifier на тестовій вибірці показала значення повноти (recall) 0.961 та F1-міри 0.971 (перший запуск згідно з протоколом експерименту, див. таблицю нижче).

Досягнуті результати свідчать про можливість подальшого підвищення точності моделі завдяки квазіоптимальному налаштуванню гіперпараметрів (для порівняння у дослідженні Kahraman Kostas recall був 0.94 та F1-мера 0.94, у роботах авторів CICIDS2017 – recall 0.97 та F1-мера 0.97).

Для апробації моделі у реальній мережевій інфраструктурі було розроблено мережевий аналізатор – сніфер на C#.

Цей інструмент дозволяє перехоплювати мережевий трафік і, використовуючи алгоритми реконструкції TCP-сесій із популярних програм Wireshark та TCP Session Reconstruction Tool, виділяти окремі сесії. Для кожної збереженої сесії сніфер формує набір ознак на основі алгоритму CICFlowMeter.

Як веб-додаток, що піддається атаці, використовувалася розроблена консоль адміністратора безпеки на PHP з єдиним активним модулем авторизації, що працює під управлінням веб-сервера Apache.

Конфігурація тестового стенду наведена на рисунку 3.12.

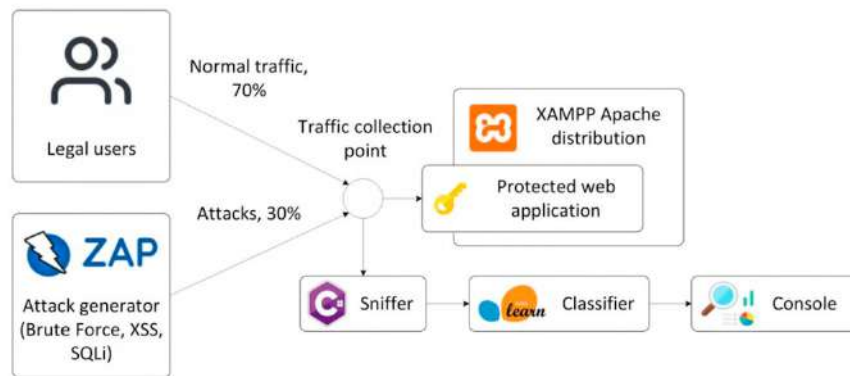


Рисунок 3.12 – Конфігурація тестового стенду

Нормальний трафік формувався запитами легітимних користувачів для підключення до консолі адміністратора та проходження авторизації. Аномальний (шкідливий) трафік імітувався за допомогою програмного засобу OWASP ZAP і включав три види атак Brute Force, XSS та SQL Injection. Співвідношення між нормальним і аномальним трафіком у реальному тестовому наборі даних становило 70% до 30% (див. рисунок 3.13).

| Task ID | Message Type | Код | Причина | RTT   | Size Resp. Header | Size Resp. Body | Highest Alert | Состояние  | Payloads |
|---------|--------------|-----|---------|-------|-------------------|-----------------|---------------|------------|----------|
| 530     | Fuzzed       | 302 | Found   | 12 ms | 237 байт          | 485 байт        |               |            | admin    |
| 0       | Original     | 200 | OK      | 18 ms | 229 байт          | 1 979 байт      | Средний       |            | root     |
| 2       | Fuzzed       | 200 | OK      | 16 ms | 229 байт          | 1 979 байт      |               |            | !admin   |
| 1       | Fuzzed       | 200 | OK      | 21 ms | 229 байт          | 1 979 байт      |               |            | * * #    |
| 4       | Fuzzed       | 200 | OK      | 26 ms | 229 байт          | 1 979 байт      |               |            | \$SRV    |
| 3       | Fuzzed       | 200 | OK      | 25 ms | 229 байт          | 1 979 байт      |               |            | 0        |
| 6       | Fuzzed       | 200 | OK      | 10 ms | 228 байт          | 1 979 байт      |               | Отраженные | *3nogutu |
| 5       | Fuzzed       | 200 | OK      | 25 ms | 229 байт          | 1 979 байт      |               |            | 000000   |
| 8       | Fuzzed       | 200 | OK      | 19 ms | 228 байт          | 1 979 байт      |               |            | 0000     |
| 7       | Fuzzed       | 200 | OK      | 19 ms | 228 байт          | 1 979 байт      |               |            | 05071992 |
| 10      | Fuzzed       | 200 | OK      | 19 ms | 228 байт          | 1 979 байт      |               |            | 00000000 |
| 9       | Fuzzed       | 200 | OK      | 27 ms | 228 байт          | 1 979 байт      |               |            |          |

Рисунок 3.13 – Висновки експериментів, виконаних із використанням розробленого набору даних

Експерименти, проведені на сформованому наборі даних (запуски №2 та №3 у протоколі експерименту), продемонстрували, що застосування моделі, навченої на наборі CICIDS2017, є неможливим з наступних причин: 1. Аналіз навчальної вибірки виявив, що характер атак, змодельованих у дослідженні авторів CICIDS2017, відрізняється від реального стану справ. Зокрема, атаки типу Brute Force в сесіях мають максимальні швидкості до 10 Кбіт/с, що не відповідає реальним випадкам використання автоматизованих засобів перебору паролів. 2. З десяти найважливіших ознак чотири – Flow Bytes/s (швидкість потоку даних), Fwd IAT Min (мінімальний міжпакетний інтервал у прямому напрямку), Flow IAT Std (середньоквадратичне відхилення міжпакетного інтервалу) та Flow IAT Mean (середній міжпакетний інтервал) – безпосередньо залежать від фізичної структури мережі, де відбувається збір трафіку, а також від налаштувань мережевого обладнання. 3. У навчальному наборі сесії з веб-атаками характеризуються низькою швидкістю потоку та великими міжпакетними інтервалами, що не відповідає особливостям реальної мережі Ethernet 100 Мбіт/с. Якісний датасет має відповідати певним вимогам. Автори CICIDS2017 у своїй роботі виділяють 11 таких критеріїв. Найважливіші з них – це забезпечення різноманітності мережевого обладнання, комп'ютерів та операційних систем у тестовій інфраструктурі; різноманітність напрямків мережевого трафіку; використання різних протоколів і типів атак; а також чітке маркування даних для атак і чистого трафіку. Поставимо додаткове завдання – оцінити можливість створення евристичного аналізатора та приблизно визначити його точність. При цьому не ставиться мета зібрати ідеальний датасет, оскільки це завдання для цілих дослідницьких інститутів. План збору датасету:

1. Етап. Запис рсар-файлів та їх очищення. Під час збору «брудного» трафіку змінюємо параметри фазера та вставляємо паузи між фазами, щоб розірвати сесії і збільшити їхню кількість у датасеті. Для «чистого» трафіку моделюємо різноманітні дії користувача.

2. Етап. Передача рсар-файлів сніферу для виділення ознак і об'єднання всіх розмічених записів у єдиний датасет.

Запуск №2. Модель навчали на вибірці WebAttacks із набору CICIDS2017 (трафік збирався в одній мережі). Після цього модель протестували на реальному трафіку іншої мережі з відмінними характеристиками, зокрема швидкістю. Результат виявився незадовільним – значення F1-міри склало лише 0.064. Обчислювальна складність оцінювалася непрямим методом розроблений у Jupyter Notebook прототип системи виявлення веб-атак запускали на ПК з процесором Intel Core i5-2300 @ 2.3 ГГц та 8 ГБ оперативної пам'яті у режимі детекції. Тестовий набір містив близько 70 000 сесій, а час виявлення склав 0,74669 секунди. Таким чином, швидкість детекції веб-атак оцінюється приблизно в 100 000 сесій на секунду.

Таблиця 3.4 – Реєстр проведених експериментів

| Експеримент/Характеристика   | Запуск 1                                                                                                                                                                                                                                         | Запуск 2                                                                             | Запуск 3                                                                                                                                                                                                                       |
|------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Етап навчання моделі         |                                                                                                                                                                                                                                                  |                                                                                      |                                                                                                                                                                                                                                |
| Використовуваний набір даних | Збалансована та передопрацьована підбірка веб-атак WebAttacks набору даних CICIDS2017. 7267 записів, з них 5087 екземплярів класу «немає атаки» та 2180 екземплярів класу «є атака».                                                             |                                                                                      | Сформований набір даних, що відповідають реальному мережному трафіку                                                                                                                                                           |
| Навчальна модель             | 70% записів використовуваного набору даних                                                                                                                                                                                                       |                                                                                      | 70% записів набору даних                                                                                                                                                                                                       |
| Ознаковий простір            | 1. Average Packet Size<br>2. Flow Bytes/s<br>3. Max Packet Length<br>4. Fwd Packet Length Mean<br>5. FwdLATMin<br>6. Total Length of Fwd Packets<br>7. FwdLATStd<br>8. Flow IAT Mean<br>9. Fwd Packet Length <b>Max</b><br>10. Fwd Header Length |                                                                                      | 1. Flow Packets/s<br>2. Flow IAT Max<br>3. Bwd Packet Length Min<br>4. Flow Duration<br>5. Flow IAT Mean<br>6. Flow IAT Std<br>7. Average Packet Size<br>8. Fwd Packet Length Max<br>9. Total Packets<br>10. Fwd Header Length |
| Етап тестування моделі       |                                                                                                                                                                                                                                                  |                                                                                      |                                                                                                                                                                                                                                |
| Тестова вибірка              | 30% записів використовуваного набору даних. Тестова і навчальна вибірка не має перетинів.                                                                                                                                                        | 100% записів сформованого набору даних, що відповідають реальному мережевому трафіку | 30% записів використовуваного набору даних. Тестова та навчальна вибірка не мають перетинів.                                                                                                                                   |
| Значення метрик якості       |                                                                                                                                                                                                                                                  |                                                                                      |                                                                                                                                                                                                                                |
| Accuracy                     | 0.983                                                                                                                                                                                                                                            | 0.456                                                                                | 0.858                                                                                                                                                                                                                          |
| Precision                    | 0.982                                                                                                                                                                                                                                            | 0.812                                                                                | 0.812                                                                                                                                                                                                                          |
| Recall                       | 0.961                                                                                                                                                                                                                                            | 0.033                                                                                | 0.966                                                                                                                                                                                                                          |
| F1                           | 0.971                                                                                                                                                                                                                                            | 0.064                                                                                | 0.882                                                                                                                                                                                                                          |

Завершено експеримент зі створення моделі «випадковий ліс» для задачі виявлення комп'ютерних атак. Модель була навчена на публічному наборі даних CICIDS2017 та протестована в реальних умовах. Налаштування параметрів класифікатора RandomForestClassifier із пакету scikit-learn дало змогу на тестовій вибірці досягти показників повноти (recall) 0.961 та F1-міри 0.971 для даних CICIDS2017, а також 0.966 та 0.882 відповідно для сформованого власного набору.

Основний висновок експерименту методи машинного навчання є практично застосовними для виявлення комп'ютерних атак.

1. Характер атак, змодельованих у навчальному наборі, відрізнявся від реальних атак.
2. Частина ключових ознак тісно пов'язана з фізичною структурою мережі, в якій здійснювався збір трафіку, а також з налаштуваннями мережевого обладнання.

Оптимальним є навчання моделі на наборі даних, розміченому на основі аналізу трафіку саме тієї мережі, що захищається. При використанні моделі, навченої на одній мережі, у іншій (проблема transfer learning), критично важливо, щоб фізична структура мережі та конфігурації обладнання відповідали оригінальним умовам навчання.

## ВИСНОВКИ

У ході наукового дослідження було здійснено всебічний аналіз систем виявлення та запобігання вторгненням (IDS/IPS), що дозволило ідентифікувати їхні переваги та обмеження, зокрема в контексті детектування прихованих атак, реалізованих через стеганоканали.

Досліджено природу стеганоканалів як інструменту прихованого перенесення даних, а також фактори, що ускладнюють їх виявлення традиційними методами. Встановлено, що ефективне виявлення таких атак потребує побудови розширеного набору індикаторів компрометації.

Особливу увагу приділено індикаторам компрометації (IoC), сформованим за допомогою методів штучного інтелекту на основі аналізу мережевого трафіку. Було показано, що такі індикатори можуть підвищити точність і швидкість виявлення складних атак.

Проведено оцінку ефективності платформи Splunk Machine як інструмента для побудови моделей детектування аномалій у мережевому трафіку. Результати свідчать про її доцільне застосування в рамках системи виявлення атак на основі міток компрометації.

Розроблено та досліджено класифікатори атак на основі IoC із використанням алгоритмів машинного навчання. Було обрано відповідний набір даних, який враховує особливості атак, пов'язаних із прихованим передаванням інформації.

Виконано попередню обробку даних, а саме проведено балансування класів, оцінено значущість ознак, здійснено їх відбір та зменшення розмірності ознакового простору. Це дозволило підвищити узагальнюючу здатність моделей.

Спроектовано та оптимізовано модель машинного навчання для виявлення атак, побудовану на мітках компрометації. Результати тестування показали високу ефективність розробленої моделі щодо виявлення як типових, так і складно детектованих атак.

Таким чином, поставлені в дослідженні завдання виконані в повному обсязі, а отримані результати можуть бути використані для удосконалення систем кіберзахисту, зокрема в частині виявлення прихованих та складних атак у комп'ютерних мережах.

## СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Zhijun W. et al. Low-rate DoS attacks, detection, defense, and challenges: a survey // IEEE Access. – 2020. – Vol. 8. – P. 43920–43943.
2. Hristov M. et al. Integration of Splunk Enterprise SIEM for DDoS Attack Detection in IoT // 2021 IEEE 20th International Symposium on Network Computing and Applications (NCA). – IEEE, 2021. – P. 1–5.
3. Gadze J. D. et al. An investigation into the application of deep learning in the detection and mitigation of DDOS attack on SDN controllers // Technologies. – 2021. – Vol. 9, No. 1. – P. 14.
4. Awan M. J. et al. Real-time DDoS attack detection system using big data approach // Sustainability. – 2021. – Vol. 13, No. 19. – P. 10743.
5. Han S., Kim H., Lee Y. S. Double random forest // Machine Learning. – 2020. – P. 1569–1586. DOI: <https://doi.org/10.1007/s10994-020-05889-1>
- Singh R., Mannepalli P.-K. Survey on Feature Reduction Techniques of Intrusion Detection System // International Journal of Engineering Research in Current Trends (IJERCT). – 2020. – Vol. 2, Issue 3. – ISSN 2582-5488.
7. Artificial intelligence and machine learning in cybersecurity: a deep dive into state-of-the-art techniques and future paradigms | Knowledge and Information Systems <https://link.springer.com/article/10.1007/s10115-025-02429-y>
8. Artificial intelligence and machine learning in cybersecurity: a deep dive into state-of-the-art techniques and future paradigms | Knowledge and Information Systems <https://link.springer.com/article/10.1007/s10115-025-02429-y>
9. Artificial intelligence and machine learning in cybersecurity: a deep dive into state-of-the-art techniques and future paradigms | Knowledge and Information Systems <https://link.springer.com/article/10.1007/s10115-025-02429-y>
10. What Is Anomaly Based Detection System? | Fidelis Security <https://fidelissecurity.com/cybersecurity-101/learn/anomaly-based-detection-system/>
11. What Is Anomaly Based Detection System? | Fidelis Security <https://fidelissecurity.com/cybersecurity-101/learn/anomaly-based-detection-system/>

12. AI and Zero-Day Attack Detection: Anticipating Unknown Threats | by Megasis Network | Medium <https://megasisnetwork.medium.com/ai-and-zero-day-attack-detection-anticipating-unknown-threats-c0a3a627a7d6>
13. What Is Anomaly Based Detection System? | Fidelis Security <https://fidelissecurity.com/cybersecurity-101/learn/anomaly-based-detection-system/>
14. What is MITRE ATT&CK®: An Explainer | Exabeam <https://www.exabeam.com/explainers/mitre-attck/what-is-mitre-attck-an-explainer/>
15. MITRE ATT&CK Framework Guide for Beginners | Picus <https://www.picussecurity.com/mitre-attack-framework-beginners-guide>
16. What is MITRE ATT&CK®: An Explainer | Exabeam <https://www.exabeam.com/explainers/mitre-attck/what-is-mitre-attck-an-explainer/>
17. MITRE ATT&CK Framework Guide for Beginners | Picus <https://www.picussecurity.com/mitre-attack-framework-beginners-guide>
18. Big Data Analytics in Cyber Security: Enhancing Threat Detection <https://www.puppygraph.com/blog/big-data-analytics-in-cybersecurity>
19. Big Data Analytics in Cyber Security: Enhancing Threat Detection <https://www.puppygraph.com/blog/big-data-analytics-in-cybersecurity>
20. Big Data Analytics in Cyber Security: Enhancing Threat Detection <https://www.puppygraph.com/blog/big-data-analytics-in-cybersecurity>
21. Big Data Analytics in Cyber Security: Enhancing Threat Detection <https://www.puppygraph.com/blog/big-data-analytics-in-cybersecurity>
22. Hadoop vs Spark: Key Differences in Big Data Analytics <https://www.veritis.com/blog/hadoop-vs-spark-all-you-need-to-know-about-big-data-analytics/>
23. Big Data Analytics in Cyber Security: Enhancing Threat Detection <https://www.puppygraph.com/blog/big-data-analytics-in-cybersecurity>
24. Meet Apache Spot, a new open source project for cybersecurity | PCWorld <https://www.pcworld.com/article/410492/meet-apache-spot-a-new-open-source-project-for-cybersecurity.html>

25. Meet Apache Spot, a new open source project for cybersecurity | PCWorld  
<https://www.pcworld.com/article/410492/meet-apache-spot-a-new-open-source-project-for-cybersecurity.html>
26. Meet Apache Spot, a new open source project for cybersecurity | PCWorld  
<https://www.pcworld.com/article/410492/meet-apache-spot-a-new-open-source-project-for-cybersecurity.html>
- Advanced threat detection with User and Entity Behavior Analytics (UEBA) in Microsoft Sentinel | Microsoft Learn <https://learn.microsoft.com/en-us/azure/sentinel/identify-threats-with-entity-behavior-analytics>
28. Update to UEBA gives a better understanding of risks and better view of your security data <https://www.logpoint.com/en/blog/product-releases/update-to-ueba-gives-a-better-understanding-of-risks-and-better-view-of-your-security-data/>
29. Zhijun, W., et al. (2020). Low-rate DoS attacks, detection, defense, and challenges: a survey. *IEEE Access*, 8, 43920–43943.
30. Hristov, M., et al. (2021). Integration of Splunk Enterprise SIEM for DDoS Attack Detection in IoT. *2021 IEEE 20th International Symposium on Network Computing and Applications (NCA)*, IEEE, 1–5.
31. Gadze, J. D., et al. (2021). An investigation into the application of deep learning in the detection and mitigation of DDOS attack on SDN controllers. *Technologies*, 9(1), 14.
32. Awan, M. J., et al. (2021). Real-time DDoS attack detection system using big data approach. *Sustainability*, 13(19), 10743.
33. Han, S., Kim, H., & Lee, Y. S. (2020). Double random forest. *Machine Learning*, 1569–1586. <https://doi.org/10.1007/s10994-020-05889-1>
34. Singh, R., & Mannepalli, P.-K. (2020). Survey on Feature Reduction Techniques of Intrusion Detection System. *International Journal of Engineering Research in Current Trends (IJERCT)*, 2(3).
35. (Оглед) Artificial intelligence and machine learning in cybersecurity: a deep dive into state-of-the-art techniques and future paradigms. *Knowledge and Information Systems* (Springer). (2025).

36. Fidelis Security. What Is Anomaly Based Detection System? (онлайн-стаття). <https://fidelissecurity.com/cybersecurity-101/learn/anomaly-based-detection-system/>
37. Megasis Network (Medium). AI and Zero-Day Attack Detection: Anticipating Unknown Threats. (стаття, 2024). <https://megasisnetwork.medium.com/...>
38. Exabeam. What is MITRE ATT&CK®: An Explainer. (онлайн-ресурс). <https://www.exabeam.com/explainers/mitre-attck/what-is-mitre-attck-an-explainer/>
39. Picus. MITRE ATT&CK Framework Guide for Beginners. (онлайн-ресурс). <https://www.picussecurity.com/mitre-attack-framework-beginners-guide>
40. PuppyGraph / блог/аналітика. Big Data Analytics in Cyber Security: Enhancing Threat Detection. (онлайн-стаття).
41. Veritis. Hadoop vs Spark: Key Differences in Big Data Analytics. (онлайн-стаття).
42. PCWorld. Meet Apache Spot, a new open source project for cybersecurity. (стаття про Apache Spot). <https://www.pcworld.com/article/410492/meet-apache-spot-a-new-open-source-project-for-cybersecurity.html>
43. Microsoft Learn. Advanced threat detection with UEBA in Microsoft Sentinel. (документація/гайд). <https://learn.microsoft.com/en-us/azure/sentinel/identify-threats-with-entity-behavior-analytics>
44. Logpoint Blog. Update to UEBA gives a better understanding of risks and better view of your security data.(продуктовий блог).
45. Mischinger, M., et al. (2024). IoC Stalker: Early detection of Indicators of Compromise. (конференційна доповідь / PDF). [suarez-tangil.networks.imdea.org](https://suarez-tangil.networks.imdea.org)
46. Alaeifar, P., et al. (2024). Current approaches and future directions for Cyber Threat Intelligence: a comprehensive survey. *ScienceDirect / Computers & Security* (огляд). [ScienceDirect](https://www.sciencedirect.com)

47. Marchiori, F., et al. (2023). STIXnet: A Novel and Modular Solution for Extracting All STIX Entities. *ACM*(conference/journal paper). [dl.acm.org](https://dl.acm.org)
48. Vasani, V., et al. (2023). Comprehensive Analysis of Advanced Techniques and Open-source EDR/Threat Detection Tools. *Electronics (MDPI)*, 12(20), 4299. [MDPI](https://www.mdpi.com/12/20/4299)
49. Dardouri, S., et al. (2025). A deep learning/machine learning approach for anomaly-based network intrusion detection. *Frontiers in Artificial Intelligence* (article). [Frontiers](https://www.frontiersin.org)
50. Schummer, P. (2024). Machine Learning-Based Network Anomaly Detection. *MDPI* (article on NIDS with ML). [MDPI](https://www.mdpi.com/12/20/4299)
51. Arkan, S. M., et al. (2024). Automating shareable cyber threat intelligence production: methods and challenges. *The Computer Journal / Springer* (paper). [SpringerLink](https://www.springerlink.com)
52. Rahman, M. M., et al. (2025). A survey on intrusion detection systems in IoT networks. *ScienceDirect* (survey). [ScienceDirect](https://www.sciencedirect.com)
53. Zhong, Y., et al. (2024). RFG-HELAD: Robust Fine-Grained Network Traffic Anomaly Detection (IEEE TIFS).(пов'язане дослідження DL-моделей для аномалій). [arXiv](https://arxiv.org)
54. Research/Review. Deep Learning-based Intrusion Detection Systems: A Survey. *arXiv* (2025) — огляд тенденцій застосування DL у виявленні шкідливого ПЗ та аномалій. [arXiv](https://arxiv.org)
55. Kodituwakku, A., et al. (2025). A Zero-Configuration Agentless Endpoint Detection and Response solution. *MDPI Electronics* (стаття). [MDPI](https://www.mdpi.com/12/20/4299)
56. Santos, P., et al. (2025). A Systematic Review of Cyber Threat Intelligence (CTI). *Sensors (MDPI)* (systematic review). [MDPI](https://www.mdpi.com/12/20/4299)
57. Park, J., et al. (2022). Evaluation of open-source EDR approaches (Google Rapid Response + osquery) — comparative study. *conference/journal article* (EDR effectiveness). [MDPI](https://www.mdpi.com/12/20/4299)

58. (Стандарт/технічна документація) STIX 2.1 specification; TAXII 2.1 protocol; OpenIOC / YARA rulesets — офіційні специфікації та керівництва (прим.: корисні доки для практики IOC-обміну й ідентифікації).
59. MITRE ATT&CK Framework – <https://attack.mitre.org/>
60. MITRE Engage Framework – <https://engage.mitre.org/>
61. CISA – Cybersecurity & Infrastructure Security Agency  
– <https://www.cisa.gov/>
62. US-CERT National Cyber Awareness System – <https://www.us-cert.gov/>
63. ENISA Threat Landscape – <https://www.enisa.europa.eu/topics/threat-risk-management/threats-and-trends>
64. VirusTotal Threat Intelligence – <https://www.virustotal.com/>
65. Hybrid Analysis Sandbox – <https://www.hybrid-analysis.com/>
66. ANY.RUN Interactive Malware Sandbox – <https://any.run/>
67. MalwareBazaar (abuse.ch) – <https://bazaar.abuse.ch/>
68. ThreatFox (abuse.ch) – <https://threatfox.abuse.ch/>
69. OpenCTI Threat Intelligence Platform – <https://www.opencti.io/>
70. MISP (Malware Information Sharing Platform) – <https://www.misp-project.org/>
71. IBM X-Force Exchange – <https://exchange.xforce.ibmcloud.com/>
72. AlienVault OTX (Open Threat Exchange) – <https://otx.alienvault.com/>
73. CrowdStrike Threat Intelligence Blog  
– <https://www.crowdstrike.com/blog/>
74. Palo Alto Networks Unit42 – <https://unit42.paloaltonetworks.com/>
75. FireEye (Trellix) Threat Research – <https://www.trellix.com/en-us/about/newsroom.html>
76. Kaspersky Securelist – <https://securelist.com/>
77. ESET Threat Research Blog – <https://www.welivesecurity.com/>
78. Cisco Talos Intelligence Group – <https://talosintelligence.com/>
79. FortiGuard Labs Threat Research  
– <https://www.fortinet.com/blog/threat-research>

- 80. SANS Internet Storm Center – <https://isc.sans.edu/>
- 81. The DFIR Report – <https://thedfirreport.com/>
- 82. Malware Traffic Analysis – <https://www.malware-traffic-analysis.net/>
- 83. BleepingComputer Cybersecurity News  
– <https://www.bleepingcomputer.com/>
- 84. The Hacker News – <https://thehackernews.com/>
- 85. DarkReading Cybersecurity Insights – <https://www.darkreading.com/>
- 86. Krebs on Security – <https://krebsonsecurity.com/>
- 87. Cybersecurity & Infrastructure Security Blog (Microsoft)  
– <https://techcommunity.microsoft.com/t5/security-compliance-and-identity/bg-p/SecurityComplianceIdentity>
- 88. Google Threat Analysis Group (TAG) Blog – <https://blog.google/threat-analysis-group/>

## ДОДАТКИ

## Словник (глосарій) термінів і аббревіатур

| №  | Термін /<br>Абревіатура                                         | Визначення                                                                                    |
|----|-----------------------------------------------------------------|-----------------------------------------------------------------------------------------------|
| 1  | Індикатор<br>компрометації<br>(IOC, Indicator of<br>Compromise) | Дані, які свідчать про можливу наявність<br>кіберінциденту або компрометацію системи.         |
| 2  | Кібербезпека<br>(Cybersecurity)                                 | Сукупність методів, технологій та процесів для<br>захисту комп'ютерних систем, мереж і даних. |
| 3  | Кіберзагроза<br>(Cyber Threat)                                  | Потенційна дія або подія, що може завдати шкоди<br>інформаційним активам.                     |
| 4  | Кіберінцидент<br>(Cyber Incident)                               | Подія, яка негативно впливає на конфіденційність,<br>цілісність або доступність інформації.   |
| 5  | Атака (Attack)                                                  | Несанкціонована дія, спрямована на порушення<br>безпеки інформаційної системи.                |
| 6  | APT (Advanced<br>Persistent Threat)                             | Тривала, цілеспрямована атака, зазвичай підтримувана<br>ресурсами державного рівня.           |
| 7  | Malware<br>(Шкідливе<br>програмне<br>забезпечення)              | Програми, створені для завдання шкоди або<br>отримання несанкціонованого доступу.             |
| 8  | Spyware                                                         | Програма, що непомітно збирає дані про користувача<br>чи систему.                             |
| 9  | Ransomware                                                      | Шкідливе ПЗ, яке блокує доступ до даних та вимагає<br>викуп.                                  |
| 10 | Trojan (Троян)                                                  | Програма, що маскується під легітимну, але виконує<br>шкідливі дії.                           |

|    |                                                        |                                                                       |
|----|--------------------------------------------------------|-----------------------------------------------------------------------|
| 11 | Phishing<br>(Фішинг)                                   | Соціотехнічна атака для викрадення конфіденційних даних через обман.  |
| 12 | Spear Phishing                                         | Цілеспрямований фішинг проти конкретної особи або організації.        |
| 13 | DDoS<br>(Distributed Denial of Service)                | Атака, що перевантажує ресурси системи численними запитами.           |
| 14 | Botnet                                                 | Мережа заражених пристроїв, які діють під контролем зловмисника.      |
| 15 | Command and Control (C2)                               | Сервер або система, через яку зловмисник керує зараженими пристроями. |
| 16 | SIEM (Security Information and Event Management)       | Система збору, кореляції та аналізу подій безпеки.                    |
| 17 | SOC (Security Operations Center)                       | Центр моніторингу та реагування на кіберінциденти.                    |
| 18 | SOAR (Security Orchestration, Automation and Response) | Інструмент автоматизації процесів реагування на інциденти.            |
| 19 | Threat Intelligence<br>(Кіберрозвідка)                 | Збір і аналіз інформації про актуальні кіберзагрози.                  |
| 20 | TTP (Tactics, Techniques, and Procedures)              | Модель опису поведінки зловмисників.                                  |

|    |                                            |                                                                        |
|----|--------------------------------------------|------------------------------------------------------------------------|
| 21 | MITRE ATT&CK                               | База знань тактик і технік, які використовують зловмисники.            |
| 22 | Kill Chain                                 | Модель життєвого циклу кібератаки від розвідки до реалізації.          |
| 23 | Network Traffic Analysis (NTA)             | Аналіз трафіку для виявлення підозрілої активності.                    |
| 24 | Endpoint Detection and Response (EDR)      | Технологія моніторингу та реагування на загрози на кінцевих пристроях. |
| 25 | XDR (Extended Detection and Response)      | Інтегрована платформа детектування загроз на різних рівнях системи.    |
| 26 | IDS (Intrusion Detection System)           | Система виявлення вторгнень у мережу.                                  |
| 27 | IPS (Intrusion Prevention System)          | Система запобігання вторгнень у режимі реального часу.                 |
| 28 | Firewall (Міжмережевий екран)              | Пристрій або програма, що контролює мережевий трафік.                  |
| 29 | Sandbox (Пісочниця)                        | Безпечне середовище для аналізу потенційно шкідливого коду.            |
| 30 | Hash (Геш)                                 | Унікальний цифровий відбиток файлу або даних.                          |
| 31 | MD5, SHA-1, SHA-256                        | Алгоритми гешування, що використовуються для ідентифікації файлів.     |
| 32 | Digital Forensics (Цифрова криміналістика) | Аналіз цифрових даних для розслідування інцидентів.                    |

|    |                                                 |                                                         |
|----|-------------------------------------------------|---------------------------------------------------------|
| 33 | Log Analysis<br>(Аналіз журналів подій)         | Вивчення логів для виявлення аномальної активності.     |
| 34 | Incident Response<br>(Реагування на інциденти)  | Сукупність дій для ліквідації наслідків кіберінциденту. |
| 35 | Playbook<br>(Плейбук)                           | Набір типових сценаріїв реагування на інциденти.        |
| 36 | Threat Hunting                                  | Проактивний пошук загроз у мережі.                      |
| 37 | Zero-Day<br>(Нульовий день)                     | Уразливість, про яку ще не знають розробники ПЗ.        |
| 38 | Exploit                                         | Код або техніка, що використовує уразливість системи.   |
| 39 | Patch Management                                | Управління оновленнями для усунення уразливостей.       |
| 40 | Vulnerability<br>(Уразливість)                  | Слабке місце в системі безпеки.                         |
| 41 | CVSS (Common Vulnerability Scoring System)      | Стандарт оцінки критичності уразливостей.               |
| 42 | IOC Feed                                        | Потік даних з індикаторами компрометації.               |
| 43 | STIX (Structured Threat Information Expression) | Формат структурованого опису даних про загрози.         |
| 44 | TAXII (Trusted Automated Exchange of            | Протокол обміну інформацією про кіберзагрози.           |

|    |                                       |                                                                |
|----|---------------------------------------|----------------------------------------------------------------|
|    | Indicator Information)                |                                                                |
| 45 | OpenIOC                               | Формат представлення індикаторів компрометації.                |
| 46 | YARA                                  | Інструмент для класифікації та виявлення шкідливих файлів.     |
| 47 | Snort / Suricata                      | Системи виявлення вторгнень із сигнатурним аналізом.           |
| 48 | Syslog                                | Стандарт обміну повідомленнями журналів подій.                 |
| 49 | DNS Poisoning                         | Маніпуляція DNS-запитами для перенаправлення трафіку.          |
| 50 | IP Reputation                         | Оцінка надійності IP-адреси на основі історії активності.      |
| 51 | TLS/SSL                               | Протоколи шифрування даних у мережі.                           |
| 52 | Encryption<br>(Шифрування)            | Перетворення даних у форму, недоступну стороннім.              |
| 53 | Decryption<br>(Розшифрування)         | Зворотне перетворення за допомогою ключа.                      |
| 54 | Public Key<br>Infrastructure<br>(PKI) | Система управління цифровими сертифікатами.                    |
| 55 | Certificate<br>Authority (CA)         | Організація, що видає цифрові сертифікати.                     |
| 56 | Threat Actor                          | Особа або група, що здійснює кібератаки.                       |
| 57 | Insider Threat                        | Загроза з боку внутрішнього користувача системи.               |
| 58 | Backdoor                              | Прихований спосіб доступу до системи.                          |
| 59 | Rootkit                               | Інструмент для приховування присутності зловмисника в системі. |
| 60 | Keylogger                             | Програма, що записує натискання клавіш.                        |
| 61 | Payload                               | Частина шкідливого коду, яка виконує атаку.                    |

|    |                               |                                                              |
|----|-------------------------------|--------------------------------------------------------------|
| 62 | Brute Force Attack            | Підбір паролів шляхом перебору комбінацій.                   |
| 63 | Credential Stuffing           | Використання викрадених облікових даних для входу в системи. |
| 64 | Data Breach                   | Несанкціоноване розкриття або викрадення даних.              |
| 65 | Data Exfiltration             | Таємне виведення конфіденційних даних із системи.            |
| 66 | Threat Modeling               | Методологія аналізу потенційних загроз системі.              |
| 67 | Security Policy               | Сукупність правил забезпечення інформаційної безпеки.        |
| 68 | Security Audit                | Перевірка ефективності заходів безпеки.                      |
| 69 | Penetration Testing (Pentest) | Імітація атаки для оцінки захищеності системи.               |
| 70 | Red Team                      | Група, що тестує безпеку шляхом симуляції атак.              |
| 71 | Blue Team                     | Група, що відповідає за захист та моніторинг.                |
| 72 | Purple Team                   | Колаборація між Red і Blue Team для підвищення ефективності. |
| 73 | Security Baseline             | Набір мінімальних вимог до безпеки системи.                  |
| 74 | Risk Assessment               | Процес оцінки ризиків для інформаційних активів.             |
| 75 | Business Continuity           | Підхід до підтримки функціонування бізнесу після інцидентів. |
| 76 | Disaster Recovery             | Відновлення систем після серйозних збоїв.                    |
| 77 | Asset Management              | Облік і контроль інформаційних активів.                      |
| 78 | Data Loss Prevention (DLP)    | Технології запобігання витоку даних.                         |
| 79 | Access Control                | Механізми управління доступом користувачів.                  |
| 80 | Identity and Access           | Система контролю автентифікації та авторизації.              |

|    |                                    |                                                                |
|----|------------------------------------|----------------------------------------------------------------|
|    | Management (IAM)                   |                                                                |
| 81 | Multi-Factor Authentication (MFA)  | Захист за допомогою кількох факторів автентифікації.           |
| 82 | Password Policy                    | Правила створення і використання паролів.                      |
| 83 | Network Segmentation               | Поділ мережі на ізольовані зони безпеки.                       |
| 84 | Zero Trust                         | Модель безпеки, що не довіряє жодному користувачу чи пристрою. |
| 85 | Cloud Security                     | Методи захисту даних у хмарних середовищах.                    |
| 86 | Container Security                 | Захист контейнеризованих додатків.                             |
| 87 | API Security                       | Захист інтерфейсів прикладного програмування.                  |
| 88 | Blockchain Security                | Забезпечення захисту децентралізованих систем.                 |
| 89 | Machine Learning in Security       | Використання ШІ для аналізу та детектування загроз.            |
| 90 | Anomaly Detection                  | Виявлення відхилень у поведінці системи.                       |
| 91 | Behavioral Analysis                | Аналіз поведінкових шаблонів користувачів.                     |
| 92 | Threat Intelligence Platform (TIP) | Система збору та обробки даних про загрози.                    |
| 93 | Security Metrics                   | Кількісні показники ефективності безпеки.                      |
| 94 | Telemetry Data                     | Технічна інформація про стан системи, зібрана автоматично.     |
| 95 | False Positive                     | Хибне спрацювання системи безпеки.                             |

|     |                                 |                                                                |
|-----|---------------------------------|----------------------------------------------------------------|
| 96  | False Negative                  | Пропущена справжня загроза.                                    |
| 97  | Alert Fatigue                   | Втома аналітиків SOC через надмір кількості сповіщень.         |
| 98  | Correlation Rule                | Правило для зв'язування подій у SIEM.                          |
| 99  | Whitelisting / Blacklisting     | Дозволені / заборонені списки об'єктів або адрес.              |
| 100 | Cyber Resilience                | Здатність системи протистояти та відновлюватись після атак.    |
| 101 | Cyber Threat Intelligence (CTI) | Інформація про поточні та потенційні кіберзагрози.             |
| 102 | Digital Footprint               | Сукупність цифрових слідів користувача в Інтернеті.            |
| 103 | Threat Feed                     | Потік даних про загрози, що надходить із зовнішніх джерел.     |
| 104 | Malware Analysis                | Процес дослідження шкідливого ПЗ для визначення його функцій.  |
| 105 | Static Analysis                 | Аналіз файлів без виконання коду.                              |
| 106 | Dynamic Analysis                | Аналіз виконання програми в ізолюваному середовищі.            |
| 107 | Reverse Engineering             | Декомпіляція для вивчення структури програмного коду.          |
| 108 | Signature-Based Detection       | Виявлення загроз за відомими сигнатурами.                      |
| 109 | Heuristic Analysis              | Виявлення невідомих загроз на основі поведінки.                |
| 110 | Behavioral Detection            | Аналіз поведінки процесів для визначення аномалій.             |
| 111 | Threat Landscape                | Загальна картина сучасних кіберзагроз.                         |
| 112 | IOC Enrichment                  | Додавання контекстної інформації до індикаторів компрометації. |

|     |                         |                                                              |
|-----|-------------------------|--------------------------------------------------------------|
| 113 | IOC Correlation         | Зіставлення індикаторів для виявлення пов'язаних інцидентів. |
| 114 | Threat Feed Aggregator  | Система об'єднання кількох джерел даних про загрози.         |
| 115 | Threat Score            | Оцінка рівня небезпеки загрози.                              |
| 116 | IOC Normalization       | Приведення індикаторів до єдиного формату.                   |
| 117 | IOC Validation          | Перевірка достовірності індикаторів компрометації.           |
| 118 | IOC Lifecycle           | Етапи створення, перевірки, поширення та видалення ІоС.      |
| 119 | Threat Attribution      | Визначення джерела або групи, відповідальної за атаку.       |
| 120 | Adversary Emulation     | Імітація дій зловмисника для перевірки захисту.              |
| 121 | Cyber Kill Chain        | Послідовність етапів здійснення кібератаки.                  |
| 122 | Data Correlation        | Виявлення зв'язків між різними джерелами даних.              |
| 123 | Event Normalization     | Уніфікація подій безпеки для аналізу в SIEM.                 |
| 124 | Log Management          | Збір, зберігання та аналіз системних журналів.               |
| 125 | Security Monitoring     | Безперервне спостереження за станом системи безпеки.         |
| 126 | Threat Detection Engine | Механізм, що виконує аналіз подій для виявлення загроз.      |
| 127 | Data Enrichment         | Додавання метаданих для покращення аналітики.                |
| 128 | Security Automation     | Автоматизація процесів моніторингу та реагування.            |
| 129 | Data Correlation Engine | Підсистема SIEM, що зіставляє події безпеки.                 |

|     |                                      |                                                        |
|-----|--------------------------------------|--------------------------------------------------------|
| 130 | Alert Management                     | Керування сповіщеннями про загрози.                    |
| 131 | Incident Prioritization              | Класифікація інцидентів за рівнем критичності.         |
| 132 | Triage                               | Початкове сортування інцидентів безпеки.               |
| 133 | Threat Actor Profiling               | Створення профілю зловмисника на основі його дій.      |
| 134 | Threat Campaign                      | Серія атак, спрямованих на спільну мету.               |
| 135 | Threat Vector                        | Шлях або метод реалізації атаки.                       |
| 136 | Attack Surface                       | Сукупність усіх можливих точок входу в систему.        |
| 137 | Attack Path                          | Ланцюг кроків, який використовує зловмисник.           |
| 138 | Breach Simulation                    | Тестування захисту шляхом симуляції інцидентів.        |
| 139 | Threat Mapping                       | Відображення загроз у контексті активів або процесів.  |
| 140 | MITRE D3FEND                         | База знань про техніки кіберзахисту.                   |
| 141 | Cyber Threat Sharing                 | Обмін інформацією між організаціями про загрози.       |
| 142 | Threat Intelligence Sharing Platform | Система обміну СТІ-даними (наприклад, MISP).           |
| 143 | SOC Automation                       | Використання скриптів і ШІ для автоматизації SOC.      |
| 144 | Incident Containment                 | Обмеження поширення інциденту.                         |
| 145 | Root Cause Analysis (RCA)            | Виявлення першопричини кіберінциденту.                 |
| 146 | Post-Incident Review                 | Аналіз ефективності реагування після інциденту.        |
| 147 | Security Posture                     | Поточний рівень кіберзахисту організації.              |
| 148 | Security Hardening                   | Посилення конфігурації системи для підвищення захисту. |

|     |                                            |                                                                                                                                  |
|-----|--------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------|
| 149 | Vulnerability Assessment                   | Перевірка системи на наявність уразливостей.                                                                                     |
| 150 | CVE (Common Vulnerabilities and Exposures) | CVE (Common Vulnerabilities and Exposures)                                                                                       |
| 151 | Vulnerability Management                   | Безперервний процес виявлення, оцінки та усунення уразливостей.                                                                  |
| 152 | Patch Tuesday                              | День регулярного випуску оновлень безпеки від Microsoft.                                                                         |
| 153 | Threat Mitigation                          | Заходи для зниження впливу кіберзагроз.                                                                                          |
| 154 | Threat Modeling Framework                  | Методика системного аналізу загроз (наприклад, STRIDE, PASTA).                                                                   |
| 155 | STRIDE                                     | Модель класифікації загроз: Spoofing, Tampering, Repudiation, Information disclosure, Denial of service, Elevation of privilege. |
| 156 | DREAD                                      | Модель оцінки ризиків за критеріями шкоди, відтворюваності, експлуатації тощо.                                                   |
| 157 | Threat Intelligence Lifecycle              | Етапи збору, аналізу та розповсюдження інформації про загрози.                                                                   |
| 158 | Cyber Threat Framework (CTF)               | Модель опису етапів дій супротивника в кіберпросторі.                                                                            |
| 159 | IOC Pivoting                               | Пошук пов'язаних ІоС шляхом аналітичних зв'язків.                                                                                |
| 160 | Indicator Scoring                          | Оцінювання надійності та актуальності індикаторів компрометації.                                                                 |
| 161 | Threat Confidence Level                    | Рівень довіри до достовірності даних про загрозу.                                                                                |
| 162 | IOC Feed Integration                       | Автоматичне підключення зовнішніх ІоС-джерел у SIEM/TIP.                                                                         |

|     |                            |                                                                           |
|-----|----------------------------|---------------------------------------------------------------------------|
| 163 | Data Retention Policy      | Політика зберігання журналів подій і артефактів безпеки.                  |
| 164 | Threat Data Normalization  | Уніфікація даних із різних джерел СТІ.                                    |
| 165 | IOC Expiration             | Механізм автоматичного видалення застарілих індикаторів.                  |
| 166 | IOC False Flag             | Навмисне створення фальшивих індикаторів для введення аналітиків в оману. |
| 167 | Threat Intelligence Report | Документ із результатами аналізу кібератак та тенденцій загроз.           |
| 168 | Incident Escalation        | Передача інциденту на вищий рівень підтримки або керівництву.             |
| 169 | Threat Correlation Matrix  | Таблиця зв'язків між різними джерелами даних загроз.                      |
| 170 | IOC Repository             | Централізоване сховище індикаторів компрометації.                         |
| 171 | IOC Extraction             | Автоматичне вилучення ІоС із текстових або технічних звітів.              |
| 172 | IOC Classification         | Категоризація індикаторів за типом (IP, домен, URL, файл, хеш тощо).      |
| 173 | Threat Simulation          | Перевірка стійкості системи до атак за допомогою моделювання.             |
| 174 | Adversary Simulation       | Відтворення поведінки відомих груп зловмисників.                          |
| 175 | Cyber Range                | Навчальне середовище для практики кіберзахисту.                           |
| 176 | IOC Aggregation            | Об'єднання кількох джерел ІоС у єдину базу.                               |
| 177 | Threat Visibility          | Рівень прозорості та обізнаності про стан безпеки в мережі.               |

|     |                        |                                                                       |
|-----|------------------------|-----------------------------------------------------------------------|
| 178 | Threat Telemetry       | Дані, що надходять із систем моніторингу для аналітики загроз.        |
| 179 | Anomaly Scoring        | Оцінка ймовірності відхилення від нормальної поведінки.               |
| 180 | Data Lake              | Централізоване сховище неструктурованих даних для аналізу.            |
| 181 | Security Data Fabric   | Архітектура інтегрованої обробки даних безпеки з різних систем.       |
| 182 | Threat Analytics       | Використання аналітичних методів для виявлення та прогнозування атак. |
| 183 | Cyber Threat Dashboard | Візуалізація поточного стану загроз у системі моніторингу.            |
| 184 | IOC Intelligence Score | Числова оцінка рівня ризику індикатора компрометації.                 |
| 185 | IOC False Correlation  | Помилковий зв'язок між непов'язаними ІоС.                             |
| 186 | IOC Confidence Score   | Оцінка достовірності ІоС за джерелами.                                |
| 187 | Threat Dissemination   | Поширення інформації про загрози серед користувачів і партнерів.      |
| 188 | Threat Validation      | Перевірка точності даних про кіберзагрози.                            |
| 189 | Threat Deconfliction   | Усунення дублювання або конфліктів між даними про загрози.            |
| 190 | Threat Indicator Decay | Зменшення актуальності індикаторів із часом.                          |
| 191 | IOC Mapping            | Відображення індикаторів на відомі техніки MITRE ATT&CK.              |
| 192 | Threat Signature       | Унікальний шаблон або опис поведінки загрози.                         |

|     |                            |                                                                       |
|-----|----------------------------|-----------------------------------------------------------------------|
| 193 | IOC Propagation            | Автоматичне розповсюдження ІоС між системами безпеки.                 |
| 194 | IOC Reputation             | Рейтинг надійності ІоС на основі минулих інцидентів.                  |
| 195 | Threat Database            | Централізована база знань про типові загрози.                         |
| 196 | IOC Tagging                | Маркування індикаторів за категоріями або рівнями загрози.            |
| 197 | Threat Feed API            | Інтерфейс обміну даними між платформами кіберрозвідки.                |
| 198 | IOC Indicator Type         | Класифікація типу артефакту (URL, домен, IP, хеш, сертифікат).        |
| 199 | IOC Extraction Tool        | Програмний засіб для автоматичного вилучення ІоС із тексту чи звітів. |
| 200 | IOC Deduplication          | Процес видалення дубльованих індикаторів у базі даних.                |
| 201 | Threat Scorecard           | Короткий звіт із кількісними показниками загроз.                      |
| 202 | IOC Timeline               | Хронологія появи, використання та блокування індикатора.              |
| 203 | Cyber Threat Correlation   | Аналіз зв'язків між кількома подіями та індикаторами.                 |
| 204 | Threat Attribution Report  | Аналітичний документ про походження атаки.                            |
| 205 | Threat Research            | Систематичне дослідження нових методів атак і захисту.                |
| 206 | Cyber Threat Observatory   | Аналітичний центр збору даних про загрози.                            |
| 207 | IOC Expiry Policy          | Політика автоматичного видалення неактуальних ІоС.                    |
| 208 | Indicator Validation Score | Показник достовірності ІоС.                                           |

|     |                           |                                                           |
|-----|---------------------------|-----------------------------------------------------------|
| 209 | IOC Contextualization     | Додавання контексту (час, джерело, тип загрози) до IoC.   |
| 210 | IOC Versioning            | Ведення історії змін індикаторів у базі.                  |
| 211 | Threat Taxonomy           | Система класифікації кіберзагроз.                         |
| 212 | Threat Ontology           | Формальна модель взаємозв'язків між об'єктами загроз.     |
| 213 | IOC Feed Normalization    | Приведення потоків даних до єдиного формату STIX/TAXII.   |
| 214 | IOC Confidence Weight     | Вага довіри до конкретного джерела IoC.                   |
| 215 | Threat Vector Mapping     | Зіставлення шляхів атаки з конкретними активами.          |
| 216 | IOC Sharing Policy        | Політика обміну індикаторами між організаціями.           |
| 217 | Threat Source Reliability | Рейтинг надійності джерела даних про загрози.             |
| 218 | IOC Aggregator            | Сервіс збору IoC із різних систем.                        |
| 219 | Threat Hunting Query      | Запит для пошуку підозрілої активності у логах.           |
| 220 | Threat Event Correlation  | Виявлення залежностей між кількома інцидентами.           |
| 221 | Security Orchestration    | Координація інструментів безпеки через автоматизацію.     |
| 222 | Threat Feed Subscription  | Підписка на регулярне оновлення IoC-даних.                |
| 223 | IOC Confidence Threshold  | Поріг достовірності, нижче якого IoC не використовується. |
| 224 | IOC Correlation Graph     | Візуальне відображення зв'язків між індикаторами.         |

|     |                            |                                                     |
|-----|----------------------------|-----------------------------------------------------|
| 225 | Threat Cluster             | Група схожих інцидентів або атак.                   |
| 226 | IOC Aging Policy           | Політика зниження пріоритету старих IoC.            |
| 227 | IOC Provenance             | Джерело походження індикатора компрометації.        |
| 228 | Threat Context             | Сукупність метаданих, що пояснюють природу загрози. |
| 229 | IOC Risk Score             | Числова оцінка ризику, який несе індикатор.         |
| 230 | IOC Distribution           | Механізм поширення IoC між системами.               |
| 231 | IOC Deduplication Engine   | Компонент, що автоматично видаляє дублікати IoC.    |
| 232 | Threat Validation Pipeline | Процес перевірки і фільтрації даних СТІ.            |
| 233 | Threat Feed Parser         | Інструмент для розбору даних із потоків загроз.     |
| 234 | Cyber Threat Enrichment    | Додавання додаткових атрибутів до даних загроз.     |
| 235 | IOC False Attribution      | Хибне приписування атаки невинному джерелу.         |
| 236 | Threat Correlation Engine  | Механізм аналізу взаємозв'язків між подіями.        |
| 237 | Threat Contextual Data     | Інформація, що описує середовище атаки.             |
| 238 | Threat Risk Assessment     | Аналіз рівня ризику кіберзагроз.                    |
| 239 | IOC Confidence Metric      | Показник надійності IoC.                            |
| 240 | Threat Scoring Algorithm   | Математичний метод для розрахунку рівня загрози.    |

|     |                                   |                                                    |
|-----|-----------------------------------|----------------------------------------------------|
| 241 | Threat Response Plan              | Документ із заходами реагування на інциденти.      |
| 242 | Threat Recovery Plan              | План відновлення після атаки.                      |
| 243 | IOC Lifecycle Management          | Управління повним життєвим циклом ІоС.             |
| 244 | Threat Data Warehouse             | Централізоване сховище даних кіберрозвідки.        |
| 245 | Threat Data Fusion                | Об'єднання даних із різних джерел для аналітики.   |
| 246 | Threat Enrichment API             | Інтерфейс для автоматичного додавання контексту.   |
| 247 | IOC Automation Script             | Скрипт для обробки або імпорту індикаторів.        |
| 248 | Threat Metric                     | Кількісний показник стану безпеки.                 |
| 249 | IOC Quality Control               | Перевірка точності й актуальності індикаторів.     |
| 250 | Threat Intelligence Fusion Center | Центр об'єднання даних із різних СТІ-джерел.       |
| 251 | IOC Management Platform           | Платформа для зберігання, обміну та перевірки ІоС. |
| 252 | Cyber Threat Lifecycle            | Повний цикл розвитку та усунення кіберзагрози.     |
| 253 | IOC Registry                      | Реєстр індикаторів компрометації.                  |
| 254 | IOC Threat Mapping                | Прив'язка індикаторів до відповідних атак MITRE.   |
| 255 | IOC Visualization                 | Графічне представлення зв'язків ІоС.               |

|     |                            |                                                                   |
|-----|----------------------------|-------------------------------------------------------------------|
| 256 | Threat Knowledge Graph     | Семантична модель зв'язків між загрозами.                         |
| 257 | IOC Repository Index       | Каталог індикаторів за типом і часом.                             |
| 258 | Threat Pattern Recognition | Розпізнавання типових шаблонів атак.                              |
| 259 | IOC Data Integrity         | Перевірка цілісності бази індикаторів.                            |
| 260 | Threat Timeline            | Хронологічна послідовність подій загрози.                         |
| 261 | Threat Prediction Model    | Модель прогнозування ймовірності атак.                            |
| 262 | Threat Trend Analysis      | Дослідження тенденцій у розвитку кіберзагроз.                     |
| 263 | Threat Forecasting         | Прогнозування майбутніх атак на основі даних.                     |
| 264 | IOC Sharing Framework      | Стандартизована модель обміну IoC між системами.                  |
| 265 | Threat Knowledge Base      | База знань про відомі техніки, індикатори та акторів.             |
| 266 | Threat Deception           | Використання пасток і фальшивих цілей для виявлення зловмисників. |
| 267 | Honeypot                   | Система-приманка для виявлення атак.                              |
| 268 | Honeytoken                 | Фіктивний об'єкт для спостереження за зловмисною активністю.      |
| 269 | Decoy Network              | Імітована мережа для виявлення вторгнень.                         |
| 270 | Deception Technology       | Технології створення обманних середовищ у безпеці.                |
| 271 | Threat Signal              | Ознака, яка свідчить про наявність кібератаки.                    |

|     |                                        |                                                        |
|-----|----------------------------------------|--------------------------------------------------------|
| 272 | IOC Expiration Date                    | Термін актуальності індикатора компрометації.          |
| 273 | Threat Collaboration                   | Спільна робота команд безпеки над аналізом інцидентів. |
| 274 | Threat Knowledge Sharing               | Обмін знаннями про атаки та методи захисту.            |
| 275 | Threat Fusion Platform                 | Система інтеграції даних про загрози.                  |
| 276 | Threat Maturity Model                  | Рівні зрілості процесів кіберрозвідки в організації.   |
| 277 | Threat Detection Framework             | Концептуальна модель процесу виявлення атак.           |
| 278 | IOC Alert Correlation                  | Зіставлення IoC із подіями для формування сповіщення.  |
| 279 | Threat Governance                      | Управління процесами безпеки на рівні політик.         |
| 280 | IOC Confidence Analysis                | Статистичний аналіз достовірності IoC.                 |
| 281 | Threat KPI (Key Performance Indicator) | Ключові показники ефективності захисту від загроз.     |
| 282 | Cyber Threat Taxonomy                  | Ієрархічна структура класифікації типів загроз.        |
| 283 | IOC Data Quality                       | Оцінка повноти та точності індикаторів.                |
| 284 | Threat Feed Validation                 | Перевірка якості потоків загроз.                       |
| 285 | Threat Detection Model                 | Алгоритм або правило для виявлення атак.               |

|     |                             |                                                         |
|-----|-----------------------------|---------------------------------------------------------|
| 286 | Threat Analytics Pipeline   | Послідовність етапів обробки даних загроз.              |
| 287 | IOC Feedback Loop           | Механізм зворотного зв'язку для вдосконалення бази IoC. |
| 288 | Threat Mitigation Strategy  | Комплекс заходів для зменшення наслідків атак.          |
| 289 | Cyber Resilience Framework  | Модель забезпечення стійкості організації до атак.      |
| 290 | Threat Recovery Testing     | Перевірка ефективності плану відновлення після атаки.   |
| 291 | IOC Prioritization          | Визначення важливості індикаторів для реагування.       |
| 292 | Threat Readiness Assessment | Оцінка готовності організації до реагування на загрози. |
| 293 | Threat Response Workflow    | Послідовність дій реагування в SOC.                     |
| 294 | Threat Learning System      | Самонавчальна система з виявлення атак.                 |
| 295 | Threat Signature Update     | Оновлення сигнатур у системах виявлення.                |
| 296 | IOC Dataset                 | Набір структурованих індикаторів для досліджень.        |
| 297 | Threat Vector Analysis      | Аналіз шляхів і методів реалізації атак.                |
| 298 | Threat Resilience Testing   | Перевірка здатності системи протистояти загрозам.       |
| 299 | Threat Data Governance      | Політика управління даними кіберрозвідки.               |
| 300 | Threat Response Playbook    | Набір інструкцій для реагування на конкретні типи атак. |

|     |                             |                                                                                     |
|-----|-----------------------------|-------------------------------------------------------------------------------------|
| 301 | Threat Response Playbook    | Набір інструкцій для реагування на конкретні типи атак.                             |
| 302 | Incident Containment        | Процес ізоляції інциденту безпеки для запобігання його поширенню.                   |
| 303 | Log Correlation             | Аналіз та зв'язування подій із різних журналів для виявлення закономірностей.       |
| 304 | Threat Simulation           | Моделювання кіберзагроз для перевірки готовності системи безпеки.                   |
| 305 | Endpoint Isolation          | Відключення зараженого кінцевого пристрою від мережі для обмеження збитків.         |
| 306 | Alert Fatigue               | Стан перевантаження аналітика великою кількістю сповіщень безпеки.                  |
| 307 | Threat Mitigation           | Дії, спрямовані на зниження або усунення впливу загрози.                            |
| 308 | Malware Reverse Engineering | Аналіз шкідливого ПЗ для розуміння його функцій і механізмів.                       |
| 309 | Network Baseline            | Нормальна поведінка мережі, що використовується для виявлення аномалій.             |
| 310 | Security Hardening          | Посилення безпеки системи шляхом мінімізації вразливостей.                          |
| 311 | Data Exfiltration           | Несанкціоноване виведення конфіденційних даних із системи.                          |
| 312 | Command and Control (C2)    | Канал управління, який використовують зловмисники для контролю заражених пристроїв. |
| 313 | Fileless Malware            | Шкідливе ПЗ, що діє без запису файлів на диск, використовуючи пам'ять системи.      |
| 314 | Privilege Escalation        | Отримання зловмисником вищих прав доступу, ніж передбачено.                         |
| 315 | Threat Actor                | Особа або група, що здійснює кіберзлочинні дії.                                     |

|     |                             |                                                                              |
|-----|-----------------------------|------------------------------------------------------------------------------|
| 316 | Cyber Threat Hunting        | Активний пошук ознак компрометації в системах без очікування сповіщень.      |
| 317 | Threat Intelligence Sharing | Обмін інформацією про загрози між організаціями.                             |
| 318 | Security Posture            | Поточний рівень захищеності організації від кіберзагроз.                     |
| 319 | Data Integrity Check        | Перевірка цілісності даних для виявлення несанкціонованих змін.              |
| 320 | Forensic Imaging            | Створення точної копії цифрового носія для аналізу.                          |
| 321 | Threat Attribution          | Визначення джерела або відповідальної сторони за кібератаку.                 |
| 322 | Cyber Range                 | Тренувальне середовище для практичного відпрацювання сценаріїв кіберзахисту. |
| 323 | Insider Threat Detection    | Виявлення шкідливої діяльності всередині організації.                        |
| 324 | Threat Landscape            | Сукупність усіх поточних кіберзагроз і тенденцій.                            |
| 325 | Patch Management            | Система управління оновленнями безпеки для усунення вразливостей.            |
| 326 | Threat Scoring              | Кількісна оцінка ризику або серйозності загрози.                             |
| 327 | Risk Register               | Документ, у якому фіксуються всі відомі ризики та заходи для їх зменшення.   |
| 328 | Behavior Analytics          | Аналіз поведінкових даних користувачів для виявлення аномалій.               |
| 329 | Attack Chain                | Послідовність дій, що здійснює зловмисник під час атаки.                     |
| 330 | Threat Surface              | Сукупність усіх можливих точок для атак у системі.                           |
| 331 | Identity Federation         | Об'єднання систем ідентифікації кількох організацій.                         |

|     |                                     |                                                                          |
|-----|-------------------------------------|--------------------------------------------------------------------------|
| 332 | Cloud Access Security Broker (CASB) | Посередник між користувачами та хмарними сервісами для контролю безпеки. |
| 333 | Encryption Key Management           | Керування життєвим циклом криптографічних ключів.                        |
| 334 | Phishing Simulation                 | Тренування співробітників через імітацію фішингових атак.                |
| 335 | Threat Feed                         | Потік актуальної інформації про нові загрози.                            |
| 336 | Data Loss Prevention (DLP)          | Технології для запобігання витоку конфіденційних даних.                  |
| 337 | Threat Lifecycle                    | Повний цикл існування загрози — від створення до нейтралізації.          |
| 338 | Digital Fingerprinting              | Ідентифікація систем або користувачів за унікальними цифровими ознаками. |
| 339 | Zero-Day Vulnerability              | Уразливість, яка ще не відома виробникові або громадськості.             |
| 340 | Cyber Threat Modeling               | Моделювання потенційних атак для оцінки ризиків.                         |
| 341 | Security Orchestration              | Автоматизація та координація процесів реагування на інциденти.           |
| 342 | Network Traffic Analysis            | Аналіз мережевого трафіку для виявлення аномалій і атак.                 |
| 343 | Credential Dumping                  | Викрадення облікових даних із пам'яті системи.                           |
| 344 | Threat Indicator Enrichment         | Розширення інформації про індикатор компрометації з додаткових джерел.   |
| 345 | Application Whitelisting            | Дозвіл на виконання лише перевірених додатків у системі.                 |

|     |                                       |                                                                           |
|-----|---------------------------------------|---------------------------------------------------------------------------|
| 346 | Lateral Movement                      | Переміщення злоумисника всередині мережі після початкового вторгнення.    |
| 347 | Threat Context                        | Додаткова інформація, що допомагає зрозуміти значення індикатора загрози. |
| 348 | Host-based Detection                  | Виявлення загроз на рівні окремого пристрою або сервера.                  |
| 349 | Network Segmentation                  | Поділ мережі на ізольовані зони для підвищення безпеки.                   |
| 350 | Intrusion Prevention System (IPS)     | Система, яка блокує підозрілі дії у мережі в режимі реального часу.       |
| 351 | Security Audit Trail                  | Хронологічний запис усіх подій безпеки в системі.                         |
| 352 | Sandboxing                            | Ізоляція виконання потенційно небезпечного коду для аналізу.              |
| 353 | Security Benchmark                    | Еталонні показники безпеки, що використовуються для оцінки стану системи. |
| 354 | Threat Response Automation            | Автоматичне виконання дій із реагування на виявлені загрози.              |
| 355 | Deception Technology                  | Використання пасток і фіктивних ресурсів для виявлення злоумисників.      |
| 356 | Endpoint Detection and Response (EDR) | Інструмент моніторингу й реагування на атаки на кінцевих пристроях.       |
| 357 | Malware Sandbox                       | Середовище для безпечного аналізу підозрілих файлів.                      |
| 358 | Threat Analytics Platform             | Система для збору, кореляції та аналізу даних про загрози.                |
| 359 | Incident Coordination                 | Організація взаємодії команд під час реагування на інцидент.              |

|     |                                    |                                                                      |
|-----|------------------------------------|----------------------------------------------------------------------|
| 360 | Threat Intelligence Platform (TIP) | Платформа для централізованого управління даними про загрози.        |
| 361 | Network Forensics                  | Розслідування інцидентів шляхом аналізу мережових даних.             |
| 362 | Cyber Deception                    | Стратегія введення зловмисників в оману для виявлення їхніх дій.     |
| 363 | Breach Notification                | Офіційне повідомлення про інцидент витоку даних.                     |
| 364 | Threat Correlation Engine          | Модуль, що поєднує дані з різних джерел для виявлення складних атак. |
| 365 | Anomaly-based Detection            | Виявлення загроз шляхом пошуку відхилень від норми.                  |
| 366 | Threat Exposure                    | Рівень схильності системи до певних типів атак.                      |
| 367 | Log Retention Policy               | Політика зберігання журналів подій для аудиту безпеки.               |
| 368 | Incident Retrospective             | Аналіз інциденту після його завершення для вдосконалення процедур.   |
| 369 | Security Governance Framework      | Структура управління кібербезпекою в організації.                    |
| 370 | Threat Prioritization              | Визначення пріоритетів реагування залежно від серйозності загроз.    |
| 371 | Network Intrusion Analysis         | Вивчення вторгнень у мережу з метою ідентифікації джерела.           |
| 372 | Threat Actor Profiling             | Побудова профілю зловмисника за його поведінковими характеристиками. |
| 373 | Host Forensics                     | Аналіз окремого пристрою для збору доказів компрометації.            |

|     |                            |                                                                          |
|-----|----------------------------|--------------------------------------------------------------------------|
| 374 | Automated Threat Detection | Використання алгоритмів для автоматичного виявлення загроз.              |
| 375 | Malware Persistence        | Механізми, за допомогою яких шкідливе ПЗ зберігає присутність у системі. |
| 376 | Threat Reconnaissance      | Початковий етап збору даних про ціль перед атакою.                       |
| 377 | Breach Containment         | Дії для обмеження наслідків витоку даних.                                |
| 378 | Vulnerability Assessment   | Процес пошуку та оцінки вразливостей у системі.                          |
| 379 | Threat Pattern Recognition | Виявлення повторюваних шаблонів атак.                                    |
| 380 | Exploit Kit                | Набір інструментів для автоматизації використання вразливостей.          |
| 381 | Threat Triage              | Початкова оцінка й класифікація загроз.                                  |
| 382 | Security Awareness Program | Освітня програма для підвищення обізнаності персоналу з безпеки.         |
| 383 | Threat Emulation           | Імітація дій зловмисника для тестування захисту.                         |
| 384 | Access Control List (ACL)  | Список дозволів для управління доступом до ресурсів.                     |
| 385 | Threat Heatmap             | Візуальне представлення розподілу кіберзагроз за пріоритетами.           |
| 386 | Post-Incident Review       | Аналіз післяінцидентних дій для покращення реакції.                      |
| 387 | Threat Classification      | Систематизація загроз за типами, джерелами чи методами.                  |
| 388 | Endpoint Telemetry         | Збір телеметричних даних із пристроїв для моніторингу стану безпеки.     |

|     |                                  |                                                                         |
|-----|----------------------------------|-------------------------------------------------------------------------|
| 389 | Threat Data Normalization        | Стандартизація даних про загрози для подальшого аналізу.                |
| 390 | Adaptive Security Architecture   | Гнучка архітектура безпеки, що динамічно реагує на загрози.             |
| 391 | Attack Path Analysis             | Аналіз шляхів, якими зловмисник може досягти цілі.                      |
| 392 | Threat Readiness Assessment      | Оцінка готовності організації до реагування на атаки.                   |
| 393 | Data Breach Simulation           | Імітація витоку даних для перевірки процедур реагування.                |
| 394 | Threat Source Identification     | Визначення походження конкретної загрози.                               |
| 395 | Security Log Aggregation         | Збір журналів подій із різних систем у єдине сховище.                   |
| 396 | Threat Correlation Matrix        | Матриця взаємозв'язків між різними типами загроз.                       |
| 397 | Advanced Persistent Threat (APT) | Тривала цілеспрямована атака, що діє приховано протягом тривалого часу. |
| 398 | Threat Mitigation Plan           | План заходів щодо зниження впливу загроз.                               |
| 399 | Incident Simulation Exercise     | Навчальна вправа з відпрацювання сценаріїв кіберінцидентів.             |
| 400 | Threat Detection Workflow        | Послідовність дій для виявлення та аналізу кіберзагроз.                 |
| 401 | Threat Validation                | Перевірка достовірності та актуальності отриманих індикаторів загроз.   |

|     |                              |                                                                     |
|-----|------------------------------|---------------------------------------------------------------------|
| 402 | Security Event Correlation   | Зіставлення подій безпеки з різних джерел для виявлення інцидентів. |
| 403 | Threat Detection Rule        | Набір умов, за якими система ідентифікує потенційну атаку.          |
| 404 | Cyber Kill Chain             | Модель етапів кібератаки від розвідки до дій із компрометацією.     |
| 405 | Threat Alert Prioritization  | Визначення рівня важливості повідомлень про загрози.                |
| 406 | Data Breach Response         | Дії з ліквідації наслідків витоку даних.                            |
| 407 | Threat Trend Analysis        | Аналіз тенденцій розвитку кіберзагроз у часі.                       |
| 408 | Malware Variant              | Модифікація відомого шкідливого програмного забезпечення.           |
| 409 | Threat Aggregation           | Об'єднання схожих загроз у групи для спрощення аналізу.             |
| 410 | Insider Risk Management      | Управління ризиками, пов'язаними з внутрішніми користувачами.       |
| 411 | Threat Detection Coverage    | Рівень охоплення системи моніторингу потенційних атак.              |
| 412 | Botnet Command Server        | Сервер, що керує мережею заражених пристроїв.                       |
| 413 | Threat Landscape Analysis    | Системний огляд поточного стану кіберзагроз.                        |
| 414 | Vulnerability Prioritization | Визначення черговості усунення знайдених вразливостей.              |
| 415 | Threat Hunting Hypothesis    | Припущення, яке тестується під час активного пошуку загроз.         |

|     |                              |                                                                      |
|-----|------------------------------|----------------------------------------------------------------------|
| 416 | Security Incident Lifecycle  | Етапи розвитку інциденту від виявлення до відновлення.               |
| 417 | Threat Modeling Framework    | Методологія для ідентифікації, аналізу та управління загрозами.      |
| 418 | Data Integrity Violation     | Порушення цілісності даних через несанкціоновані зміни.              |
| 419 | Threat Intel Report          | Аналітичний звіт про поточні кіберзагрози.                           |
| 420 | Cyber Defense Strategy       | Сукупність заходів для захисту інформаційних ресурсів.               |
| 421 | Threat Propagation           | Поширення загрози через мережу або систему.                          |
| 422 | Network Telemetry            | Збір технічних даних про активність у мережі.                        |
| 423 | Threat Vector                | Шлях, через який загроза може проникнути у систему.                  |
| 424 | Security Playbook Automation | Автоматизація дій із реагування за заздалегідь визначеним сценарієм. |
| 425 | Threat Cluster               | Група загроз, пов'язаних спільними ознаками або індикаторами.        |
| 426 | Threat Attribution Report    | Документ, що описує джерело та мотиви кібератаки.                    |
| 427 | Security Data Lake           | Централізоване сховище даних безпеки для аналітики.                  |
| 428 | Threat Reduction Strategy    | План заходів для зменшення кількості та впливу загроз.               |
| 429 | Threat Contextualization     | Додавання додаткових даних для розуміння суті загрози.               |
| 430 | Threat Actor Tactics         | Тактики, що застосовуються зловмисниками під час атак.               |

|     |                                    |                                                               |
|-----|------------------------------------|---------------------------------------------------------------|
| 431 | Cloud Threat Detection             | Виявлення загроз у хмарних середовищах.                       |
| 432 | Threat Lifecycle Management        | Управління всіма етапами життєвого циклу загрози.             |
| 433 | Threat Escalation                  | Підвищення рівня серйозності атаки або її наслідків.          |
| 434 | Security Event Management          | Централізований збір та аналіз подій безпеки.                 |
| 435 | Threat Intelligence Feed           | Потік структурованих даних про загрози.                       |
| 436 | Threat Hunting Maturity Model      | Модель оцінки рівня розвитку процесів пошуку загроз.          |
| 437 | Threat Propagation Vector          | Механізм поширення шкідливого коду в мережі.                  |
| 438 | Endpoint Protection Platform (EPP) | Комплексне рішення для захисту кінцевих пристроїв.            |
| 439 | Threat Validation Process          | Процес перевірки достовірності виявленої загрози.             |
| 440 | Threat Actor Motivation            | Причини, що спонукають до здійснення кібератак.               |
| 441 | Threat Remediation                 | Усунення або нейтралізація наслідків загрози.                 |
| 442 | Threat Pattern Analysis            | Аналіз закономірностей у діях зловмисників.                   |
| 443 | Security Automation Workflow       | Послідовність автоматизованих дій із виявлення та реагування. |

|     |                                 |                                                                    |
|-----|---------------------------------|--------------------------------------------------------------------|
| 444 | Threat Event Classification     | Категоризація подій, пов'язаних із кіберзагрозами.                 |
| 445 | Cyber Threat Framework          | Структура для опису й аналізу загроз.                              |
| 446 | Threat Containment Strategy     | Стратегія обмеження поширення загроз.                              |
| 447 | Threat Validation Framework     | Методологія перевірки достовірності інформації про загрози.        |
| 448 | Threat Response Orchestration   | Координація дій команд під час реагування на інциденти.            |
| 449 | Threat Remediation Workflow     | Послідовність етапів усунення наслідків атак.                      |
| 450 | Threat Visualization Dashboard  | Інтерактивна панель для відображення інформації про загрози.       |
| 451 | Threat Intelligence Correlation | Поєднання різних джерел даних для отримання повної картини загроз. |
| 452 | Threat Lifecycle Analysis       | Оцінка розвитку загрози від виявлення до ліквідації.               |
| 453 | Threat Signature Database       | База сигнатур відомих шкідливих програм.                           |
| 454 | Threat Risk Scoring             | Присвоєння загрозам числового показника ризику.                    |
| 455 | Threat Scenario Planning        | Створення сценаріїв можливих атак для підготовки до реагування.    |

|     |                                 |                                                                |
|-----|---------------------------------|----------------------------------------------------------------|
| 456 | Threat Intel Correlation Engine | Модуль для зіставлення інформації з різних аналітичних джерел. |
| 457 | Security Monitoring Framework   | Архітектура для побудови систем моніторингу кібербезпеки.      |
| 458 | Threat Intelligence Repository  | Сховище даних про загрози для подальшої обробки.               |
| 459 | Threat Analysis Report          | Документ із результатами розслідування кіберінцидентів.        |
| 460 | Threat Data Enrichment          | Додавання контекстної інформації до даних про загрози.         |
| 461 | Threat Response Framework       | Узгоджена система заходів реагування на інциденти.             |
| 462 | Threat Detection Matrix         | Таблиця відповідності типів загроз і методів їх виявлення.     |
| 463 | Security Telemetry Analysis     | Аналіз технічних показників для виявлення кіберінцидентів.     |
| 464 | Threat Campaign Analysis        | Дослідження тривалих і багатоетапних атак.                     |
| 465 | Threat Actor Group              | Організоване угруповання, яке здійснює кібероперації.          |
| 466 | Threat Response Policy          | Визначені правила реагування на інциденти безпеки.             |
| 467 | Threat Hunting Procedure        | Покроковий процес пошуку прихованих атак.                      |

|     |                               |                                                                       |
|-----|-------------------------------|-----------------------------------------------------------------------|
| 468 | Threat Intelligence Lifecycle | Етапи створення, обробки та розповсюдження даних про загрози.         |
| 469 | Threat Landscape Report       | Огляд сучасних тенденцій у сфері кіберзагроз.                         |
| 470 | Threat Mitigation Workflow    | Алгоритм дій для зниження впливу кіберризиків.                        |
| 471 | Threat Data Fusion            | Об'єднання різних потоків інформації для підвищення точності аналізу. |
| 472 | Threat Risk Assessment        | Оцінка рівня ризику, пов'язаного з конкретною загрозою.               |
| 473 | Threat Remediation Plan       | План заходів для усунення впливу інциденту.                           |
| 474 | Threat Management Dashboard   | Панель моніторингу стану безпеки та активних загроз.                  |
| 475 | Threat Information Exchange   | Обмін даними про загрози між організаціями.                           |
| 476 | Threat Data Pipeline          | Процес збору, обробки та передачі даних безпеки.                      |
| 477 | Threat Response Coordination  | Взаємодія команд при ліквідації наслідків атаки.                      |
| 478 | Threat Correlation Dashboard  | Візуалізація взаємозв'язків між різними загрозами.                    |
| 479 | Threat Actor Attribution      | Ідентифікація групи або особи, відповідальної за атаку.               |
| 480 | Threat Knowledge Base         | База знань про відомі типи атак і техніки зловмисників.               |

|     |                            |                                                            |
|-----|----------------------------|------------------------------------------------------------|
| 481 | Threat Behavior Modeling   | Моделювання поведінки загроз для прогнозування атак.       |
| 482 | Threat Impact Analysis     | Аналіз наслідків кіберінциденту для організації.           |
| 483 | Threat Response Drill      | Практичне тренування команди реагування.                   |
| 484 | Threat Intelligence Portal | Вебплатформа для доступу до аналітики про загрози.         |
| 485 | Threat Remediation Report  | Документ із описом заходів після інциденту.                |
| 486 | Threat Detection Playbook  | Набір правил для автоматизації процесу виявлення атак.     |
| 487 | Threat Response Metrics    | Ключові показники ефективності реагування на інциденти.    |
| 488 | Threat Detection Engine    | Програмний компонент для ідентифікації кіберзагроз.        |
| 489 | Threat Response Center     | Координаційний центр із ліквідації кіберінцидентів.        |
| 490 | Threat Hunting Report      | Аналітичний документ про результати пошуку загроз.         |
| 491 | Threat Risk Matrix         | Матриця зіставлення рівнів ризику та типів загроз.         |
| 492 | Threat Management Process  | Сукупність дій для контролю всього життєвого циклу загроз. |
| 493 | Threat Analysis Framework  | Методологічна основа для дослідження загроз.               |

|     |                                |                                                                  |
|-----|--------------------------------|------------------------------------------------------------------|
| 494 | Threat Response Strategy       | Комплексна стратегія реагування на кібератаки.                   |
| 495 | Threat Intelligence API        | Інтерфейс для інтеграції даних про загрози в інші системи.       |
| 496 | Threat Correlation Framework   | Система логічних зв'язків між подіями кібербезпеки.              |
| 497 | Threat Response Checklist      | Перелік дій, які слід виконати при інциденті безпеки.            |
| 498 | Threat Visualization Framework | Інструмент для графічного представлення даних про загрози.       |
| 499 | Threat Analysis Platform       | Комплексна система збору та аналізу інформації про кіберзагрози. |
| 500 | Threat Intelligence Dashboard  | Інтерактивна панель моніторингу даних про поточні загрози.       |

Для нотаток:

Для нотаток:

Для нотаток:

Для нотаток: