

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ ДОНЕЦЬКИЙ
НАЦІОНАЛЬНИЙ УНІВЕРСИТЕТ ІМЕНІ ВАСИЛЯ СТУСА

ЛАСКАВЧУК МАКСИМ АНДРІЙОВИЧ

Допускається до захисту:
В.о. Завідувача кафедри
прикладної математики та
кібербезпеки, д-р філософії з
математики,

_____ Луценко А.В.
« ____ » _____ 20__ р.

**МЕТОДИ ПОПЕРЕДЖЕННЯ КОМП'ЮТЕРНИХ АТАК
ТИПУ *MAN IN THE MIDDLE*, ЩО ЗДІЙСНЮЮТЬСЯ З
ВИКОРИСТАННЯМ ГЕНЕРАТИВНОГО ШТУЧНОГО ІНТЕЛЕКТУ**

Спеціальність 125 Кібербезпека

Кваліфікаційна (бакалаврська) робота

Науковий керівник:
Загоруйко Л.В.,
к.т.н., доцент, доцент кафедри
прикладної математики та кібербезпеки

(підпис)

Оцінка : ____ / ____ / ____
(бали/за шкалою ЕКТС/за національною шкалою)

Голова ЕК: _____
(підпис)

АНОТАЦІЯ

Ласкавчук М.А. Методи попередження комп'ютерних атак типу Man In The Middle, що здійснюються з використанням генеративного штучного інтелекту. Спеціальність 125 «Кібербезпека». Донецький національний університет імені Василя Стуса, Вінниця 2025 рік.

У кваліфікаційній (бакалаврській) роботі досліджено комп'ютерну атаку типу «Man In The Middle», яка здійснюється за допомогою генеративного штучного інтелекту, а також проаналізовано сучасні методи її попередження. У межах дослідження реалізовано програмне рішення для виявлення атак типу «Man In The Middle», проведено його налаштування, тестування в реальному середовищі та здійснено порівняння ефективності з існуючим засобом захисту від подібних загроз.

Ключові слова: кібербезпека, комп'ютерні атаки, Man In The Middle, порушник, попередження атак, генеративний штучний інтелект.

53 с., 4 табл., 27 рис., 1 дод., 45 джерел.

ABSTRACT

Laskavchuk M.A. Methods of preventing computer attacks of the Man In The Middle type, carried out using generative artificial intelligence. Specialty 125 «Cybersecurity». Vasyl' Stus Donetsk National University, Vinnytsia, 2025.

The qualification (bachelor's) thesis investigates a computer attack of the «Man In The Middle» type, which is carried out using generative artificial intelligence, and analyzes modern methods of its prevention. As part of the study, a software solution for detecting «Man In The Middle» attacks was implemented, configured, tested in a real environment, and compared with the existing means of protection against such threats.

Keywords: cybersecurity, computer attacks, Man In The Middle, intruder, attack prevention, generative artificial intelligence.

53 p., 4 tables, 24 figures, 1 appendix, 45 sources.

ЗМІСТ

ПЕРЕЛІК СКОРОЧЕНЬ	4
ВСТУП.....	6
РОЗДІЛ 1 ЗАГАЛЬНІ ВІДОМОСТІ ПРО КОМП'ЮТЕРНІ АТАКИ З ВИКОРИСТАННЯМ ШТУЧНОГО ІНТЕЛЕКТУ	8
1.1. Відомості про комп'ютерні атаки, які здійснюються з використанням штучного інтелекту	8
1.2. Поняття комп'ютерної атаки Man In The Middle.....	15
РОЗДІЛ 2 МЕТОДИ ПОПЕРЕДЖЕННЯ ВТОРГНЕНЬ ТИПУ «MAN IN THE MIDDLE»	31
2.1. Методики запобігання комп'ютерним атакам типу «Man In The Middle». 31	
2.2 Попередження атак в бездротовій мережі Wi-Fi	39
2.3 Попередження атак в бездротовій мережі Bluetooth	45
РОЗДІЛ 3 ПРАКТИЧНА РЕАЛІЗАЦІЯ ПОПЕРЕДЖЕННЯ АТАК ТИПУ «MAN IN THE MIDDLE», ЩО ЗДІЙСНЮЮТЬСЯ З ВИКОРИСТАННЯМ ГЕНЕРАТИВНОГО ШТУЧНОГО ІНТЕЛЕКТУ	49
3.1 Налаштування існуючого засобу попередження атаки	49
3.2. Розробка програмного коду для попередження атак «MITM»	54
ВИСНОВКИ	58
СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ	59
ДОДАТКИ.....	64

ПЕРЕЛІК СКОРОЧЕНЬ

КЗЗ – комплекс засобів захисту.

ОС – операційна система.

ІНМ – штучні нейронні мережі.

ІІ – штучний інтелект.

AES (англ. Advanced Encryption Standard) – розширений стандарт шифрування.

AML (англ. Adversarial Machine Learning) – змагальне машинне навчання.

ARP (англ. Address Resolution Protocol) – протокол визначення адрес.

CCMP (англ. Counter Mode with Cipher Block Chaining Message Authentication Code Protocol) – протокол блочного шифрування з кодом автентичності повідомлення і режим зчеплення блоків і лічильника.

DFI (англ. Deep Flow Inspection) – глибока перевірка потоку.

DPI (англ. Deep Packet Inspection) – глибока перевірка пакетів.

DNS (англ. Domain Name System) – система доменних імен.

DNSSEC (англ. Domain Name System Security Extensions) – розширення безпеки системи доменних імен.

DoS (англ. Denial-of-Service Attack) – відмова в обслуговуванні.

EAPOL (англ. Extensible Authentication Protocol over LAN) – розширюваний протокол аутентифікації по локальній мережі.

GCMP (англ. Galois/Counter Mode Protocol) – протокол Галуа/лічильника.

GTK (англ. Group Transient Key) – груповий перехідний ключ.

IDPS (англ. Intrusion Detection and Prevention System) – система виявлення і запобігання вторгненням.

IGTK (англ. Integrity Group Temporary Key) – тимчасовий ключ групи цілісності.

IPS (англ. Intrusion Prevention System) – система запобігання вторгненням.

MIC (англ. Message Integrity Code) – код цілісності повідомлення.

MITM (англ. Man in the Middle) – людина посередині.

OWASP (англ. Open Web Application Security Project) – проєкт безпеки відкритих веб-додатків.

PMF (англ. Protected Management Frames) – захищені рамки управління.

PMK (англ. Pairwise Master Key) – парний майстер-ключ.

PSK (англ. Pre-Shared Key) – попередньо наданий ключ.

PTK (англ. Pairwise Transient Key) – парний перехідний ключ.

RSN (англ. Robust Security Network) – надійна мережа безпеки.

S-ARP (англ. Secure Address Resolution Protocol) – протокол безпечного визначення адрес.

SA (англ. Security Association) – асоціація безпеки.

SQL (англ. Structured Query Language) – мова структурованих запитів.

SSID (англ. Service Set Identifier) – ідентифікатор набору послуг.

SSL (англ. Secure Sockets Layer) – рівень захищених сокетів.

TKIP (англ. Temporal Key Integrity Protocol) – протокол цілісності тимчасового ключа.

TLS (англ. Transport Layer Security) – безпека транспортного рівня.

VPN (англ. Virtual Private Network) – віртуальна приватна мережа.

WLAN (англ. Wireless Local Area Network) – бездротова локальна мережа.

WPA (англ. Wi-Fi Protected Access) – захищений доступ через Wi-Fi.

WPA2 (англ. Wi-Fi Protected Access II) – захищений доступ через Wi-Fi II.

WPA3 (англ. Wi-Fi Protected Access III) – захищений доступ через Wi-Fi III.

XSS (англ. Cross Site Scripting) – міжсайтовий скриптинг.

ВСТУП

Актуальність роботи. З кожним роком спостерігається зростання інтенсивності комп'ютерних інцидентів. За інформацією Державної служби спеціального зв'язку та захисту інформації України, урядова команда реагування на комп'ютерні надзвичайні події CERT-UA, яка діє при Держспецзв'язку, в 2024 році опрацювала 4315 кіберінцидентів [1]. Це на 69,8% більше, ніж роком раніше, коли кіберзлочинці атакували український кіберпростір 2541 раз. Найчастіше зловмисники атакують місцеві органи влади, уряд та урядові організації, сектор безпеки та оборони, енергетичний сектор, комерційні організації, телекомунікації. Найпоширенішими типами інцидентів є розповсюдження шкідливого програмного забезпечення, фішинг, несанкціоноване підключення, компрометація облікових записів або систем. Метою зловмисників є викрадення чутливої інформації, а також знищення даних та інформаційних систем [1]. Використовуються різні методи проведення комп'ютерних атак, одним із таких методів є атака типу «Man In The Middle».

Порушники мають у своєму розпорядженні різноманітні інструменти, що дозволяють їм отримати несанкціонований доступ до даних. Крім того, порушники можуть мати в своєму розпорядженні значні ресурси для здійснення кібератак, що ураховано при класифікації порушників відповідно до класифікації [2]. Так, відповідно до п. 8 розділу II [2] введено другий рівень - порушник корпоративного типу має змогу створення спеціальних технічних засобів, вартість яких співвідноситься з можливими фінансовими збитками, що виникатимуть від порушення конфіденційності, цілісності та підтвердження авторства інформації, та третій рівень – порушник, який має науково-технічний ресурс, який прирівнюється до науково-технічного ресурсу спеціальної служби економічно розвинутої держави.

Поширення кіберзагроз на всі сфери суспільного життя та державного управління, а також удосконалення методів їх реалізації, зокрема із застосуванням штучного інтелекту, актуалізують необхідність глибокого аналізу механізмів та сценаріїв кібератак. Це вимагає перегляду існуючих стратегій і

тактик протидії, а також пошуку нових підходів підвищення ефективності кіберзахисту.

Мета роботи полягає в розробці методики попередження комп'ютерних атак типу «Man In The Middle», що здійснюються з використанням генеративного штучного інтелекту.

Для досягнення поставленої мети поставлено наступні **завдання**:

- провести пошук, збір, систематизацію та аналіз інформації про комп'ютерні атаки, що здійснюються з використанням штучного інтелекту;
- здійснити аналіз методів виявлення вторгнень комп'ютерних атак типу «Man In The Middle»;
- практично реалізувати попередження атак типу «Man In The Middle», що здійснюються з використанням генеративного штучного інтелекту.

Об'єкт дослідження: системи передачі даних, які можуть стати об'єктами атак типу «Man In The Middle», що здійснюються із використанням генеративного штучного інтелекту.

Предмет дослідження: методи попередження комп'ютерних атак типу «Man In The Middle», що здійснюються з використанням генеративного штучного інтелекту.

Апробація результатів: Ласкавчук М.А., Загоруйко Л.В. Попередження комп'ютерної атаки типу «Man In The Middle», що здійснюється за допомогою генеративного штучного інтелекту. Вісник студентського наукового товариства Донецького національного університету імені Василя Стуса. Вінниця: ДонНУ імені Василя Стуса, 2025. Вип. 17, т. 1. 248 с.

М. Ласкавчук Специфіка атак типу «Man In The Middle». Proceedings of the 2nd International Scientific and Practical Conference "Achievements of Science and Applied Research" 91-95 p.

Структура та обсяг роботи: кваліфікаційна (бакалаврська) робота складається зі вступу, трьох розділів, висновків, списку використаних джерел із 45 найменувань, додатків. Обсяг основного тексту становить 53 сторінки.

РОЗДІЛ 1

ЗАГАЛЬНІ ВІДОМОСТІ ПРО КОМП'ЮТЕРНІ АТАКИ З ВИКОРИСТАННЯМ ШТУЧНОГО ІНТЕЛЕКТУ

1.1. Відомості про комп'ютерні атаки, які здійснюються з використанням штучного інтелекту

Відповідно до статті 1 Закону України «Про основні засади забезпечення кібербезпеки України» [3], кібератака (комп'ютерна атака) – це спрямовані (навмисні) дії в кіберпросторі, які здійснюються за допомогою засобів електронних комунікацій (включаючи інформаційно-комунікаційні технології, програмні, програмно-апаратні засоби, інші технічні та технологічні засоби і обладнання) та спрямовані на досягнення однієї або сукупності таких цілей: порушення конфіденційності, цілісності, доступності електронних інформаційних ресурсів, що обробляються (передаються, зберігаються) в комунікаційних та/або технологічних системах, отримання несанкціонованого доступу до таких ресурсів, порушення безпеки, сталого, надійного та штатного режиму функціонування комунікаційних та/або технологічних систем; використання комунікаційної системи, її ресурсів та засобів електронних комунікацій для здійснення кібератак на інші об'єкти кіберзахисту.

До класичних комп'ютерних атак можна віднести атаки типу DoS, атаки MITM, фішинг, цільові фішингові атаки, атаки шляхом впровадження (SQL-ін'єкція та XSS-вразливість), а також атаки із використанням шкідливого програмного забезпечення [4-5].

Однією з новітніх загроз є комп'ютерні атаки, що використовують технології штучного інтелекту, наслідками яких можуть бути: помилкова класифікація даних, створення синтетичних даних, несанкціонований доступ до інформації та її аналіз [5].

Нові типи кібератак на основі ШІ поділяються на [5]:

- неправильну класифікацію даних;
- синтетичну генерацію даних;
- несанкціонований аналіз даних.

Класифікацію кібератак на основі ШІ та класичних кібератак представлено на рисунку 1.1.

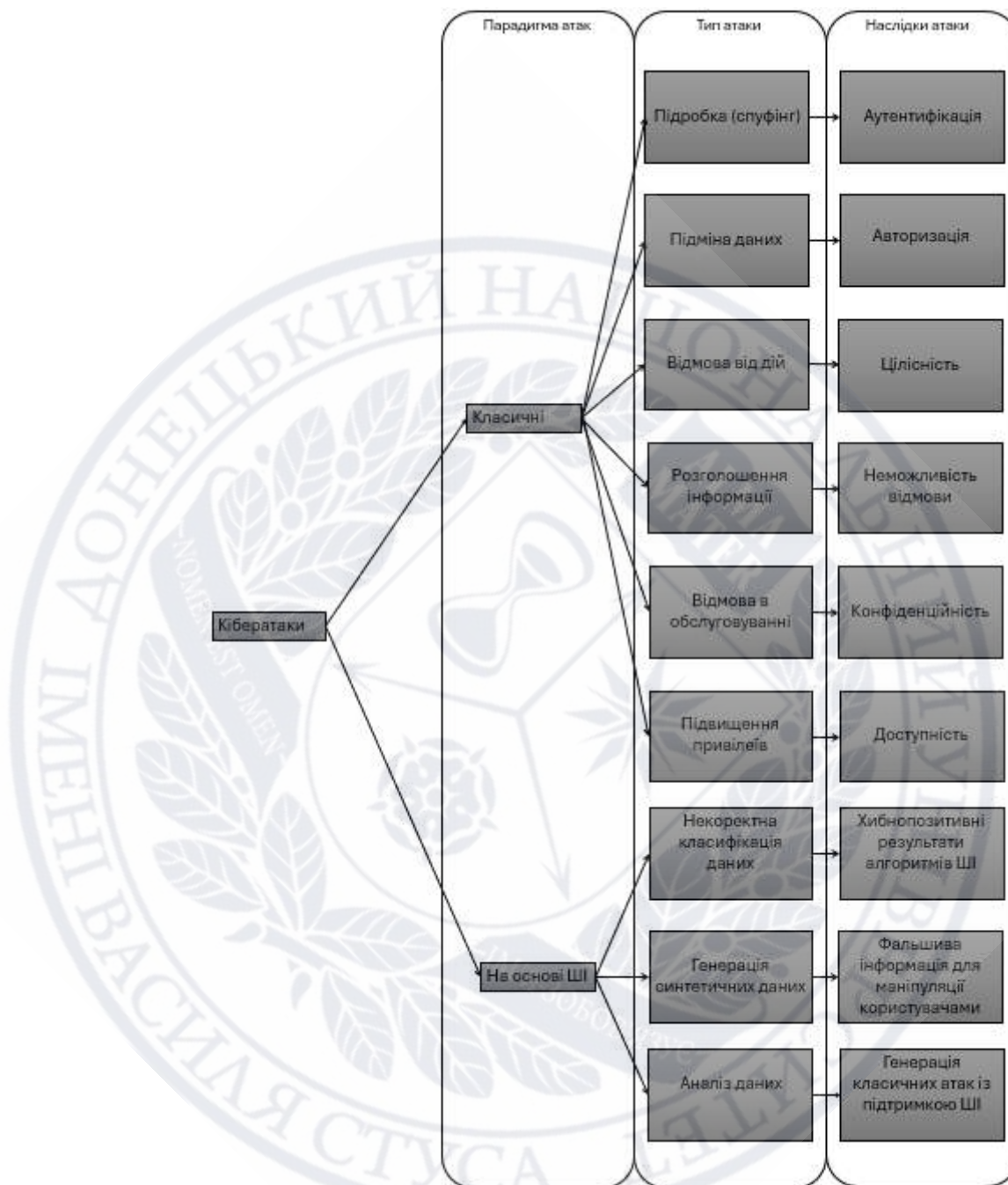


Рисунок 1.1 - Класифікація кібератак [8]

Генеративний ШІ став потужним інструментом для зловмисників, що дозволяє створювати високотехнологічні кібератаки. Одним з прикладів його використання є автоматизація створення фішингового контенту. Генеративні моделі можуть створювати реалістичні, схожі на людські електронні листи, які вводять в оману навіть найсучасніші системи виявлення. Крім того, зловмисники

використовують генеративний ШІ для створення так званих «глибоких підробок»: підроблених зображень, відео чи аудіозаписів, які використовують для того, щоб видавати себе за людей, вчиняти шахрайські дії або маніпулювати довірою [6].

Генеративний ШІ також використовується для створення складних шкідливих програм. Створюючи варіанти шкідливого програмного забезпечення, які обходять системи виявлення на основі сигнатур, генеративні моделі роблять традиційні засоби захисту застарілими та неефективними. Такі методи, як генерація фальшивих вхідних даних, передбачають навчання моделей для виявлення вразливостей у системах безпеки, що уможливлює автоматизовані та широкомасштабні атаки. Ці можливості демонструють, як генеративний ШІ використовується зловмисниками як ефективний інструмент, що полегшує реалізацію складних кіберзагроз. Оскільки можливості сучасних моделей генеративного ШІ продовжують розвиватися, це вимагає удосконалення «оборонних стратегій» для протидії атакам із його широким використанням [6].

Інструменти генеративного ШІ допомагають кіберзловмисникам здійснювати атаки через вразливості системи безпеки та надмірну залежність від трансформаційних технологій [7]. Приклад використання ШІ для здійснення порушень у сфері кібербезпеки показаний рисунку 1.2:



Рисунок 1.2 – Порушення із застосуванням ШІ

Генеративний ШІ сприяє проведенню складних кібератак та багаторівневих кібератак, які охоплюють широкий спектр векторів впливу – від маніпуляцій на рівні апаратного забезпечення до втручання в роботу операційної системи та програмного забезпечення. Це призводить до значного послаблення

ефективності комплексу засобів захисту та підвищення ризику компрометації критично важливих інформаційних систем та ресурсів [8].

Атаки, які здійснюються за допомогою генеративного ШІ, та найчастіше здійснюються останнім часом [8]:

- а) апаратні атаки;
- б) атаки на ОС;
- в) атаки на програмне забезпечення.

Апаратні атаки. Зазвичай характеризуються необхідністю взаємодії з фізичними пристроями. Хоча моделі генеративного ШІ не мають прямої взаємодії з апаратними компонентами, вони генерують атаки, обробляючи дані або інформацію, пов'язану з апаратним забезпеченням. Однією з найвідоміших атак на апаратному рівні є атака побічного каналу. У дослідженні [9] автори представляють використання умовно-генеративних мереж для вдосконалення стратегій атак на побічні канали. В ній підкреслюється здатність умовно-генеративних мереж створювати додаткові потоки даних, збільшуючи розмір набору профілів і підвищуючи загальну ефективність атак профілювання. У дослідженні [10] запропоновано підхід з використанням умовно-генеративних мереж та «сіамської мережі» для генерації фальшивих сигналів для аналізу побічних каналів на основі глибокого навчання. Запропонована модель генерує реалістичні сліди витоку, які можуть бути успішно використані для аналізу побічних каналів, навіть якщо набір даних, доступний з апаратного пристрою, не є оптимальним.

Атаки на рівні ОС. Націлені на операційну систему апаратної платформи з метою використання вразливостей, виконання шкідливих дій або отримання несанкціонованого доступу. Найчастіші типи атак включають використання руткіти (англ. rootkit), підвищення привілеїв користувачів та використання системних ресурсів. Ці атаки можуть порушити цілісність системи, конфіденційність даних або контролювати системи. Незважаючи на те, що моделі генеративного ШІ працюють на вищому рівні абстракції без низькорівневого доступу до системи, вони можуть аналізувати інформацію з операційних систем,

щоб допомогти в атаках на рівні ОС. У дослідженні [11] автори продемонстрували подібну атаку, встановивши зворотній зв'язок між GPT 3.5 та вразливою віртуальною машиною через SSH. Модель GPT аналізує стан машини, виявляє вразливості та пропонує стратегії атаки, що виконуються всередині віртуальної машини.

Атаки на програмне забезпечення. Атаки програмного рівня на комплекс засобів захисту використовують вразливості в програмних компонентах, що контролюють фізичні процеси. Однак найпоширеніший випадок використання на програмному рівні пов'язаний зі зловмисниками, які використовують моделі генеративного ШІ для створення шкідливого програмного забезпечення. Ці атаки можуть маніпулювати логікою програмного забезпечення, порушувати алгоритми управління або впроваджувати шкідливий код, що призводить до несанкціонованого контролю, крадіжки даних або пошкодження фізичної інфраструктури. Наприклад, модель ШІ ChatGPT, здатна генерувати шкідливий код, використовуючи своє машинне навчання на різноманітних наборах даних, включаючи сховища коду та доступну технічну документацію виробників (розробників). За допомогою спеціальних інструкцій або навчання на наборах даних, що містять приклади шкідливого програмного забезпечення, ці моделі можуть синтезувати нові зразки шкідливого програмного забезпечення, код для експлуатації або фішингові скрипти, які часто неможливо ідентифікувати за допомогою типового антивірусу. Крім того, моделі, подібні до ChatGPT, можуть бути використані зловмисниками для застосування існуючих методів SQL-ін'єкцій, генеруючи адаптоване штатне навантаження SQL для експлуатації вразливості систем баз даних. У роботі [12] представлено доказ концепції, в якій ChatGPT використовується для розповсюдження шкідливого програмного забезпечення, уникаючи при цьому виявлення. У дослідженні [13] автори демонструють потенційне використання ChatGPT для створення декількох шкідливих програм (наприклад, програм-вимагачів, черв'яків, клавіатурних шпигунів, шкідливого програмного забезпечення для перебору паролів, та навіть безфайлового шкідливого програмного забезпечення).

Нині розроблено і застосовують низку систем ШІ, які призначені для здійснення комп'ютерних атак різного типу, деякі з них, які найчастіше застосовують, наведено в таблиці 1.1:

Таблиця 1.1 - Інструменти на базі штучного інтелекту, що використовують аналіз даних для вчинення комп'ютерних злочинів

DeerHack	Хакерський ШІ з відкритим вихідним кодом, навчився зламувати веб-додатки за допомогою нейронної мережі, яка може експлуатувати різні види вразливостей, відкриваючи можливості для безлічі хакерських систем на основі ШІ в майбутньому [14].
EagleEye	Інструмент на базі штучного інтелекту для розвідки інформації в соціальних мережах [15].
DeerLocker	Новий варіант шкідливого програмного забезпечення зі штучним інтелектом. У поєднанні зі штучним інтелектом DeerLocker може ідентифікувати свої цілі за допомогою розпізнавання голосу та обличчя, а також геолокації. Шкідливе програмне забезпечення буде приховане в нешкідливих операторських додатках, поки не будуть виконані умови запуску; ШІ робить майже неможливим зворотне проектування цих умов запуску [16].
GyoiThon	Інструмент на базі штучного інтелекту для збору інформації та автоматичної експлуатації [17].
Malware-GAN	Інструмент на базі штучного інтелекту, що використовується для створення шкідливого програмного забезпечення, яке може обходити механізми безпеки [18].

Щодо кожної із систем проводилися низка досліджень, які були спрямовані на вивчення особливостей та можливостей її застосування для здійснення

несанкціонованого аналізу даних і нецільового використання інформаційних систем.

Основним методом здійснення представлених атак є змагальне машинне навчання, як метод, що ґрунтується на машинному навчанні, суть якого базується на використанні наявних «сліпих зон» між сукупністю даних, що обробляються в процесі навчання моделі. У процесі навчання «шкідливого» ШІ визначаються слабкі сторони системи, що захищається, і вносяться невеликі зміни в її масиви даних [19]. У зв'язку з цим, у моделі, що захищається, не формуються стійкі зв'язки між цільовими значеннями, що надалі призводить до неправильних класифікацій із перетином межі ухвалення рішення та помилкового віднесення даних до іншого класу. Надалі інший ШІ під час аналізу системи, що захищається, не покаже наявності шкідливих даних, оскільки відбулася підміна класифікації даних [20].

Низка дослідників [21-26], аналізуючи етап машинного навчання, виділяють чотири основні напрямки атак AML:

- атаки на рішення класифікатора, включно з отруйними (причинними) атаками на етапі навчання та дослідницькими (такими, що ухиляються) атаками навченої моделі на етапах тестування;

- атаки на цілісність моделі, що призводять до неправильної класифікації, або на придатність моделі за наявності високої частоти неправильних класифікацій;

- цілеспрямовані атаки, коли змагальні вибірки націлені на досягнення певного цільового значення, або невибіркові атаки, коли вибірки не націлені на певне цільове значення;

- атаки на конфіденційність, метою якої виступає вилучення інформації з класифікатора.

Інший підхід [27] до класифікації атак пропонується здійснювати на основі:

- складності за критерієм наслідків, що варіюються від незначного зниження достовірності прогнозів моделі до неправильної класифікації всіх невидимих точок даних;

- отриманого противником знання, наприклад, типу «атака білої скриньки» з метою отримання зловмисником знань про навчальну модель (її архітектуру, мережевий трафік, який вона аналізує, і її функції, які використовуються для підтримки навчання).

У даній роботі основна увага зосереджена на комп'ютерній атаці MITM, у межах якої атакуючим суб'єктом може виступати генеративний ШІ.

1.2. Поняття комп'ютерної атаки Man In The Middle

Комп'ютерна атака типу MITM - це атака, коли суб'єкт (генеративний ШІ) використовує різні методики та технічні рішення, спрямовані на отримання доступу до трафіку та його декодування. Дана комп'ютерна атака реалізується із застосуванням алгоритмів перехоплення трафіку, що передається між двома кінцевими пристроями [28]. На рисунку 1.3 наведено, як відбувається типова атака типу MITM.

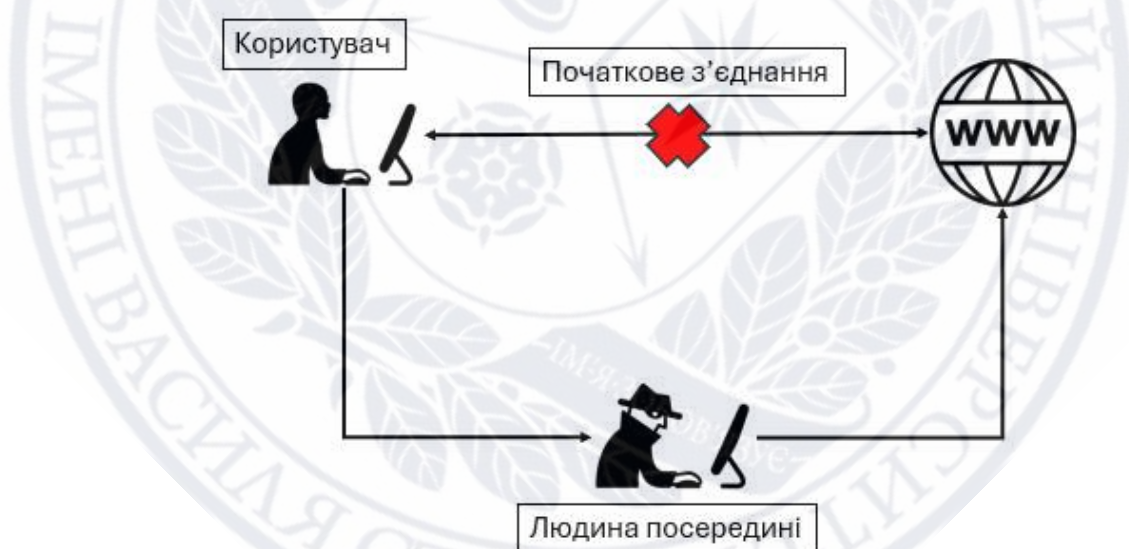


Рисунок 1.3 – Процес атаки «людина посередині»

Класична MITM-атака проходить у два етапи. Перший - перехоплення даних, коли злочинець інтегрується в середовище передачі даних. Далі за допомогою спуфінгу реалізується підміна IP-адрес, ARP-повідомлень, сервера доменних імен тощо. Атаки MITM найчастіше використовують ARP-кеш, який являє собою локальний кеш із призначеними IP-адресами і зіставленими фізичними унікальними ідентифікаторами пристроїв у мережі (MAC-адресами).

У результаті реалізуються завдання отримання відомостей про структуру досліджуваної мережі та зіставлення локальних ідентифікаторів і універсальних ідентифікаторів мережі (MAC-адрес).

Один з випадків атак типу MITM відомий як «динамічне підслуховування», при якому зловмисники створюють незалежні асоціації з жертвами і впливають на комунікацію між користувачами, модифікуючи повідомлення між двома користувачами, щоб вплинути на довіру між ними. Зловмисник має можливість перехоплювати повідомлення між ними з метою їх модифікації. Більшість криптографічних угод здійснюють автентифікацію кінцевих точок переважно для запобігання MITM-атакам.

Другий етап – дешифрування, отримання доступу до зашифрованих даних.

Ще один випадок – це атака типу «людина в браузері», це особливий підтип MITM-атаки, яка здійснюється на стороні кінцевої точки зв'язку [29]. В основі «людина в браузері» лежить ідея розміщення шкідливого прозорого браузера між браузером жертви та веб-сервером, до якого жертва звертається для отримання послуги (наприклад, банківської послуги, соціальної мережі тощо). Метою атаки є перехоплення, запис і маніпулювання будь-яким обміном даними між жертвою і постачальником послуг [29-30].

Сценарій атаки типу «людина в браузері» наведено на рисунку 1.4:

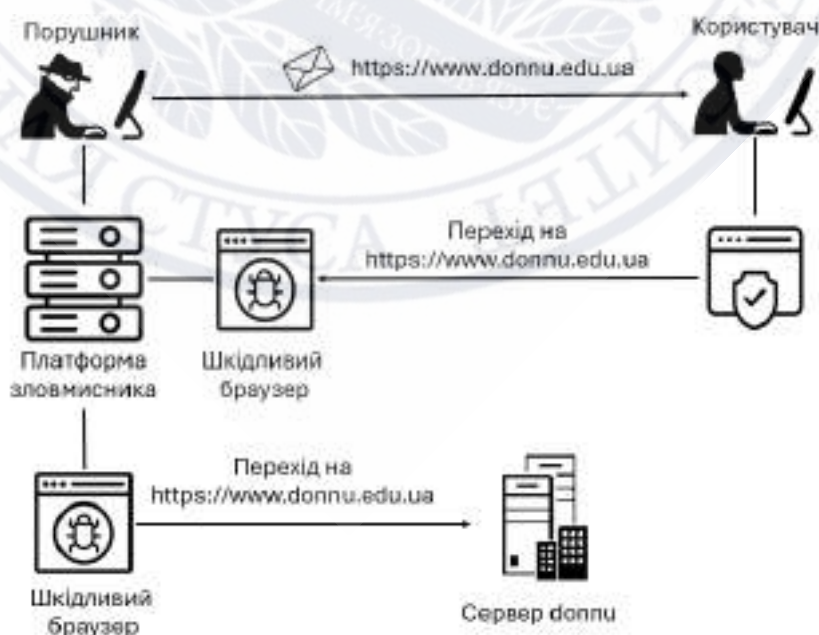


Рисунок 1.4 – Сценарій атаки типу «людина в браузері»

В процесі реалізації атаки типу «людина в браузері» задіяні:

- **Порушник:** комп'ютерний злочинець, метою якого є порушення приватності жертви, а також конфіденційності, цілісності та доступності даних, якими обмінюються дві кінцеві точки. Зловмисник встановлює і контролює шкідливий сервер, який слугує прозорим веб-браузером для кінцевого користувача [30];

- **Авторизований користувач:** потерпілий, який за допомогою фішингових методів отримує доступ до шкідливого сервера (на якому розміщений прозорий веб-браузер), налаштованого зловмисником, і починає користуватися веб-додатком [30];

- **Ціль:** веб-додаток, здатний надавати конфіденційні та/або цінні послуги (наприклад, банківські послуги, освітні послуги, поштові сервіси тощо), доступ до яких здійснюється за допомогою механізму автентифікації [30].

Наступні кроки описують техніку (сценарій) атаки «людина в браузері» [30]:

- а) порушник створює підроблене гіперпосилання, що вказує на шкідливий сервер. Посилання надсилається жертві, яку за допомогою фішингових методів заманюють натиснути на нього та пройти автентифікацію у веб-додатку;

- б) натиснувши на шкідливе гіперпосилання у власному браузері, користувач підключається до шкідливого сервера, де знаходиться прозорий веб-браузер разом з низкою програм (наприклад, веб-проксі, сніффер, кейлоггер, надбудови тощо), які зловмисник використовує для перехоплення, запису та маніпулювання даними, якими обмінюються жертва та цільова веб-програма;

- в) «шкідливий» сервер отримує доступ до цільового веб-додатку (наприклад, банківського додатку, соціальної мережі тощо) через шкідливий веб-браузер абсолютно прозорим способом. Під час атаки користувач може використовувати всі функціональні можливості цільового додатка (плюс ті, які тепер може запропонувати зловмисний сервер). Всі операції жертви тепер можуть бути перехоплені, а отже, записані і використані зловмисником для маніпуляцій. Наприклад, порушник зможе перехопити облікові дані, надані

жертвою в якійсь формі автентифікації, незалежно від того, які протоколи безпеки налаштовані у веб-додатку.

Атаки MITM можуть починатися з [31]:

- а) отруєння кешу протоколу дозволу адрес;
- б) підробки DNS;
- в) підміни IP-адреси;
- г) підміна сеансу;
- д) перехоплення SSL.

Отруєння кешу протоколу дозволу адрес. ARP є надійною та життєво важливою процедурою для інфраструктури локальних мереж. Зловмисники вносять зміни в кеш-таблицю локального ARP і підпорядковують MAC-адресу хоста цільовій IP-адресі. Атака здійснюється з метою отримання доступ до конфіденційної інформації користувача. Ці підміни ARP можна розділити на два елементарних типи, а саме: підміна хоста та шлюзу внутрішньої мережі. Коли користувач бажає з'єднатися з іншим користувачем, що належить до подібної мережі з неідентифікованою MAC-адресою, це призводить до широкомовної розсилки ARP-запиту в мережі. Доступ до кешу легко здійснити, оскільки відсутні будь-які відповідні методи верифікації. Пристрій-джерело може зберегти IP-адресу в MAC-записі для швидкої передачі даних наступного разу. Варіант сценарію атаки типу «Отруєння ARP» наведено на рисунку 1.5:



Рисунок 1.5 – Сценарій атаки типу «Отруєння ARP»

Атака типу «Отруєння APR» (представлена на рисунку 1.5) відбувається за таким сценарієм:

- щоб виконувати зіставлення MAC- і IP-адрес, кожен хост у локальній мережі підтримує локальну таблицю, звану ARP Cache;
- через «безвідмовність» протоколу, зловмисник, який перебуває за пристроєм користувач_1, здатний фальсифікувати пакети, оскільки кожний пристрій у мережі може відповісти на ARP-запит, навіть якщо запит не був адресований йому, а жертва атаки прийме таку відповідь за дійсну;
- перш ніж відбудеться ця інкапсуляція, хосту-відправнику потрібна MAC-адреса хоста-одержувача, з огляду на IP-адресу, ARP може динамічно визначати MAC-адресу відповідного хоста, а користувач_1 змушувати хост Router «підмінити» істинну MAC-адресу на свою, що порушує логіку зіставлення адрес і ARP-таблиць у мережі.

Після проникнення зловмиснику доступні різні операції з трафіком, він може переглядати і змінювати пакети даних, перш ніж відправити його в істинний пункт призначення, або зовсім перекрити шляхи комунікації.

Підробка DNS. Ця атака є найнебезпечнішою атакою, що виконується з використанням кешування для підвищення продуктивності. Вона може складатися з трьох основних типів атак. Перший, підслуховування пакетів даних під час процесу відповіді. Другий, використання атаки на день народження для розміщення кешу. Третій, злом несанкціонованого DNS. Щоб завершити підміну DNS, зловмисник контролює локальний доступ до DNS, змушуючи ціль використовувати неавторизований сервер. DNS відповідає за час очікування записів, і DNS повторно вводить свій кеш. Зловмисники використовують його для атак підміни DNS. Щоб знайти кеш, необхідно використовувати два методи. Перший, вставити в мережу зловмисний DNS, який генерує шахрайські дані, а другий, надіслати фальшиву DNS-відповідь перед відправленням авторизованої DNS-відповіді.

Атаки на DNS можна умовно розділити на дві категорії:

а) перша категорія - це атаки на вразливості в DNS-серверах. Внаслідок DNS-атак користувач ризикує не потрапити на потрібну сторінку. Під час введення адреси сайту атакований DNS перенаправлятиме запит на підставні сторінки. У результаті переходу користувача на помилкову IP-адресу зловмисник може отримати доступ до його особистої інформації;

б) друга категорія - це DoS-атаки, що призводять до непрацездатності DNS-сервера. У разі недоступності DNS-сервера користувач не зможе потрапити на потрібну йому сторінку, оскільки його браузер не зможе знайти IP-адресу, що відповідає введеній адресі сайту.

Для здійснення такої атаки, необхідно передбачити правильний ID-запит. Якщо атака проводиться всередині локальної мережі, то не складає труднощів прослуховувати трафік.

Підміна IP-адреси. Атака, при якій зловмисник маніпулює вихідною IP-адресою, щоб видати себе за іншу особу або пристрій у мережі. Оскільки IP-адреса визначає відправника та отримувача пакета, її підробка дозволяє зловмиснику втручатися у зв'язок, отримувати або змінювати інформацію без відома справжніх сторін. Це дає змогу контролювати потік даних і приховувати реальне джерело атаки.

Підміна сеансу. Зловмисник викрадає або видає себе за токен сеансу користувача, щоб отримати несанкціонований доступ або виконати дії від його імені.

Так, відомий метод т. зв. «сніффінгу» (перехоплення даних) - один із способів атаки, за допомогою якого зловмисник перехоплює мережу і викрадає ідентифікатор сеансу. Зловмисник продовжує стежити за мережевим трафіком жертви і намагається виявити будь-який пакет, що проходить в незашифрованому вигляді, якщо зловмисник знаходить пакет без шифрування, то він намагається з'ясувати, чи містить він ідентифікатор сеансу чи ні. Якщо зловмисник отримує ідентифікатор сеансу в пакеті, то, використовуючи цей ідентифікатор, він перехоплює вже створений сеанс між жертвою і сервером, зловмисник має

можливість створити новий сеанс і отримати всю інформацію про цільову віртуальну машину.

Також широко використовується метод т. зв. «перебору грубої сили» - атаки, яка здатна «зламати» (підібрати) будь-який символ, цифру, символ і спеціальний символ, комбіноване ім'я користувача або пароль, або будь-яке слово. Ця атака займає багато часу, але дає гарантію отримання бажаного результату зловмиснику. У цій атаці перебором, зловмисник може встановити комбінацію символів, спеціальних символів, символів і чисел відповідно до своїх вимог, навіть зловмисник може встановити кількість слів. Припустимо, що жертва має 8 цифр пароля, здогадуючись, що зловмисник встановив довжину пароля до 8, тож цей перебір перевіряє всі комбінації слів, встановлені зловмисником, до 8-ї цифри. У цьому типі атаки зловмисник повинен мати багато терпіння і часу, щоб завершити її. Для злому сеансу зловмисник проводить атаку грубою силою в цільовій мережі і перевіряє всі комбінації, задані користувачем, щоб зламати ідентифікатор сеансу і передати правильний ідентифікатор сеансу зловмиснику, а зловмисник, використовуючи його, перехоплює дійсний вже створений сеанс.

Міжсайтовий скриптинг (XSS) - це ще один спосіб атаки для викрадення ідентифікатора сеансу. Цей XSS також відомий, як атака впровадження коду на стороні клієнта, цей код має можливість виконати шкідливий сценарій (шкідливе корисне навантаження) на будь-якому веб-сайті або веб-додатку. Ця атака в основному використовує вразливості веб-сайтів, і незалежно від того, що вводить користувач на сайті, вона робить це введення незашифрованим або закодованим і надсилає цю інформацію зловмиснику у вигляді звичайного тексту. Але є великий недолік для зловмисника - ця атака працює тільки для вразливої цілі.

Перехоплення SSL. Зловмисник перехоплює та розшифровує зашифровані SSL/TLS-повідомлення, використовуючи підроблені сертифікати. Оскільки під час атаки на перехоплення SSL веб-браузери зазвичай, не видають жодних попереджень кінцевим користувачам. Атака перехоплення сеансу на основі SSL є дуже небезпечною атакою. Це зумовлено тим, що виконавши успішну атаку на перехоплення SSL, порушники можуть легко отримати облікові дані жертви і

використовувати їх не компрометуючи себе. Це можуть бути банківські рахунки, дані кредитних карток, та інша конфіденційна інформація. Сценарій базової атаки на перехоплення SSL продемонстровано на рис. 1.6:

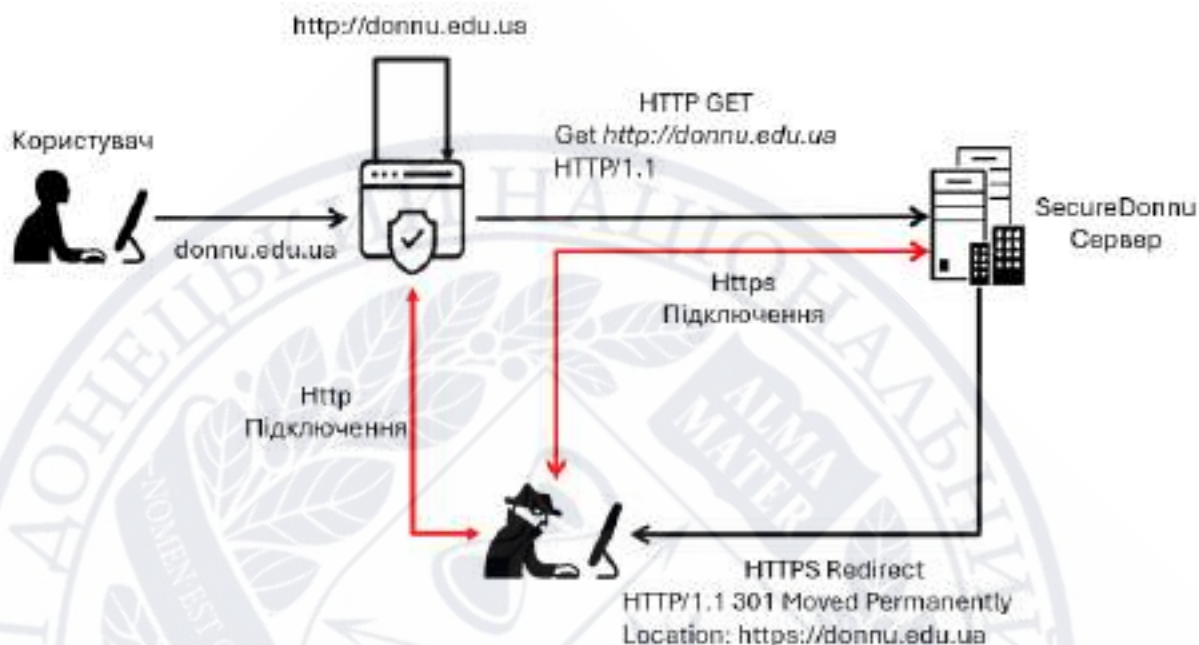


Рисунок 1.6 - Сценарій базової атаки на перехоплення SSL

Атаки на перехоплення SSL здійснюється наступним чином:

- а) користувач заходить на сайт (наприклад, університетський) за його URL-адресою, наприклад: www.donnu.edu.ua;
- б) веб-браузер автоматично додає <http://> перед www.donnu.edu.ua, таким чином роблячи <http://www.donnu.edu.ua>;
- в) тепер браузер звертається до сервера SecureDonnu методом HTTP GET;
- г) сервер SecureDonnu отримує запит і перенаправляє його на <https://www.donnu.edu.ua>;
- д) оскільки перенаправлення HTTPS проходить через незахищену мережу, можливі наступні варіанти:
 - злоумисник може відключити HTTPS для користувача;
 - злоумисник може взаємодіяти з SecureDonnu через HTTPS і робити все від імені реального користувача;
 - зазвичай не відображається жодних попереджень у веб-браузері;

- користувач, як правило, не знає про цю атаку;

д) користувач все ще може підключатися до сервера SecureDonnu через HTTP через зловмисника, навіть якщо університет фактично забезпечив HTTPS для своїх користувачів.

В дослідженні [31] приведені інші поширені приклади атак типу MITM:

- а) атаки на бездротову мережу Wi-Fi;
- б) атака «злий двійник» (від англ. evil twin);
- в) перехоплення електронної пошти;
- г) перехоплення бездротового зв'язку «Bluetooth».

Атаки на бездротову мережу Wi-Fi. Зловмисник використовує вразливості Wi-Fi для перехоплення та маніпулювання мережевим трафіком. Зловмисник, перебуваючи в зоні дії Wi-Fi мережі, може здійснити перехоплення трафіку між точкою доступу та клієнтом за допомогою технологій пасивного моніторингу. У випадках недостатнього шифрування (наприклад, використання застарілих протоколів WEP або неправильно реалізованого WPA/WPA2), ці дані можуть бути розшифровані, що дає змогу отримати доступ до конфіденційної інформації, включно з логінами, паролями, файлами, сесіями доступу. Процес (сценарій) такої атаки наведено на рисунку 1.7:

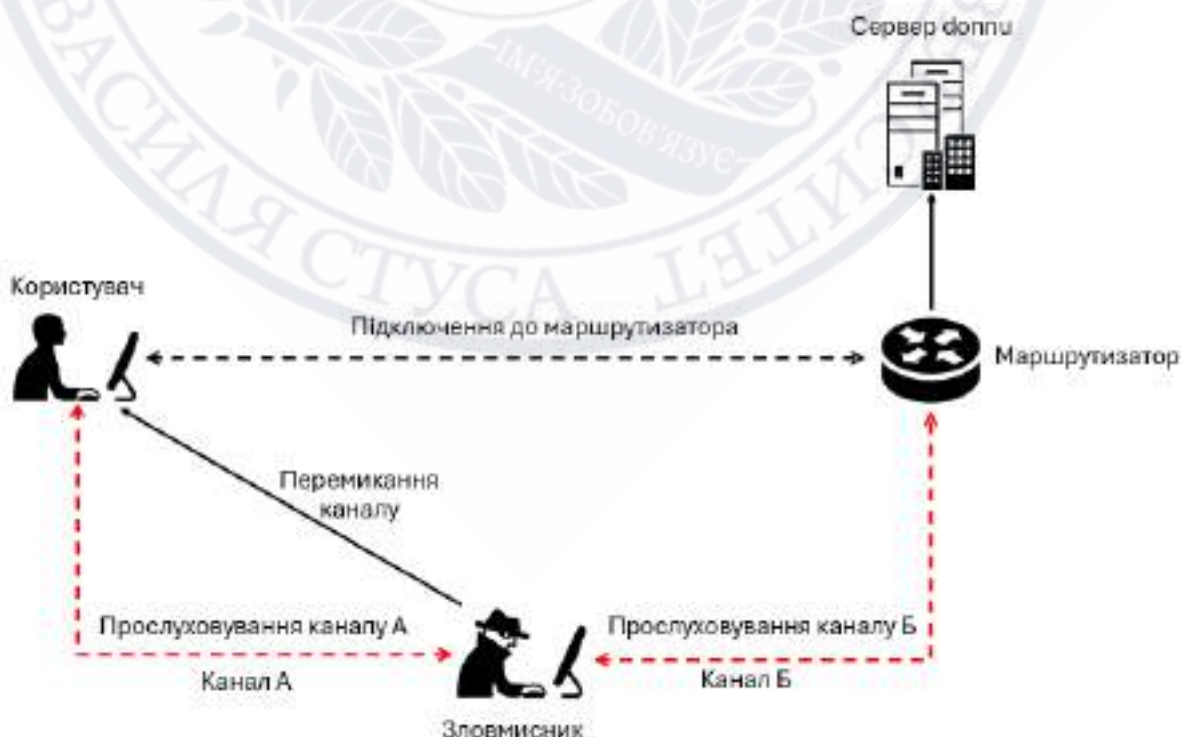


Рисунок 1.7 – Процес (сценарій) атаки на бездротову мережу Wi-Fi

На рисунку 1.7 схематично зображено реалізацію атаки типу MITM, метою якої є несанкціоноване прослуховування і, за необхідності, модифікація даних, які передаються між користувачем і сервером університету (donnu).

Учасниками взаємодії є:

- авторизований користувач – легітимний клієнт, який здійснює підключення до комутатора з метою доступу до сервера donnu;
- комутатор – мережевий пристрій, який передає трафік між користувачем та сервером;
- сервер donnu – цільовий ресурс, з яким користувач взаємодіє.
- зломисник – третя сторона, яка здійснює атаку, підключившись до тієї ж локальної мережі (наприклад, через той самий комутатор).

Атака на бездротову мережу, яка зображена на рисунку 1.7, полягає в наступному:

- перехоплення трафіку (канал А): зломисник ініціює пасивне прослуховування трафіку між користувачем і комутатором. Для цього він здійснює атаку типу ARP-спуфінгу або використання техніки дзеркалювання портів. Це дозволяє зломиснику перехоплювати пакети, що надходять до/від користувача;
- перехоплення трафіку (канал Б): зломисник перехоплює пакети, які передаються між комутатором і сервером. Це дозволяє йому бачити повний контекст з'єднання користувача із сервером;
- перемикання каналу: у випадку активного MITM, зломисник може не тільки перехоплювати, а й змінювати дані, що передаються. Він виступає посередником, отримує запит від користувача, можливо модифікує його, пересилає до сервера, а потім аналогічно передає відповідь назад користувачу. Таким чином він повністю контролює інформаційний потік.
- формування фальшивих каналів зв'язку: зломисник імітує наявність прямого з'єднання між користувачем і сервером, у той час як реальна передача даних відбувається через його комп'ютер. Це дозволяє йому непомітно

впроваджувати шкідливі скрипти, перехоплювати облікові дані, змінювати інформацію або здійснювати інші несанкціоновані дії.

Атака «злий двійник» (від англ. evil twin). Це підроблена точка доступу Wi-Fi, яка виглядає довіреною, але налаштована на підслуховування бездротового зв'язку [32]. Для проведення цієї атаки не потрібно майже ніякого додаткового обладнання, достатньо ноутбука з бездротовою картою. Можна використовувати «raspberry pi» з бездротовою картою та павербанком, щоб виконати атаку віддалено та дискретно. Для спрощення рекомендується використовувати готові інструменти, такі як «airodump-ng», «airbase-ng» і «aireplay-ng», які можна знайти в операційних системах для атак, таких як Kali Linux, Parrot OS і багатьох інших [32]. Процес (сценарій) атаки «злий двійник» наведено на рисунку 1.8:

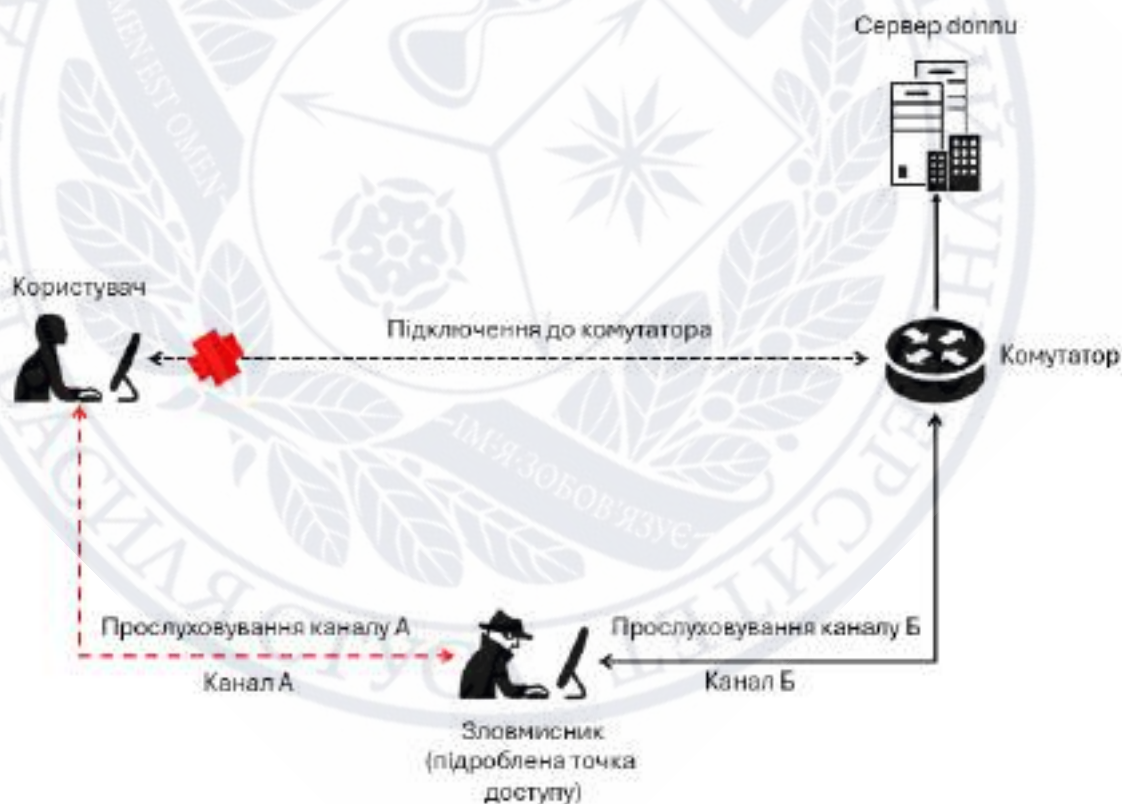


Рисунок 1.8 – Процес (сценарій) атаки «злий двійник»

На рисунку 1.8 схематично зображено атаку «злий двійник», яка полягає у створенні зловмисником фальшивої (клона) точки доступу Wi-Fi, яка імітує легітимну бездротову мережу (наприклад, університетську мережу). Авторизований користувач, не підозрюючи підміну, автоматично або вручну

підключається до цієї точки доступу. У результаті зловмисник отримує змогу перехоплювати, аналізувати, а в деяких випадках, і змінювати трафік користувача.

Учасниками взаємодії є:

- авторизований користувач – легітимний клієнт, який намагається підключитися до мережі університету;
- комутатор – мережевий пристрій, через який має відбуватись справжній обмін даними між клієнтом і сервером;
- сервер donnu – цільовий ресурс, з яким користувач взаємодіє;
- зловмисник (підроблена точка доступу) – атакуючий, який створює клон точки доступу та здійснює атаку.

Атака «злий двійник», яка зображена на рисунку 1.8 полягає в наступному:

- створення підробленої точки доступу: зловмисник налаштовує пристрій (наприклад, ноутбук або смартфон з функцією точки доступу) з ідентичними параметрами справжньої мережі: SSID (ім'я мережі), тип шифрування, MAC-адреса;
- підключення користувача до підробленої точки доступу (канал А): через автоматичне підключення або через обман користувач підключається до фальшивої мережі зловмисника замість справжньої. На рисунку це зображено як розрив прямого з'єднання до комутатора (позначка червоного X);
- перехоплення трафіку (Канал А і В): усі запити користувача спрямовуються через зловмисника (Канал А), який виконує роль "проксі". Зловмисник може пересилати запити до справжнього сервера (Канал В) та повертати відповіді, маніпулювати трафіком (вставляти шкідливі скрипти, підміняти дані), а також повністю блокувати доступ або перенаправляти користувача на фішингові сайти.
- прослуховування та збір даних: зловмисник фіксує усі запити, включно з введеними обліковими даними, запитами на сервіси, cookie-файлами, сесіями тощо.

Перехоплення електронної пошти. Зловмисник перехоплює та читає електронні повідомлення, якими обмінюються два поштових сервери.

Порушники використовують переважно фішингові атаки, щоб скомпрометувати паролі до поштових акаунтів. Після викрадення облікових даних вони створюють шкідливі переадресатори та фільтри з метою перехоплення конфіденційної інформації з електронної пошти, зокрема повідомлень, що містять фінансову інформацію, платіжні вимоги, банківські реквізити тощо. Процес (сценарій) перехоплення електронної пошти наведено на рисунку 1.9:

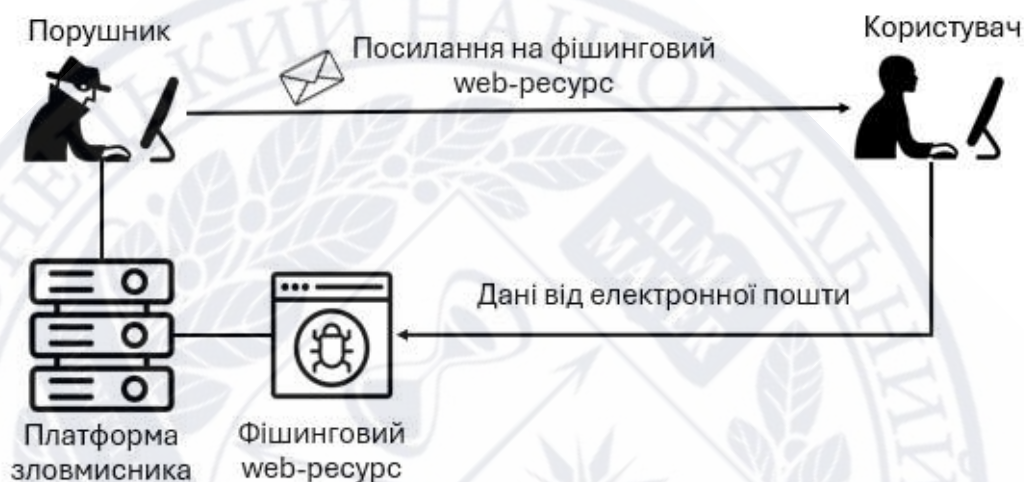


Рисунок 1.9 – Процес (сценарій) перехоплення електронної пошти за допомогою фішингу

На рисунку 1.9 відображено типову схему фішингової атаки, орієнтованої на викрадення даних доступу до електронної пошти користувача. Після атаки зловмисник може налаштувати переадресацію всіх вхідних листів і отримувати їх на свою пошту.

Учасниками взаємодії є:

- зловмисник – розміщує фішинговий web-ресурс на своїй платформі (сервері), що зовні імітує справжній інтерфейс авторизації електронної пошти;
- платформа зловмисника – обробляє та зберігає викрадені облікові дані користувачів;
- фішинговий web-ресурс – сторінка, що копіює зовнішній вигляд легітимного веб-сайту електронної пошти (наприклад, Gmail, Outlook, Ukr.net тощо). Містить форму для введення логіну і пароля, які, після заповнення, пересилаються на сервер зловмисника.

- авторизований користувач - отримує повідомлення з посиланням на фішингову сторінку, зазвичай у вигляді електронного листа, нібито з офіційного джерела (служба підтримки, системний адміністратор, знайомий тощо). Переходить за посиланням, не підозрюючи підміну сайту, і вводить свої облікові дані.

Перехоплення бездротового зв'язку «Bluetooth». Зловмисник перехоплює та маніпулює Bluetooth-зв'язком між пристроями. Процес створення пари зазвичай відбувається, коли два пристрої намагаються встановити з'єднання вперше або коли вони раніше не були пов'язані між собою. Процес створення пари Bluetooth зазвичай складається з наступних кроків [33]:

- ініціювання: один пристрій ініціює процес сполучення, повідомляючи про свою доступність і готовність встановити з'єднання. Цей пристрій часто називають «ініціатором» або «ведучим»;

- виявлення: пристрій-ініціатор сканує найближчі Bluetooth-пристрої в межах свого діапазону. Він визначає потенційні пристрої, з якими може встановити з'єднання;

- вибір пристрою: пристрій-ініціатор вибирає конкретний пристрій, з яким він хоче створити пару, зі списку виявлених пристроїв;

- автентифікація: після того, як пристрої ідентифіковані, вони проходять процес автентифікації, щоб переконатися, що вони є легітимними і мають право на зв'язок один з одним. Механізми автентифікації можуть включати ключі, PIN-коди або інші коди безпеки, що вводяться користувачами пристроїв;

- запит на створення пари: пристрій-ініціатор надсилає запит на створення пари вибраному пристрою, вказуючи на свій намір встановити безпечне з'єднання;

- прийняття: вибраний пристрій отримує запит на створення пари і пропонує користувачеві підтвердити створення пари. Якщо користувач погоджується, пристрої переходять до встановлення безпечного з'єднання;

- обмін ключами: під час процесу сполучення пристрої обмінюються ключами шифрування, які будуть використовуватися для шифрування та

розшифрування даних, що передаються між ними. Це гарантує, що дані залишаються конфіденційними та безпечними під час передачі;

- встановлення з'єднання: після успішного завершення процесу сполучення пристрої встановлюють з'єднання, що дозволяє їм безпечно обмінюватися даними.

Перехоплення бездротового зв'язку типу «Bluetooth» можливе через наступні атаки [33]:

- пасивне прослуховування: при пасивному сніфінгу зломисник перехоплює Bluetooth-зв'язок пасивно, не беручи активної участі у з'єднанні. Це передбачає моніторинг Bluetooth-трафіку за допомогою спеціалізованих апаратних або програмних засобів, здатних перехоплювати Bluetooth-пакети, якими обмінюються пристрої;

- активний сніфінг: атаки активного сніфінгу передбачають активну участь зломисника в Bluetooth-зв'язку, який видає себе за легітимний пристрій або маніпулює параметрами з'єднання. Впроваджуючи шкідливі пакети в потік зв'язку або змушуючи пристрої відновлювати з'єднання, зломисник може отримати доступ до конфіденційної інформації, якою обмінюються пристрої. Активні сніфінг-атаки є більш нав'язливими і потенційно можуть порушити легітимний зв'язок між пристроями;

- атаки грубої сили: полягають у систематичному переборі всіх можливих комбінацій облікових даних для автентифікації, таких як PIN-коди або паролі, для отримання несанкціонованого доступу до Bluetooth-пристрою. Цей тип атаки часто використовується при спробі створити пару з пристроєм, який має слабкі облікові дані або облікові дані за замовчуванням. Неодноразово підбираючи коди автентифікації, зломисник може зрештою отримати доступ до пристрою та його даних;

- блюджекінг: нешкідлива форма атаки через Bluetooth, коли зломисник надсилає небажані повідомлення або файли на пристрої з підтримкою Bluetooth, що знаходяться поблизу. Використовується така атака для порушення нормальної роботи пристрою або для збору інформації про пристрої, що знаходяться поруч;

- блюснарфінг: Bluetooth-атака, яка передбачає несанкціонований доступ до інформації, що зберігається на пристрої з підтримкою Bluetooth, наприклад, до контактів, електронних листів, текстових повідомлень або мультимедійних файлів. Зловмисники використовують вразливості в протоколі або реалізації Bluetooth, задля отримання доступу до даних пристрою без відома або згоди користувача;

- активне підслуховування: атака на Bluetooth, яка дозволяє зловмиснику отримати контроль над пристроєм з підтримкою Bluetooth без відома користувача. Використовуючи вразливості в протоколі Bluetooth або прошивці пристрою, зловмисники можуть віддалено виконувати команди, надсилати повідомлення, здійснювати дзвінки або отримувати доступ до конфіденційної інформації, яка знаходиться на пристрої.

- атаки підроблення: під час атак підроблення зловмисники намагаються видати себе за легітимні Bluetooth-пристрої або сервіси, щоб обманом змусити користувачів або інші пристрої під'єднатися до них. Маскуючись під надійний пристрій або службу, зловмисники можуть перехоплювати зв'язок, отримувати доступ до конфіденційних даних або вчиняти інші зловмисні акти.

- атаки на відмову в обслуговуванні (DoS-атаки): DoS-атаки націлені на доступність пристроїв або мереж з підтримкою Bluetooth, переповнюючи їх великим обсягом шкідливого трафіку або використовуючи вразливості для виведення пристроїв з ладу або заморожування. DoS-атаки можуть порушити зв'язок Bluetooth, позбавляючи пристрої можливості підключатися або викликаючи їх несправність;

- експлойти та шкідливе ПЗ Bluetooth : спеціально розроблені для пристроїв з підтримкою Bluetooth, можуть поставити під загрозу їхню безпеку та цілісність. Шкідливе ПЗ, яке використовує технологію Bluetooth, здатне заражати пристрої через поширення шкідливих файлів або додатків. Таке ПЗ активно експлуатує вразливості в системах безпеки пристроїв для отримання несанкціонованого доступу. Після цього зловмисники можуть викрасти конфіденційні дані користувачів, а також здійснювати інші шкідливі операції.

РОЗДІЛ 2

МЕТОДИ ПОПЕРЕДЖЕННЯ ВТОРГНЕНЬ ТИПУ «MAN IN THE MIDDLE»

2.1. Методики запобігання комп'ютерним атакам типу «Man In The Middle»

Коли мова йде про MITM-атаки, той факт, що їх важко виявити, є основним фактором, який пояснює, чому захист є настільки важливим. MITM-атака може залишатися непоміченою протягом тривалого часу, якщо її активно не шукати, даючи зловмиснику достатньо часу для здійснення атаки до того, як будуть вжиті відповідні заходи. Виробники мережевого обладнання та програмного забезпечення, такі, як Cisco, Eltex, Fortinet та інші, розробляють рішення для захисту від MITM-атак, пропонуючи вбудовані засоби безпеки у свої пристрої та системи. Також деякі організації, як OWASP, включають MITM-атаки у свої методики тестування на проникнення та оцінювання вразливостей. Сучасні системи також включають методи моніторингу трафіку для виявлення аномальних дій, які вказують на атаку MITM. Існує кілька стратегій, які виявилися ефективними для запобігання MITM-атакам [34]:

- а) VPN
- б) безпечні з'єднання;
- в) ARP;
- г) DNS;
- д) обстеження латентності;
- е) IDPS;
- є) статистичний метод;

ж) метод захисту, заснований на математичній моделі та структурі нейронної мережі для виявлення кібератак на інформаційно-комунікаційні системи.

VPN. Використання віртуальних приватних мереж є дуже ефективним підходом для зупинення MITM-атаки. Віртуальна приватна мережа - це розширення приватної мережі через публічну мережу (найчастіше Інтернет), що дозволяє приладам підключатися так, ніби вони підключені до приватної мережі

поверх публічної мережі. Зазвичай це супроводжується посиленими заходами безпеки і відповідним шифруванням, щоб запобігти потенційним спробам переглянути або змінити дані, що передаються, або іншим чином втрутитися в з'єднання.

VPN приховує шлях зв'язку, шифрує мережевий трафік і маскує IP-адресу. У сценарії MITM зломиснику буде неймовірно складно відстежити IP-адресу і запустити легальну атаку в результаті такого маскування.

VPN використовує 256-бітове шифрування AES, яке унеможливорює перегляд оригінального повідомлення. За допомогою цього шифрування, MITM-атака не може прослуховувати мережеве спілкування. Зашифровані дані не дають хакерам змоги прочитати повідомлення або виявити пункт призначення, який відвідує система, навіть якщо вони отримають доступ до мережі Wi-Fi. Ключі шифрування часто змінюються за допомогою технології Perfect Forward Secrecy, що не дозволяє зломисникам використовувати їх для розшифрування повідомлень. Захист від витіку IP-адреси запобігає розкриттю місцезнаходження користувача. Схема функціонування VPN наведено на рисунку 2.1:

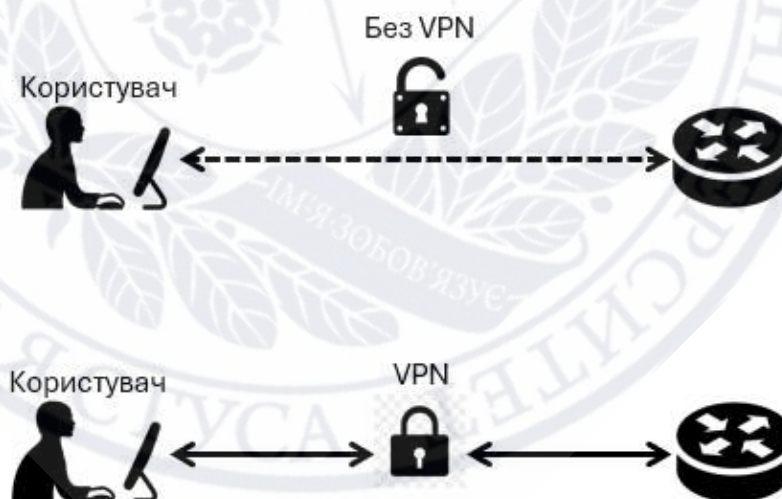


Рисунок 2.1 – Схема функціонування VPN

Безпечні з'єднання. Використання шифрування є типовою стратегією профілактики. Раціональний спосіб зупинити ці атаки - використовувати криптографічні протоколи, такі як TLS. Протокол SSL був справедливо застарілим через помилки, але TLS зайняв його місце і виявився набагато

ефективнішим у виконанні завдань шифрування даних і автентифікації. Слід зазначити, що обидва протоколи постійно модифікуються з метою виправлення помилок, розширення функціональності та впровадження методів адаптації до постійних технологічних удосконалень.

Завдяки використанню TLS, передача базується на ключовій структурі, що дозволяє використовувати криптографію з відкритим ключем для перевірки ідентичності двох або більше учасників. Дані, що передаються, захищені ключами, які створюються спеціально для кожного каналу зв'язку, що робить з'єднання безпечним.

Враховуючи, що ідентичність клієнта і сервера повністю перевірена, така взаємна автентифікація зазвичай усуває ризик MITM-атаки, оскільки виключає можливість доступу і розшифровки надісланих даних, оскільки ключі невідомі нікому іншому. Веб-сайти не повинні пропонувати альтернативи HTTP і повинні використовувати виключно HTTPS.

Впровадження HTTPS, який використовується для забезпечення безпечного з'єднання в мережі, є простою стратегією. Популярні веб-сайти та браузери попереджають користувачів, коли вони використовують незахищене з'єднання, що значно зменшило кількість MITM-атак у громадських точках доступу WI-FI. Після отримання такого попередження слід негайно розірвати WI-FI-з'єднання, щоб зменшити небезпеку.

ARP. S-ARP - це протокол, який забезпечує безпечну альтернативу ARP і може бути використаний для запобігання MITM-атакам. ARP - це протокол, який використовується комп'ютерами для зіставлення фізичної (MAC) адреси з IP-адресою в локальній мережі. Однак ARP є вразливим до MITM-атак, оскільки він не забезпечує жодного механізму автентифікації для перевірки ідентичності сторін, що спілкуються. Як наслідок, зловмисник може надсилати підроблені ARP-повідомлення, видаючи себе за легітимний мережевий пристрій, і перенаправляти мережевий трафік на власний комп'ютер, що дозволяє йому підслуховувати або модифікувати комунікацію.

S-ARP забезпечує автентифікацію і шифрування ARP-повідомлень, що робить його набагато безпечнішим, ніж ARP. У S-ARP, перш ніж клієнт надсилає ARP-запит, він спочатку автентифікується за допомогою попередньо наданого секретного ключа або цифрових сертифікатів. Вузел зв'язку перевіряє особу клієнта перед пересиланням ARP-запиту, і тільки пристрої, які мають відповідний дозвіл, можуть спілкуватися в мережі. S-ARP також використовує шифрування для захисту ARP-повідомлень від підслуховування і несанкціонованого доступу. Використовуючи S-ARP, мережеві адміністратори можуть запобігти MITM-атакам, які використовують методи підміни ARP.

S-ARP забезпечує більш безпечний метод перетворення IP-адрес в MAC-адреси, гарантуючи, що в мережі взаємодіють лише авторизовані пристрої, а також шифрування та автентифікацію зв'язку.

Другим методом профілактики є використання статичних MAC, коли певна IP-адреса прив'язується до одної MAC-адреси. Адміністратор повинен брати участь у налаштуванні IP-адреси зі статичним зв'язком MAC-адреси. Для перевірки ARP-пакетів у певній мережі використовується метод під назвою Динамічна перевірка ARP.

DNS. Використання розширення DNS DNSSEC є ще одним методом захисту від MITM. Через відсутність в DNS процедур для перевірки даних та їх походження, DNSSEC забезпечує безпеку і протидіє підміні DNS і атакам на отруєння кешу DNS. DNSSEC робить це, додаючи автентифікацію до джерела даних. DNSSEC - це розширення системи доменних імен, яке додає цифрові підписи до даних DNS для автентифікації джерела даних DNS і перевірки їхньої цілісності. Використовуючи DNSSEC, можна запобігти MITM-атакам.

DNSSEC усуває вразливість шляхом додавання цифрових підписів до даних DNS, що дозволяє перевірити дані, які надходять, що вони є автентичними і не були підроблені. Коли клієнт DNS запитує дозвіл на доменне ім'я, сервер, який надає відповідь, та підписує дані за допомогою свого приватного ключа. Потім клієнт може використовувати відкритий ключ сервера, щоб перевірити підпис і переконатися, що дані є автентичними. Використовуючи DNSSEC,

клієнти можуть бути впевнені, що дані DNS, які вони отримують, не були підроблені, і що вони зв'язуються з потрібним сервером. Це значно ускладнює зломисникам перехоплення та зміну даних DNS, ефективно запобігаючи MITM-атакам.

Однак слід зазначити, що протокол DNSSEC повинен використовуватися як на серверах, так на інших вузлах зв'язку, щоб DNSSEC був надійним методом виявлення MITM-атак і підтримував автентичність і цілісність походження даних.

Обстеження латентності. Повторюване та неочікуване переривання роботи певного сервісу є характерною ознакою MITM-атаки. Коли користувач намагається відновити з'єднання, зломисник примусово розриває його сеанс, щоб перехопити дані автентифікації. Посилання, які відрізняються від справжнього веб-сайту, що переглядаються через мережу організації, є ще однією ознакою атаки MITM. Якщо адміністратор помічає в журналах, що пристрій неодноразово намагається підключитися до адреси, наприклад, «goal.com» замість «goal.com», це означає, що відбувається активна атака типу MITM.

Атаку MITM можна виявити за допомогою аналізу затримок. Це передбачає виконання загальної, але складної задачі, наприклад, великих обчислень, необхідних для розробки хеш-функцій. Одна і та ж транзакція виконується за допомогою декількох транзакцій. Кожна з них повинна мати однаковий час відгуку. Існує ймовірність того, що зломисник втручається в цю передачу, якщо будь-яка з комунікацій займає надзвичайно тривалий час. Мітки часу в інформації заголовка TCP-пакетів використовуються для аналізування затримок. У таких ситуаціях можна застосувати виявлення несанкціонованого доступу, яке ефективно перевіряє час і затримку поточного процесу. Підвищена затримка транзакцій може бути ознакою того, що відбувається атака MITM.

IDPS. Глибока перевірка пакетів та глибока перевірка потоку також можуть бути використані для виявлення MITM-атак під час моніторингу мережі. Завдяки DPI і DFI мережеві монітори можуть отримати доступ до таких даних, як довжина пакетів. Вони можуть бути використані для виявлення аномальної

мережевої активності. Впровадження IDPS є ще одним методом виявлення MITM-атак.

У ключових вузлах мережі можуть бути розміщені спеціалізовані вузли IDPS. Для виявлення відхилень від типового потоку трафіку ці вузли використовують виявлення вторгнень на основі аномалій.

IDPS здійснює моніторинг мережевого трафіку з метою виявлення аномалій або підозрілої поведінки. У випадку MITM-атак IDPS може виявити, коли зловмисник надсилає підроблене ARP-повідомлення, видаючи себе за оригінальний мережевий пристрій і перенаправляючи трафік на свій комп'ютер. Потім IDPS може попередити мережевого адміністратора, який може вжити заходів для запобігання атаці.

Крім того, деякі рішення IDPS можуть запобігати атакам підміни ARP, використовуючи такі методи, як виявлення «отруєння ARP» або перевірка ARP-кешу. Виявлення «отруєння» ARP передбачає моніторинг ARP-кешу для виявлення будь-яких змін або невідповідностей, які можуть вказувати на атаку підміни. Перевірка ARP-кешу передбачає перевірку MAC-адреси відправника ARP-повідомлення за даними ARP-кешу, щоб переконатися, що це оригінальний пристрій.

Використовуючи рішення IDPS, яке здатне виявляти та запобігати ARP-спуффінговим MITM-атакам, мережеві адміністратори можуть забезпечити захист своєї мережі від цього типу атак. IDPS може забезпечити додатковий рівень безпеки, який доповнює інші заходи безпеки, такі як S-ARP, які також можна використовувати для запобігання MITM-атакам.

Виявлення аномальних шаблонів трафіку та проблем із затримками - це лише перший крок у виявленні MITM-атаки. Щоб встановити, чи є зібраний мережевий трафік спробою MITM, необхідно провести аналіз мережевого трафіку. Якщо атака підтверджена, необхідно визначити її джерело. Використання IDPS в поєднанні з іншими технологіями SIEM також має важливе значення.

Оскільки атака MITM спрямована на отримання несанкціонованого доступу, то для виявлення цієї атаки можуть використовуватися сигнатурні та евристичні аналізатори, що входять до складу систем виявлення атак (IDS).

Функцію IDS можна визначити як моніторинг мережевих транзакцій та виявлення зловмисних дій. Вхідні параметри для IDS можуть бути у двох формах: аномалії та сигнатури. У випадку аномалій фіксується поведінковий шаблон транзакцій, і якщо будь-яка транзакція відхиляється від звичайного шаблону, то вона може бути віднесена до підозрілої активності. У випадку з сигнатурами, кожній дії присвоюється унікальний шаблон на основі ідентифікатора, який використовується для розрізнення користувачів від шахрайських [35].

Зазвичай мережевий трафік містить великі обсяги даних, що, поряд з необхідністю аналітики в режимі реального часу, висуває жорсткі вимоги до ефективності методів пошуку та виявлення мережевих аномалій. Тому статистичний аналіз мережевого трафіку та пошук прикладних ознак, що формують його нормальну та аномальну поведінку трафіку, є важливим завданням.

Статистичний метод. Статистичний метод на основі контрольних карт CUSUM, як ефективний статистичний метод виявлення аномалій трафіку. Контрольні карти CUSUM, також називають алгоритмом кумулятивних сум (CUSUM cumulative sum) [36].

Під час моніторингу IP-адрес у вхідному трафіку метод CUSUM використовує кількість нових IP-адрес за сеанс. Припускається, що IP-адреси, які раніше не надходили, може свідчити про атаку MITM. Такий метод потребує попередньої підготовки - знайомства з організацією бази даних IP-адрес [36].

Моніторинг нових IP-адрес у вхідному трафіку складається з двох частин: навчання та виявлення. Навчання полягає у додаванні довірених IP-адрес до бази даних. На етапі виявлення здійснюється підрахунок кількості пристроїв, що вперше звернулися до IP-адреси, та отримання відхилення від середнього значення за допомогою тесту CUSUM [36].

Метод захисту, заснований на математичній моделі та структурі нейронної мережі для виявлення кібератак на інформаційно-комунікаційні системи. У роботі [37], авторами запропоновано та досліджено математичну модель та структуру нейронної мережі, здатної виявляти кібератаки. Для перевірки ефективності запропонованої математичної моделі та архітектури ШНМ, авторами [37] були проведені експериментальні дослідження, в яких розроблена мережа була протестована для виявлення кібератак. Під час перевірки ефективності використовувалась база даних, яка містила 22 типи кібератак, які поділяються на 4 основні класи: відмова в обслуговуванні (DoS), отримання несанкціонованих прав доступу незареєстрованим користувачем (R2L), несанкціоноване підвищення привілеїв (U2R) зареєстрованим користувачем та сканування портів (Probe). Кількість записів, що відповідають відсутності атаки, становила 972 781. У таблиці 2.1 наведено результати виявлення різних типів атак різними методами:

Таблиця 2.1 - Результати виявлення мережових атак [37]

Attack detection methods	DoS, %	Probe, %	R2L, %	U2R, %
Neural network	98.0	92.8	36.7	30.8
Gaussian Classifier	82.4	90.2	9.6	22.8
K –NN	97.3	87.6	6.4	29.8
Nearest Neighbour Method	97.1	88.8	3.4	2.2
Leader algorithm	97.2	83.3	1.0	6.6
Hypersphere algorithm	97.2	84.8	1.0	8.3
Fuzzy Art Map	97.0	77.2	3.7	6.1
Decision Tree	97.0	80.8	4.6	1.8

Як видно з представлених результатів, запропонований підхід показав кращі результати у виявленні кібератак. Слід зазначити, що відсоток хибних спрацьовувань, коли легітимне з'єднання приймається за атаку, становить менше 10%. Варто також зазначити, що в таблиці наведено середні результати для чотирьох класів мережових атак [37].

Запропонований підхід, що базується на застосуванні розробленої математичної моделі та архітектури ШНМ як детектора мережевих атак на інформаційно-комунікаційні системи, дозволяє підвищити рівень виявлення мережевих вторгнень на комп'ютерні системи, веб- та інтернет-ресурси. Виявлення деяких типів атак відбувається з вірогідністю 98% при незначному рівні хибних спрацьовувань. Крім того, такий підхід не вимогливий до системних ресурсів і здатний виявляти деякі невідомі типи атак на інформаційно-комунікаційні системи (ШНМ, навчений на одному типі атак, часто показує хороші результати при виявленні інших типів атак, тобто на даних, на яких навчання не проводилося) [37].

2.2 Попередження атак в бездротовій мережі Wi-Fi

В ході дослідження розглянуті різні протоколи безпеки стандартів 802.11, включаючи PMF, що використовується для захисту Wi-Fi зв'язку від атаки MITM.

Стандарт IEEE 802.11i визначає захищені мережі Wi-Fi з більш надійними рішеннями безпеки, ніж стандарт 802.11. IEEE 802.11i також відомий як Robust Security Network. Для забезпечення захисту на каналному рівні для зв'язку Wi-Fi за стандартом 802.11i, Wi-Fi Alliance підтримує три сертифікати безпеки, а саме Wi-Fi Protected Access, Wi-Fi Protected Access II і Wi-Fi Protected Access III. Як протокол шифрування або конфіденційності даних використовується протокол цілісності тимчасового ключа Temporal Key Integrity Protocol протоколи WPA2 і WPA3 вимагають використання розширеного стандарту шифрування Advanced Encryption Standard. Як протокол цілісності даних, WPA використовує перевірку цілісності повідомлень, відому як алгоритм Майкла, WPA2 і WPA3 вимагають використання протоколу CBC-MAC з лічильником для персональних мереж, а новий протокол Галуа з лічильником посилює безпеку WPA3 від атаки MITM для корпоративних мереж [38].

Згідно зі стандартом 802.11i, коли клієнт підключається до маршрутизатора або точки доступу в базовому наборі послуг або WLAN, він проходить чотири етапи [38]:

- 1) виявлення мережі;
- 2) аутентифікація;
- 3) асоціація;
- 4) 4-сторонній механізм «рукоштовування», як показано на рис. 2.2:

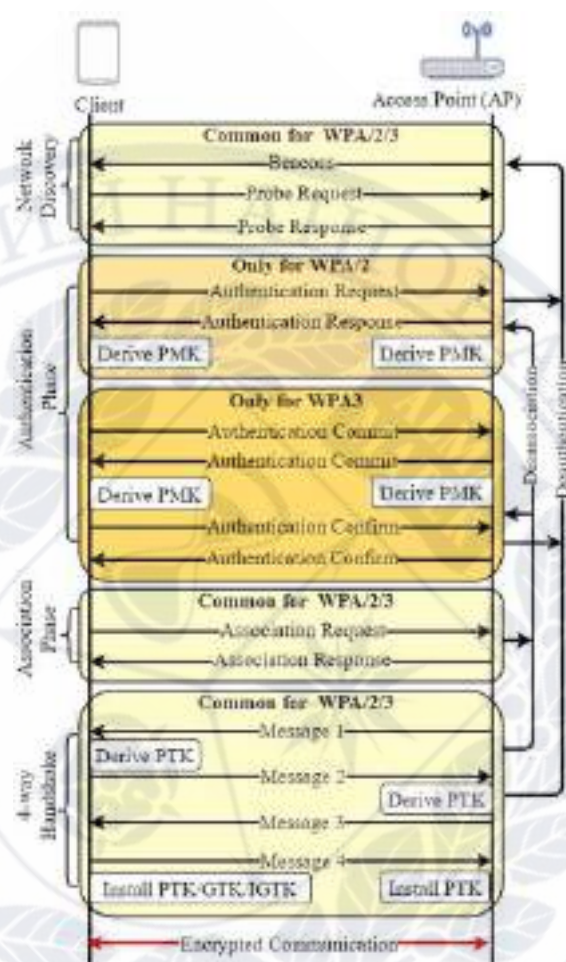


Рисунок 2.2 - Узагальнене встановлення з'єднання між точкою доступу та клієнтом [38]

У бездротових локальних мережах точки доступу показують свою присутність у мережі, періодично транслюючи маячки. Маячок містить ідентифікатор набору послуг або ім'я мережі, MAC-адресу, канал та інші можливості точки доступу. Далі клієнтський пристрій сканує і складає список доступних мереж, щоб користувач міг вибрати відповідну мережу і вручну ввести парольну фразу або попередньо наданий ключ, який вже був налаштований в точці доступу. Важливо відзначити, що ця парольна фраза зберігається або кешується в мікросхемі Wi-Fi клієнта і ніколи не передається або

не обмінюється в жодному з кадрів, гарантуючи захист від атаки MITM. З обраним SSID клієнт надсилає кадр запит, щоб перевірити, чи доступна певна мережа. У відповідь на це точка доступу надсилає відповідь зондування, підтверджуючи доступність SSID. На цьому виявлення мережі завершується. Кроки цього етапу є спільними для WPA, WPA2 і WPA3 [38].

Під час автентифікації точка доступу перевіряє ідентичність клієнта (MAC-адресу) і реєструє її у своєму кеші. Як показано на рис. 2.1, фаза автентифікації має різні етапи залежно від версії протоколу безпеки. У WPA або WPA2 клієнт і точка доступу обмінюються відкритими запитами на автентифікацію і відповідями. Після успішної автентифікації, попередній головний ключ генерується з PSK з обох сторін. З іншого боку, протокол WPA3 виконує нове рукостискання (так звана одночасна автентифікація рівних), обмінюючись чотирма кадрами автентифікації. Перед цим клієнт і точка доступу генерують секретний елемент, відомий як елемент пароля, і два секретних значення. Під час автентифікації (повідомлень-фіксаторів) клієнт і точка доступу узгоджують РМК за допомогою методу обміну ключами еліптичної кривої Діффі-Хеллмана. В останніх двох автентифікаційних повідомленнях вони підтверджують, що обидва обробляли один і той самий РМК. Таким чином, РМК обчислюється за допомогою відповідного протоколу безпеки і кешується на обох пристроях, що захищає сеанс від атаки MITM. Згенерований РМК буде використовуватися в 4-сторонньому рукостисканні. Після аутентифікації, деаутентифікація від клієнта або точки доступу призведуть до відключення від мережі [38].

Як тільки автентифікація завершується, клієнт готується до асоціювання з точкою доступу, надсилаючи запит для узгодження необхідних наборів шифрів, таких як TKIP/CCMP/GCMP. Під час асоціації клієнта точка доступу зберігає ідентифікатор асоціації та надсилає відповідь про асоціацію. Клієнт може бути автентифікований у багатьох мережах, але може бути асоційований лише з однією мережею за один раз. З кешованим РМК клієнту буде дозволено знову приєднатися до вже асоційованої точки доступу навіть після виходу з мережі, або

клієнт може бути швидко перепідключений до точки доступу. Таким чином, користувачеві або клієнту не потрібно знову вводити парольну фразу до мережі Wi-Fi, оскільки точка доступу зберігає свій стан або асоціацію безпеки, тим самим зловмисник не зможе перехопити пароль під час атаки MITM. Ця процедура однакова для WPA, WPA2 і WPA3. Як і при деаутентифікації, на цьому етапі може відбутися роз'єднання, тобто від'єднання клієнта. 4-стороннє рукостискання починається після успішної асоціації [38].

Механізм 4-стороннього «рукостискання» (однаковий у протоколах WPA, WPA2 і WPA3) передбачає обмін 4-ма EAPOL. Під час цього «рукостискання» точка доступу і клієнт отримують парний перехідний ключ, також відомий як сеансовий ключ, який потім використовується для шифрування фактичного зв'язку між ними. Для отримання РТК використовується РМК з іншими параметрами:

- АМАС - MAC-адреса точки доступу, СМАС - MAC-адреса клієнта,
- AN - випадкове число точки доступу;
- CN - випадкове число клієнта;
- RC - лічильник повторів,
- PRF - псевдовипадкова функція.

МІС містить код цілісності повідомлення, створений для вмісту в дужках з похідним РТК, щоб точка доступу або клієнт Wi-Fi були спроможні перевірити, чи є це повідомлення пошкодженим чи ні. GTK генерується незалежно на кожній точці доступу і є однаковим для всіх підключених клієнтів. Аналогічно, IGTK генерується, якщо увімкнено PMF. Процес обміну відповідними 4-сторонніми повідомленнями «рукостискання» здійснюються [38], як показано на рисунку 2.3:

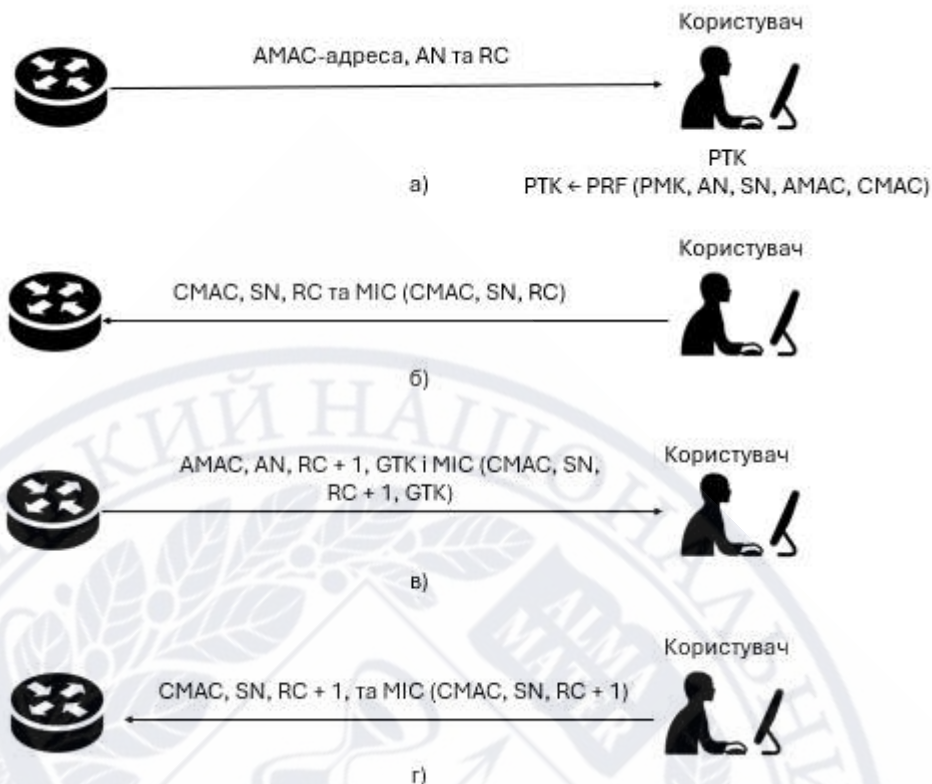


Рисунок 2.3 – Процес обміну 4-сторонніми повідомленнями «рукоштовання»:

а - точка доступу надсилає запит клієнту; б - клієнт надсилає параметри до точки доступу; в - точка доступу перевіряє МІС і виводить РТК; г - точка доступу та клієнт встановлюють РТК і GTK.

За допомогою 4-стороннього «рукоштовання» точка доступу і клієнт завершують роботу і залишаються на зв'язку. На цьому етапі з різних причин може відбутись деаутентифікація або роз'єднання. Після того, як кінцеві пристрої встановлять ключі безпеки, парний обмін даними між точкою доступу і клієнтом буде зашифровано на каналному рівні сеансовим ключем РТК з використанням узгоджених шифрів. Точка доступу використовує GTK для шифрування широкомовних або багатоадресних кадрів для зв'язку з кожним асоційованим клієнтом.

Що стосується WPA і WPA2, то головна проблема полягає в тому, що вони вразливі до атак грубої сили або атак за словником, які допомагають зломисникам отримати ключі безпеки і потенційно розшифрувати раніше

зашифровані сеанси зв'язку. Це відбувається тому, що згенерований РМК однаковий для всіх клієнтів. Однак WPA3 вирішує цю проблему за допомогою рукописання «Dragonfly», яке не тільки збільшує ентропію РМК, але й забезпечує надійну автентифікацію/обмін ключами за допомогою криптографії еліптичних кривих і стійкого шифрування за допомогою AES-GCM. Таким чином, офлайн-атаки за словником і компрометація попередніх сесій (пряма секретність) запобігаються, оскільки похідний РМК не залежить від PSK, і кожен клієнт має свій власний РМК [38].

З іншого боку, хоча кадри даних (фактичний зв'язок між кінцевими пристроями) захищені за допомогою протоколів безпеки, всі кадри управління на етапах виявлення мережі, автентифікації та асоціації залишаються незахищеними, оскільки вони обмінюються перед узгодженням ключів безпеки. Тому зломисники можуть підробляти такі кадри, видавати себе за точку доступу, встановлюючи несанкціоновані пристрої, і організовувати кілька MITM атак. Наприклад, підмінивши MAC-адресу точки доступу, зломисник може надсилати клієнту кадри деаутентифікації або роз'єднання. Аналогічно, підмінивши клієнта, він може відправити на точку доступу кадр повторної асоціації. В обох випадках клієнт від'єднується від мережі, що призводить до DoS-атаки. Щоб протистояти цим проблемам, був запроваджений стандарт PMF.

Стандарт PMF або IEEE 802.11w був ратифікований у 2009 році і став частиною стандарту 802.11i у 2014 році. Хоча PMF існує вже довший час, його прийняття на ринку було відносно низьким, оскільки він був додатковою функцією для існуючих сертифікатів WPA2. З 2018 року WFA зробила PMF обов'язковою вимогою безпеки для нових сертифікатів як WPA2, так і WPA3. Коли PMF увімкнено, він захищає деякі специфічні надійні фрейми управління, такі як відокремлення, деаутентифікація та фрейми дій (наприклад, управління спектром). Дві основні поправки до PMF полягають у наступному:

- 1) код цілісності повідомлення генерується з використанням спільного секретного IGTK, який шифрує широкомовні або багатоадресні стійкі кадри управління (наприклад, деаутентифікація) для забезпечення автентифікації та

захисту від повторного відтворення. Обчислення МІС виконується за допомогою протоколу цілісності широкомовної/багатоадресної передачі.

2) захист від руйнування асоціації безпеки (SA) додано як захист від підміни асоціацій, щоб запобігти атакам підміни, які руйнують існуючі асоціації клієнтів. Це досягається за допомогою процедури SA Query, яка забезпечує захист від несанкціонованих точок доступу або клієнтів. Це криптографічно захищене пробне повідомлення, ініційоване будь-якою стороною для перевірки автентичності запитів на (роз)асоціювання.

PMF буде ефективною лише тоді, коли і точка доступу, і клієнт підтримують її, або кожен пристрій у бездротовій локальній мережі підтримує її. Однак, більшість доступних в даний час пристроїв WPA2 не підтримують цю функцію. Хоча деякі маршрутизатори Cisco підтримують PMF, вона рідко використовується в інфраструктурних мережах через величезні проблеми з сумісністю. Він також майже не використовується в пристроях Інтернету речей через ресурсомісткі криптографічні операції. З іншого боку, хоча PMF забезпечує автентичність походження даних для певних надійних фреймів управління, таких як фрейми деаутентифікації або роз'єднання, він не захищає інші попередньо автентифіковані фрейми управління, такі як маяки, відповіді на зондування, автентифікацію або фрейми (ре)асоціації. Цей фундаментальний недолік все ще кидає виклик безпеці не лише пристроїв WPA2, але й WPA3, і дозволяє зломисникам проводити атаки MITM.

2.3 Попередження атак в бездротовій мережі Bluetooth

Створення пари передбачає перевірку автентичності двох пристроїв, шифрування з'єднання за допомогою короткострокового ключа а потім розподіл довгострокових ключів, які використовуються для шифрування. Розподіл довгострокових ключів зберігається для швидшого відновлення з'єднання в майбутньому, що називається зв'язуванням [39]. Для забезпечення безпеки бездротової технології Bluetooth потрібен захист від пасивних атак прослуховування та захист від атак MITM [40]. Для захисту від пасивного

прослуховування, трафік між двома пристроями повинен бути зашифрований. У версіях Bluetooth (4.0 і 4.1) шифрування було вразливим до атаки грубої сили. Тимчасові ключі, які використовуються для подальшої генерації та розповсюдження ключів, недостатньо стійкі. Ця проблема була виправлена в Bluetooth Secure Connection. Хоча пасивному прослуховуванню можна запобігти за допомогою надійного шифрування каналного рівня, шифрування каналного рівня може бути вразливим до MITM-атак. Деякі моделі асоціацій можуть ефективно захищати від MITM-атак [40]. Моделі асоціацій - це механізми, які визначають, як два пристрої можуть автентифікувати один одного і безпечно обмінюватися даними один з одним. У Bluetooth моделі асоціації використовуються під час процесу сполучення. Наразі існує чотири моделі асоціації. Це Numeric Comparison, Just Works, Passkey Entry та Out-Of-Band [40]. Який тип моделі асоціації буде використано, залежить від можливостей вводу/виводу обох пристроїв. В таблиці 2.2 наведено відображення моделей асоціацій відповідно до різних можливостей вводу/виводу:

Таблиця 2.2 - Модель асоціації відповідно до різних можливостей вводу/виводу [40]

Моделей асоціацій	Дисплей	Дисплей, Так/Ні	Клавіатура	Без введення/ виведення	Клавіатура, дисплей
Дисплей	Just Works	Just Works	Passkey Entry	Just Works	Just Works
Тільки клавіатура	Passkey Entry	Passkey Entry	Passkey Entry	Just Works	Passkey Entry
Без введення/ виведення	Just Works	Just Works	Just Works	Just Works	Just Works
Клавіатура, дисплей	Passkey Entry	Numeric Comparison	Passkey Entry	Just Works	Numeric Comparison

Рядки та стовпчики в табл. 2.2 представляють можливості вводу/виводу двох пристроїв відповідно. Дисплей означає, що у пристрою наявний дисплей. А клавіатура означає, чи є на пристрої клавіатура. «Так/ні» це може бути клавіатура, яка може бути використана користувачем для підтвердження або відхилення значення, що відображається на екрані.

Моделі асоціацій відповідно до таблиці 2.1 [40]:

- *Numeric Comparison (тільки для захищених з'єднань)*. Застосовується, коли обидва пристрої мають можливість відображати та підтверджувати інформацію. Наприклад, два смартфони хочуть створити пару один з одним. У цій моделі об'єднання, обидва пристрої відображають на екрані шестизначне число. Якщо ці два числа збігаються, вони автентифіковані належним чином. Числа генеруються з відкритого ключа кожного пристрою, 8-бітного коду за допомогою функції AES-CMAC. Під час MITM-атаки відкриті ключі обох пристроїв не збігаються. Таким чином, ефективно уникати MITM-атак;

- *Just Works*. Призначений для сценаріїв, коли принаймні один з пристроїв не має можливості вводу/виводу. Наприклад, миша і ноутбук хочуть об'єднатися в пару один з одним. У цій моделі об'єднання вони використовують той самий протокол, що й у числовому порівнянні, але шестизначне число ніколи не з'являється. Оскільки «Just Works» не надає користувачеві інтерфейсу, щоб дізнатися, чи збігається число, він вразливий до MITM-атак;

- *Passkey Entry*. Функція призначена для ситуацій, коли один пристрій підтримує введення шестизначного номера (пароля), але не має можливості виводити його на екран. Наприклад, клавіатура і ноутбук хочуть об'єднатися в пару один з одним. У цій моделі об'єднання, остаточне значення підтвердження генерується з відкритих ключів обох пристроїв і однієї частини Passkey за допомогою функції AES-CMAC.

У застарілих парах (відомих як 4.1 і 4.0), «Just Works» і «Passkey Entry» уразливі до пасивного підслуховування, оскільки вони не використовують еліптичну криву Діффі-Хеллмана, як це робить Secure Simple Pairing;

- *Out-of-Band*. Режим призначений для сценаріїв, коли обидва пристрої мають додатковий канал зв'язку WLAN або NFC, наприклад, два телефони, оснащені NFC. У цій моделі об'єднання можливість захисту від атак MITM або пасивного прослуховування залежить від типу каналу зв'язку, що використовується, та рівня безпеки цього каналу зв'язку.

Захист від атак, пов'язаних з Bluetooth, вимагає багаторівневого підходу, який поєднує технічні засоби контролю, найкращі практики безпеки та обізнаність користувачів. Ось деякі ключові засоби захисту для зменшення ризиків, пов'язаних з атаками через Bluetooth [40-41]:

- *автентифікація та шифрування*. Bluetooth Secure Simple Pairing (SSP) гарантує, що тільки авторизовані пристрої можуть підключатися один до одного і що дані, передані через з'єднання Bluetooth, залишаються конфіденційними та безпечними;

- *оновлення мікропрограми та програмного забезпечення*. Регулярне оновлення прошивки та програмного забезпечення пристрою, виправить відомі вразливості та забезпечить захист пристроїв від найновіших загроз безпеці. Це включає оновлення операційних систем, драйверів Bluetooth та прикладного програмного забезпечення на всіх пристроях з підтримкою Bluetooth;

- *вимикання невикористовуваних функцій Bluetooth*. Вмикання непотрібних служб або функцій Bluetooth, які не потрібні для роботи пристрою. Це зменшує поверхню атаки та мінімізує ризик;

- *використання надійного PIN-коду та паролю*. Під час сполучення Bluetooth-пристроїв використання надійного та унікального персонального ідентифікаційного номеру (PIN-коди) або паролю, допоможе запобігти атакам грубої сили;

- *виявлення сніфінгу Bluetooth*. Інструменти виявлення прослуховування Bluetooth для моніторингу підозрілої активності та спроб несанкціонованого доступу, можуть допомогти виявити та попередити адміністраторів про потенційні MITM атаки в режимі реального часу, що дасть змогу швидко відреагувати та усунути загрозу;

- *сегментація мережі*. Розділення пристроїв з підтримкою Bluetooth на окремі мережеві зони або віртуальні локальні мережі, для того щоб ізолювати їх від критично важливих систем і конфіденційних даних. Це допомагає стримувати потенційні Bluetooth-атаки та запобігає несанкціонованому доступу злоумисників до важливих ресурсів.

РОЗДІЛ 3

ПРАКТИЧНА РЕАЛІЗАЦІЯ ПОПЕРЕДЖЕННЯ АТАК ТИПУ «MAN IN THE MIDDLE», ЩО ЗДІЙСНЮЮТЬСЯ З ВИКОРИСТАННЯМ ГЕНЕРАТИВНОГО ШТУЧНОГО ІНТЕЛЕКТУ

3.1 Налаштування існуючого засобу попередження атаки

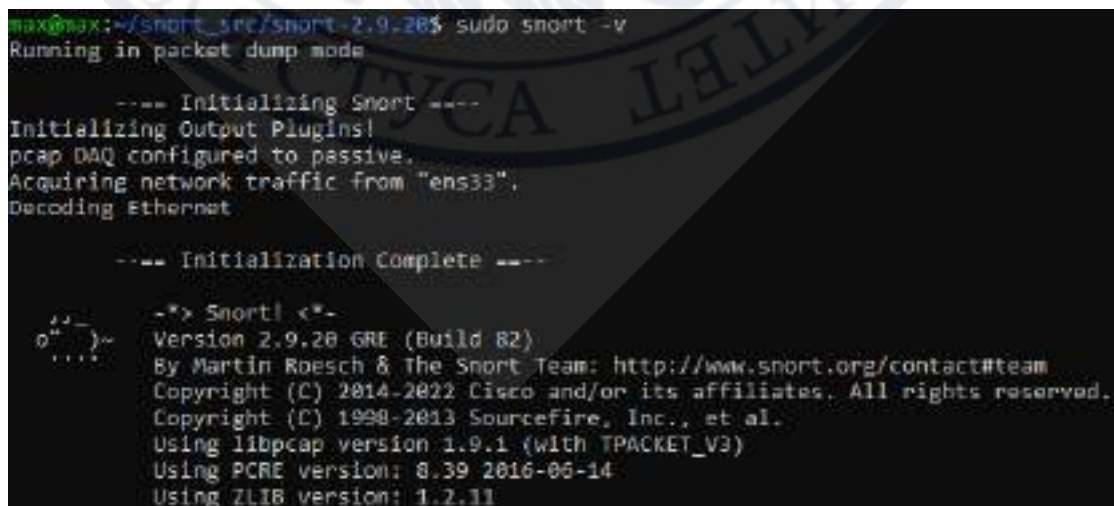
На сьогоднішній час, найпростішим та ефективним засобом для попередження атак типу «MITM» є IPS Snort. Snort – це система запобігання вторгненням з відкритим вихідним кодом. Snort IPS використовує низку правил, які допомагають визначити зловмисну мережеву активність, і на основі цих правил знаходить пакети, які їм відповідають, та генерує сповіщення для користувачів [43].

Snort має три основні застосування [43]:

- сніфер пакетів, подібний до tcpdump;
- логгер пакетів;
- корисний для налагодження мережевого трафіку, або як повноцінна система попередження мережевим вторгненням. Snort можна завантажити і налаштувати як для особистого, так і для корпоративного використання.

Для проведення експерименту буде використовуватись дві віртуальні машини з Ubuntu Server: перша із встановленим Snort, а інша буде використовуватись для атаки MITM, а саме ARP-спуффінгу.

На рисунку 3.1 наведено встановлена версія snort командою «sudo snort -v»:



```

max@max:~/snort_src/snort-2.9.28$ sudo snort -v
Running in packet dump mode

---- Initializing Snort ----
Initializing Output Plugins!
pcap DAQ configured to passive.
Acquiring network traffic from "ens33".
Decoding Ethernet

---- Initialization Complete ----

~_~
o" }~ -*) Snort! <*-
  ""  ~ Version 2.9.28 GRE (Build 82)
      ~ By Martin Roesch & The Snort Team: http://www.snort.org/contact#team
      ~ Copyright (C) 2014-2022 Cisco and/or its affiliates. All rights reserved.
      ~ Copyright (C) 1998-2013 Sourcefire, Inc., et al.
      ~ Using libpcap version 1.9.1 (with TPACKET_V3)
      ~ Using PCRE version: 8.39 2016-06-14
      ~ Using ZLIB version: 1.2.11
  
```

Рисунок 3.1 – Встановлена версія Snort


```

mac@mac: ~/snort_src/snort-2.9.28
$ ./configure
Version 2.9.28 GRE (Build 82)
By Martin Roesch & The Snort Team: http://www.snort.org/contact#team
Copyright (C) 2014-2022 Cisco and/or its affiliates. All rights reserved.
Copyright (C) 1998-2013 Sourcefire, Inc., et al.
Using libpcap version 1.9.1 (with TPACKET_V3)
Using PCRE version: 8.39 2010-06-14
Using ZLIB version: 1.2.11

Rules Engine: SF_SNORT DETECTION ENGINE version 3.2 <Build 1>
Preprocessor Object: SF_SSH Version 1.1 <Build 3>
Preprocessor Object: SF_DNP3 Version 1.1 <Build 1>
Preprocessor Object: SF_MODBUS Version 1.1 <Build 1>
Preprocessor Object: SF_SIP Version 1.1 <Build 1>
Preprocessor Object: SF_IMAP Version 1.0 <Build 1>
Preprocessor Object: SF_DCEIPC2 Version 1.0 <Build 3>
Preprocessor Object: SF_GTP Version 1.1 <Build 1>
Preprocessor Object: SF_S7COMMPPLUS Version 1.0 <Build 1>
Preprocessor Object: appid Version 1.1 <Build 5>
Preprocessor Object: SF_SSLPP Version 1.1 <Build 4>
Preprocessor Object: SF_DNS Version 1.1 <Build 4>
Preprocessor Object: SF_SMTP Version 1.1 <Build 9>
Preprocessor Object: SF_FTPTELNET Version 1.2 <Build 13>
Preprocessor Object: SF_POP Version 1.0 <Build 1>
Preprocessor Object: SF_SDF Version 1.1 <Build 1>
Preprocessor Object: SF_REPUTATION Version 1.1 <Build 1>

Total snort Fixed Memory Cost - MaxBss:46292
Snort successfully validated the configuration!
Snort exiting
mac@mac: ~/snort_src/snort-2.9.28$

```

Рисунок 3.4 – Успішний кінець перевірки

За допомогою штучного інтелекту, був згенерований сценарій, як провести атаку ARP-спуффінгу, результат якого наведений на рисунку 3.5:

Кроки:

1. Встановити `arp spoof` (якщо немає):


```
bash
sudo apt install dsniff
```
2. Запустити атаку ARP-спуффінгу:


```
bash
sudo arpspoof -i eth0 -t 192.168.0.5 192.168.0.1
```

Де:

- `eth0` – ваша мережева карта (перевірити можна `ip a`).
- `192.168.0.5` – жертва.
- `192.168.0.1` – шлюз.

Рисунок 3.5 – Згенерована атака ARP-спуффінгу

За допомогою команди «`sudo arpspoof -i ens33 -t 192.168.137.134 192.168.137.2`», буде здійснено атаку ARP-спуффінгу, на першу віртуальну машину, на якій запущений Snort в режимі моніторингу за допомогою команди

«sudo snort -A console -q -u snort -g snort -c /etc/snort/snort.conf -i ens33».

Результати наведено на рисунках 3.6-3.8:

```

mac@mac: ~/snort_src/snort-2.9.20$
****
By Martin Roesch & The Snort Team: http://www.snort.org/contact#team
Copyright (C) 2014-2012 Cisco and/or its affiliates. All rights reserved.
Copyright (C) 1998-2013 Sourcefire, Inc., et al.
Using libpcap version 1.9.1 (with TPACKET_V3)
Using PCRE version: 8.39 2016-06-14
Using ZLIB version: 1.2.11

Rules Engine: SF_SNORT DETECTION ENGINE Version 3.2 <Build 1>
Preprocessor Object: SF_SSH Version 1.1 <Build 3>
Preprocessor Object: SF_DNP3 Version 1.1 <Build 1>
Preprocessor Object: SF_MQDBUS Version 1.1 <Build 1>
Preprocessor Object: SF_SIP Version 1.1 <Build 1>
Preprocessor Object: SF_IMAP Version 1.0 <Build 1>
Preprocessor Object: SF_DCERPC2 Version 1.0 <Build 3>
Preprocessor Object: SF_GTP Version 1.1 <Build 1>
Preprocessor Object: SF_S7COMMPLUS Version 1.0 <Build 1>
Preprocessor Object: appid Version 1.1 <Build 5>
Preprocessor Object: SF_SSLPP Version 1.1 <Build 4>
Preprocessor Object: SF_DNS Version 1.1 <Build 4>
Preprocessor Object: SF_SMTP Version 1.1 <Build 9>
Preprocessor Object: SF_FTPTELNET Version 1.2 <Build 13>
Preprocessor Object: SF_PDP Version 1.0 <Build 1>
Preprocessor Object: SF_SDF Version 1.1 <Build 1>
Preprocessor Object: SF_REPUTATION Version 1.1 <Build 1>

Total snort Fixed Memory Cost - MaxRss:46292
Snort successfully validated the configuration!
Snort exiting
mac@mac:~/snort_src/snort-2.9.20$ sudo snort -A console -q -u snort -g snort -c /etc/snort/snort.conf -i ens33

```

Рисунок 3.6 – Запущений Snort в режимі сканування

```

mac@ubuntu:~$
* Introducing Expanded Security Maintenance for Applications.
  Receive updates to over 25,000 software packages with your
  Ubuntu Pro subscription. Free for personal use.

  https://ubuntu.com/pro

Expanded Security Maintenance for Applications is not enabled.

0 updates can be applied immediately.

Enable ESM Apps to receive additional future security updates.
See https://ubuntu.com/esm or run: sudo pro status

The list of available updates is more than 3 week old.
To check for new updates run: sudo apt update
New release '22.04.5 LTS' available.
Run 'do-release-upgrade' to upgrade to it.

Last login: Wed Apr 30 06:48:55 2025
mac@ubuntu:~$ sudo arpspoof -i ens33 -t 192.168.137.134 192.168.137.2
[sudo] password for max:
Sorry, try again.
[sudo] password for max:
0:c:29:e3:62:a8 0:c:29:88:ef:33 0006 42: arp reply 102.168.137.1 is-at 0:c:29:e3:62:a8
0:c:29:e3:62:a8 0:c:29:88:ef:33 0006 42: arp reply 192.168.137.2 is-at 0:c:29:e3:62:a8
0:c:29:e3:62:a8 0:c:29:88:ef:33 0006 42: arp reply 192.168.137.2 is-at 0:c:29:e3:62:a8
0:c:29:e3:62:a8 0:c:29:88:ef:33 0006 42: arp reply 192.168.137.2 is-at 0:c:29:e3:62:a8

```

Рисунок 3.7 – Початок атаки ARP-спуфінгу

```

mac@mac: ~/snort_src/snort-2.9.20$
Rules Engine: SF_SMORT DETECTION ENGINE Version 3.2 <Build 1>
Preprocessor Object: SF_SSH Version 1.1 <Build 3>
Preprocessor Object: SF_DNP3 Version 1.1 <Build 1>
Preprocessor Object: SF_MODBUS Version 1.1 <Build 1>
Preprocessor Object: SF_SIP Version 1.1 <Build 1>
Preprocessor Object: SF_IMAP Version 1.0 <Build 1>
Preprocessor Object: SF_DCERPC Version 1.0 <Build 3>
Preprocessor Object: SF_GTP Version 1.1 <Build 1>
Preprocessor Object: SF_S7COMMPLUS Version 1.0 <Build 1>
Preprocessor Object: appid Version 1.1 <Build 5>
Preprocessor Object: SF_SSLPP Version 1.1 <Build 4>
Preprocessor Object: SF_DNS Version 1.1 <Build 4>
Preprocessor Object: SF_SMTP Version 1.1 <Build 9>
Preprocessor Object: SF_FTPTELNET Version 1.2 <Build 13>
Preprocessor Object: SF_POP Version 1.0 <Build 1>
Preprocessor Object: SF_SDF Version 1.1 <Build 1>
Preprocessor Object: SF_REPUTATION Version 1.1 <Build 1>

Total snort Fixed Memory Cost - MaxRss:46292
Snort successfully validated the configuration!
Snort exiting
mac@mac: ~/snort_src/snort-2.9.20$ sudo snort -A console -q -u snort -g snort -c /etc/snort/snort.conf -i ens33
04/38-06:54:20.026810  [**] [112:2:1] (spp_arpspoof) Ethernet/ARP Mismatch request for Source [**]
04/38-06:54:21.027438  [**] [112:2:1] (spp_arpspoof) Ethernet/ARP Mismatch request for Source [**]
04/38-06:54:22.028215  [**] [112:2:1] (spp_arpspoof) Ethernet/ARP Mismatch request for Source [**]
04/38-06:54:23.029130  [**] [112:2:1] (spp_arpspoof) Ethernet/ARP Mismatch request for Source [**]
04/38-06:54:24.030072  [**] [112:2:1] (spp_arpspoof) Ethernet/ARP Mismatch request for Source [**]
^C*** Caught Int-Signal
mac@mac: ~/snort_src/snort-2.9.20$

```

Рисунок 3.8 – Успішне попередження за допомогою Snort атаки ARP-спуфінгу

Проведений експеримент із використанням системи виявлення вторгнень Snort продемонстрував її ефективність у типових сценаріях моніторингу мережевого трафіку. Проте, незважаючи на позитивні результати, слід відзначити низку критичних обмежень, що можуть становити потенційну загрозу для інформаційної безпеки.

Snort є системою з закритим ядром, що не надає користувачам можливості вносити зміни у внутрішні механізми виявлення атак. У результаті, ефективність роботи цієї системи повністю залежить від розробників і частоти виходу оновлень сигнатур та функціоналу. Якщо зловмисник знайде спосіб обійти поточний захист, адміністратор інформаційної безпеки не має можливості оперативно адаптувати систему до нових загроз, оскільки змінити механізми виявлення самостійно неможливо.

Крім того, існують додаткові ризики, пов'язані з людським фактором та інфраструктурними обмеженнями. Наприклад, у разі відсутності стабільного доступу до Інтернету або через банальну неуважність, оновлення системи можуть бути несвоєчасно встановлені, що створює вразливість.

З огляду на вищезазначені недоліки, постає необхідність розробки власної програми виявлення загроз, яка забезпечує рівень функціональності, не нижчий

за Snort, але при цьому залишається відкритою до змін. Такий підхід дозволить оперативно реагувати на нові загрози та мінімізувати залежність від сторонніх розробників, що є особливо важливим у середовищах з обмеженим доступом (неможливість швидкого встановлення онлайн оновлень) або підвищеними вимогами до гнучкості засобів захисту.

3.2. Розробка програмного коду для попередження атак «MITM»

Для розробки програмного коду та проведення практичного експерименту з його апробації використано:

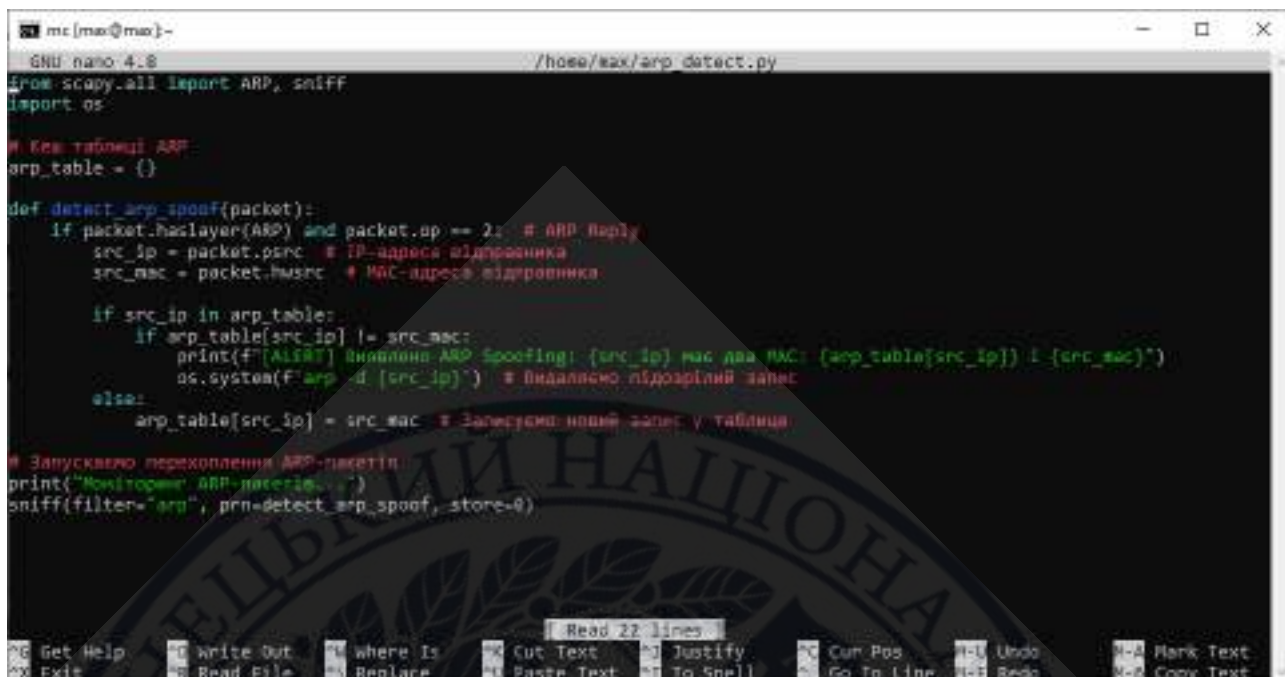
- мову програмування Python;
- дві віртуальні машини з Ubuntu Server.

Python - це інтерпретована, об'єктно-орієнтована мова програмування високого рівня з динамічною семантикою. Високорівневі вбудовані структури даних у поєднанні з динамічною типізацією та динамічним зв'язуванням роблять її дуже ефективною для швидкої розробки додатків, а також для використання в якості мови сценаріїв або для з'єднання існуючих компонентів між собою. Простий, легкий для вивчення синтаксис Python робить акцент на читабельності, а отже, зменшує витрати на підтримку програм. Python підтримує модулі та пакети, що заохочує модульність програм та повторне використання коду. Інтерпретатор Python та велика стандартна бібліотека доступні у вигляді вихідних кодів або двійкових файлів безкоштовно для всіх основних платформ і можуть вільно розповсюджуватися [44].

Для розробки використано бібліотеку Scapy.

Scapy - це інтерактивна бібліотека для маніпулювання пакетами, написана на мові Python. Scapy вміє підробляти або декодувати пакети широкого спектру протоколів, відправляти їх по дроту, перехоплювати, зіставляти запити і відповіді тощо [44].

Для встановлення Scapy використовуються команда «`pip3 install scapy`». Код програми наведений на рисунку 3.9 (лістинг програмного коду наведений в додатку Б):



```

mc [mac@mac] -
GNU nano 4.1.8 /home/mx/arp_detect.py
from scapy.all import ARP, sniff
import os

# Кеш таблиці ARP
arp_table = {}

def detect_arp_spoof(packet):
    if packet.haslayer(ARP) and packet.op == 2: # ARP Reply
        src_ip = packet.psrc # IP-адреса відправника
        src_mac = packet.hwsrc # MAC-адреса відправника

        if src_ip in arp_table:
            if arp_table[src_ip] != src_mac:
                print(f"[ALERT] виявлено ARP Spoofing: {src_ip} має два MAC: {arp_table[src_ip]} і {src_mac}")
                os.system(f"arp -d {src_ip}") # Видалимо підозрілий запис
            else:
                arp_table[src_ip] = src_mac # Записемо новий запис у таблицю
        else:
            arp_table[src_ip] = src_mac # Записемо новий запис у таблицю

# Запускаємо перехоплення ARP-пакетів
print("Моніторинг ARP-пакетів...")
sniff(filter="arp", prn=detect_arp_spoof, store=0)

```

Рисунок 3.9 – Програма для виявлення атак ARP-спуфінгу

Принцип роботи програми полягає в наступному:

- імпортується функціональність з scapy.all, зокрема класи ARP для роботи з ARP-пакетами та sniff для їх перехоплення. Також використовується модуль os для взаємодії з системними командами;
- створюється словник arp_table, який зберігає відповідність між IP-адресами та MAC-адресами, що були виявлені у мережі;
- функція detect_arp_spoof викликається кожного разу, коли перехоплюється ARP-пакет. Вона перевіряє: чи є пакет відповіддю (ARP Reply); чи IP-адреса вже була зареєстрована з іншим MAC-адресом. Якщо виявлено невідповідність (одна IP-адреса з різними MAC-адресами), виводиться попередження про можливу атаку типу ARP Spoofing. У відповідь на це також виконується системна команда «arp -d», яка видаляє підозрілий запис з кешу ARP;
- функція sniff запускає безперервне прослуховування ARP-пакетів. Кожен перехоплений пакет передається на обробку до функції detect_arp_spoof.

Результат проведення експерименту наведено на рисунках 3.10-3.12:

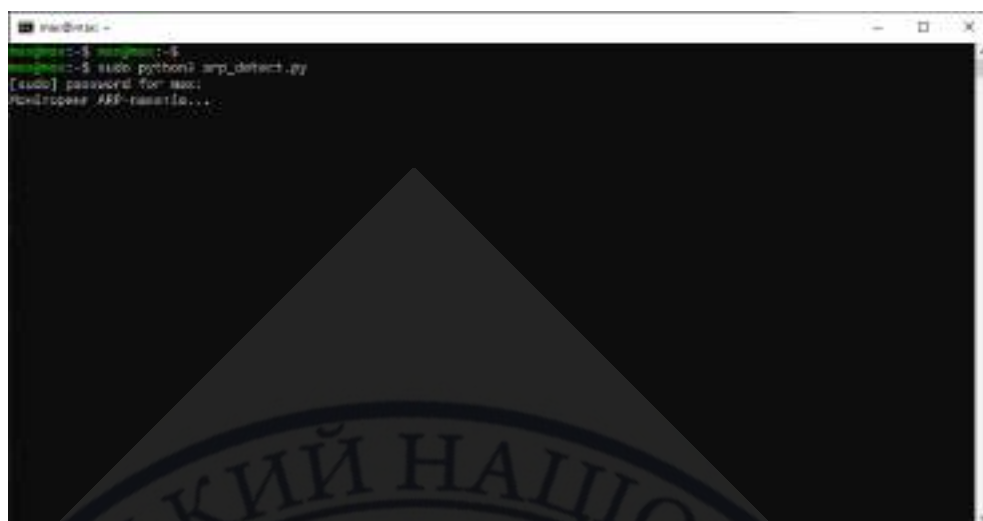


Рисунок 3.10 – Запуск скрипта

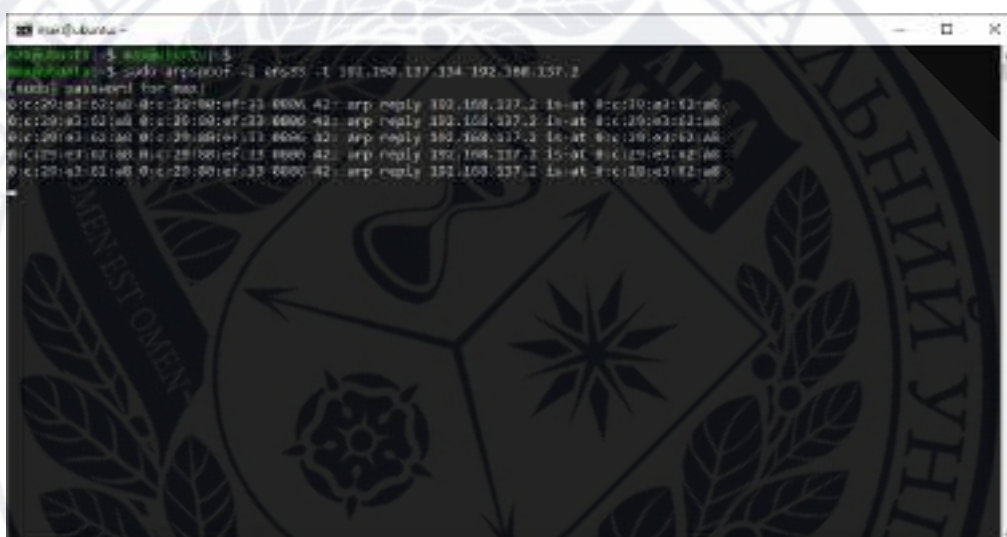


Рисунок 3.11 – Атака ARP-спуфінгу

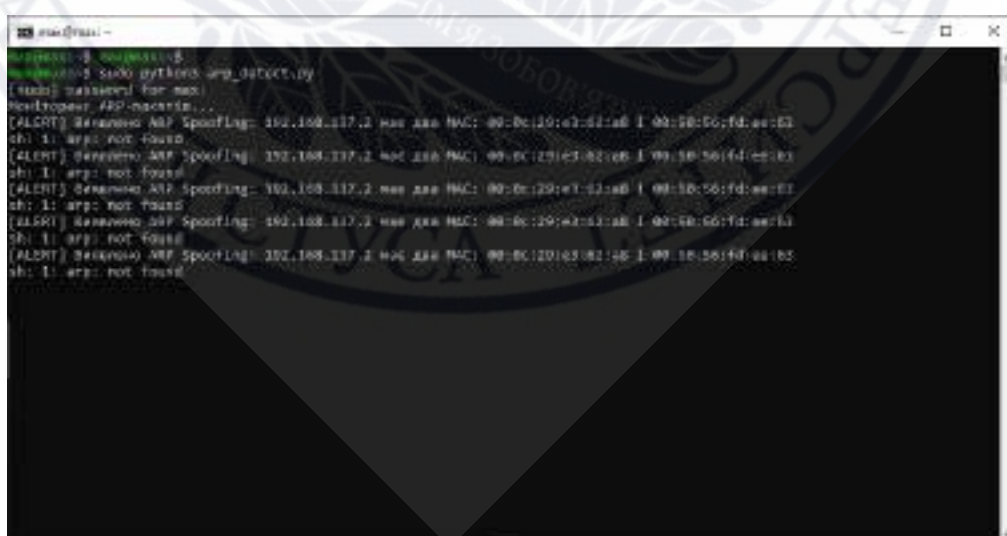


Рисунок 3.12 – Попередження атаки ARP-спуфінгу

Результати експерименту визнані успішними, розроблена програма виявила стільки ж атак скільки і Snort.

Узагальнені результати порівняння функціональності авторської розробки та існуючого рішення Snort при виявленні ARP-спуффінгу наведено у таблиці 3.1:

Таблиця 3.1 – Порівняння авторської розробки та існуючого рішення Snort

Критерій	Snort	Власний програмний код
Відкритість логіки виявлення	Закрита (вбудовані механізми фільтрації та обробки)	Повністю відкрита, код можна змінювати в будь-який момент
Можливість адаптації	Лише через написання нових правил, складна модифікація ядра	Можна вносити зміни у логіку виявлення без обмежень
Оновлення	Залежить від розробників Snort	Оновлення виконує сам адміністратор або розробник
Залежність від Інтернету	Потрібен доступ для отримання оновлень	Не потрібен (логіка локальна, оновлюється вручну)
Рівень гнучкості	Високий, але в межах правил і плагінів	Максимальний - можна реалізувати будь-яку логіку
Час реакції на нову загрозу	Відтермінований (до появи оновлень від розробників)	Миттєвий — достатньо змінити код
Простота розгортання	Потребує встановлення, налаштування, написання правил	Простий скрипт, легко розгортається у середовищі Python

Розроблений в ході дослідження програмний код є гнучкішим за існуючого рішення Snort у виявленні ARP-спуффінгу, оскільки дозволяє реалізувати більш точні алгоритми аналізу трафіку з урахуванням особливостей конкретного середовища. Існуюче рішення Snort має обмежений сигнатурний підхід та для розширення функціональності потребує задіяння додаткових модулів (та відповідно, їх налаштування), що є складним процесом.

ВИСНОВКИ

У результаті проведеної роботи вирішені поставлені завдання та досягнута її мета, що дозволяє зробити наступні висновки:

1. Проведено комплексний пошук, збір, систематизацію та аналіз інформації щодо комп'ютерних атак, які реалізуються з використанням технологій штучного інтелекту. У процесі дослідження запропоновано класифікацією кібератак із використанням штучного інтелекту та зроблено акцент на атаках типу «Man in the Middle». Розглянуто основні різновиди атак цього типу, за допомогою ілюстративного матеріалу наочно представлено сценарії її здійснення для кожного підтипу;

2. Проведено детальний аналіз сучасних методів виявлення атак типу «Man in the Middle» та узагальнено основні способи захисту від таких вторгнень. Особливу увагу приділено вразливостям у бездротових з'єднаннях, зокрема Bluetooth та Wi-Fi. Описано відповідні мережеві протоколи та актуальні механізми захисту, що застосовуються для запобігання MITM-атакам у бездротовому середовищі;

3. Здійснено розробку, тестування та практичну реалізацію програмного рішення попередження атак типу «Man in the Middle», що здійснюються з використанням генеративного штучного інтелекту. За результатами практичного експерименту здійснено порівняння ефективності виявлення таких атак за допомогою відомої системи Snort та власної розробки на мові Python. Успішне тестування розробленого програмного коду в реальному середовищі підтверджує доцільність запропонованого рішення, його простоту та ефективність.

У бакалаврській роботі теоретично охарактеризовано загрози, пов'язані з використанням штучного інтелекту в реалізації комп'ютерних атак типу «Man in the Middle», здійснено практичну перевірку ефективності методів їх виявлення. Отримані результати підтверджують актуальність проблеми, особливо в контексті швидкого розвитку генеративних технологій. Запропонований підхід може слугувати основою для подальших досліджень у напрямку створення адаптивних систем виявлення та попередження складних кібератак.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. CERT-UA минулого року опрацювала 4315 кіберінцидентів, 2025 URL: <https://cip.gov.ua/ua/news/cert-ua-minulogo-roku-opracyuvala-4315-kiberincidentiv> (Дата звернення: 23.01.2025).
2. Положення про порядок розроблення, виробництва та експлуатації засобів криптографічного захисту інформації, затверджене наказом Адміністрації Держспецзв'язку від 20.07.2007 року № 141. (Дата звернення: 23.01.2025).
3. Закон України «Про основні засади забезпечення кібербезпеки України», 2017, № 45. URL: <https://zakon.rada.gov.ua/laws/show/2163-19#Text> (Дата звернення: 23.01.2025).
4. A survey on classification of cyber-attacks on iot and iiot devices URL: <https://par.nsf.gov/servlets/purl/10295058> (Дата звернення: 25.01.2025).
5. Weaponized AI for cyber attacks URL: https://ntnuopen.ntnu.no/ntnu-xmlui/bitstream/handle/11250/3021130/Weaponized_AI_for_Cyber_Attacks_2_.pdf?sequence=1 (Дата звернення: 25.01.2025).
6. Generative AI in cyber security: new threats and solutions for adversarial attacks. URL: <https://surli.cc/tchayy> (Дата звернення: 28.01.2025).
7. From ChatGPT to ThreatGPT: impact of generative AI in cybersecurity and privacy. URL: <https://ieeexplore.ieee.org/stamp/stamp.jsp?arnumber=10198233> (Дата звернення: 28.01.2025).
8. Generative AI in cyber security of cyber physical systems: benefits and threats URL: <https://surl.gd/frwwwwt> (Дата звернення: 28.01.2025).
9. Ping Wang, Ping Chen, Zhimin Luo, Gaofeng Dong, Mengce Zheng, Nenghai Yu, and Honggang Hu. Enhancing the performance of practical profiling side-channel attacks using conditional generative adversarial networks. arXiv:2007.05285, 2020. (Дата звернення: 29.01.2025).
10. Naila Mukhtar, Lejla Batina, Stjepan Picek, and Yinan Kong. Fake it till you make it: Data augmentation using generative adversarial networks for all the crypto

you need on small devices. In Cryptographers' Track at the RSA Conference, pages 297–321. Springer, 2022. (Дата звернення: 29.01.2025).

11. Andreas Happe and Jurgen Cito. Getting pwn'd by ai: Penetration " testing with large language models. In Proceedings of the 31st ACM Joint European Software Engineering Conference and Symposium on the Foundations of Software Engineering, ESEC/FSE '23. ACM, November 2023. (Дата звернення: 29.01.2025).

12. Mika Beckerich, Laura Plein, and Sergio Coronado. Ratgpt: Turning online llms into proxies for malware attacks. arXiv:2308.09183, 2023. (Дата звернення: 29.01.2025).

13. Yin Minn Pa Pa, Shunsuke Tanizaki, Tetsui Kou, Michel van Eeten, Katsunari Yoshioka, and Tsutomu Matsumoto. An attacker's dream? exploring the capabilities of chatgpt for developing malware. In Proceedings of the 16th Cyber Security Experimentation and Test Workshop, CSET '23, page 10–18, New York, NY, USA, 2023. Association for Computing Machinery. (Дата звернення: 30.01.2025).

14. AI-Powered GUI attack and its defensive methods. URL: <https://arxiv.org/pdf/2001.09388> (Дата звернення: 30.01.2025).

15. EagleEye: fast sub-net evaluation for efficient neural network pruning URL: <https://arxiv.org/pdf/2007.02491> (Дата звернення: 31.01.2025).

16. The role of artificial intelligence in cyber operations: an analysis of ai and its application to malware-based cyberattacks and proactive cybersecurity URL: <https://www.proquest.com/openview/3c6a2b9ac66b24723d521943d178a49a/1?pq-origsite=gscholar&cbl=18750&diss=y> (Дата звернення: 05.02.2025).

17. Impact of AI on the cyber kill chain: a systematic review URL: [https://www.cell.com/heliyon/pdf/S2405-8440\(24\)16730-8.pdf](https://www.cell.com/heliyon/pdf/S2405-8440(24)16730-8.pdf) (Дата звернення: 05.02.2025).

18. Enhancing robustness of malware detection using synthetically-adversarial samples URL: <https://surl.it/vpuace> (Дата звернення: 05.02.2025).

19. Adversarial machine learning attacks and defense methods in the cyber security domain URL: <https://arxiv.org/pdf/2007.02407> (Дата звернення: 06.02.2025).

20. Adversarial machine learning in the context of network security challenges and solutions URL: <https://surl.lu/tksens> (Дата звернення: 07.02.2025).

21. NIST trustworthy and responsible AI NIST AI 100-2e2023. Adversarial machine learning a taxonomy and terminology of attacks and mitigations URL: <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-2e2023.pdf> (Дата звернення: 07.02.2025).

22. Adversarial machine learning - industry perspectives. URL: <https://arxiv.org/pdf/2002.05646> (Дата звернення: 07.02.2025).

23. Adversarial machine learning attacks and defences in multi-agent reinforcement learning URL: <https://dl.acm.org/doi/pdf/10.1145/3708320> (Дата звернення: 08.02.2025).

24. Adversarial attacks on machine learning cybersecurity defences in Industrial Control Systems URL: <https://surl.li/lfallf> (Дата звернення: 08.02.2025).

25. A Survey on machine learning adversarial attacks URL: <https://enigma.unb.br/index.php/enigma/article/viewFile/76/57> (Дата звернення: 08.02.2025).

26. A System-Driven taxonomy of attacks and defenses in adversarial machine learning. URL: <https://pmc.ncbi.nlm.nih.gov/articles/PMC7971418/pdf/nihms-1621972.pdf> (Дата звернення: 08.02.2025).

27. NAttack! Adversarial attacks to bypass a GAN based classifier trained to detect network intrusion URL: <https://arxiv.org/pdf/2002.08527> (Дата звернення: 08.02.2025).

28. Danish Javeed, Umar MohammedBadamasi, Cosmas Obiora Ndubuisi, Faiza Soomro and Muhammad Asif. Man in the Middle Attacks: Analysis, Motivation and

Prevention. International Journal of Computer Networks and Communications Security VOL. 8, NO. 7, July 2020, 52–58 (Дата звернення: 10.02.2025).

29. Man-in-the-browser attacks against IoT devices: a study of smart homes. URL: <https://surl.li/iiazic> (Дата звернення: 10.02.2025).

30. Browser-in-the-Middle (BitM) attack. URL: <https://surl.li/vjtwhd> (Дата звернення: 10.02.2025).

31. Security risks from the modern man-in-the-middle attacks. URL: https://www.meste.org/mest/MEST_Najava/XXV_Cekerevac.pdf (Дата звернення: 10.02.2025).

32. Evil twin attack mitigation techniques in 802.11 networks. URL: https://www.researchgate.net/profile/Raja_Muthalagu2/publication/352868901_Evil_Twin_Attack_Mitigation_Techniques_in_80211_Networks/links/60e2df57a6fdccb74506ceb3/Evil-Twin-Attack-Mitigation-Techniques-in-80211-Networks.pdf (Дата звернення: 10.02.2025).

33. Securing bluetooth low energy: a literature review URL: <https://arxiv.org/pdf/2404.16846> (Дата звернення: 15.02.2025).

34. Evaluation of Man-in-the-Middle attacks and countermeasures on autonomous vehicles. URL: <https://surl.li/qbgutp> (Дата звернення: 15.02.2025).

35. Handling of Man-In-The-Middle attack in WSN through intrusion detection system. URL: <https://surl.li/icegdj> (Дата звернення: 01.03.2025).

36. Statistical techniques for detecting cyberattacks on computer networks based on an analysis of abnormal traffic behavior. URL: <https://www.proquest.com/openview/572467e811add0949d579e0646aadf2f/1?pq-origsite=gscholar&cbl=2026671> (Дата звернення: 01.03.2025).

37. Lubov Zahoruiko, Tetiana Martianova, Mohammad Al-Hiari, Lyudmyla Polovenko, Maiia Kovalchuk, Svitlana Merinova, Volodymyr Shakhov, Bakhyt Yeraliyeva. Mathematical model and structure of a neural network for detection of

cyber attacks on information and communication systems. Informatyka Automatyka, Pomiary. ELMECO 2024 Vol. 14 No. 3, page 49-55. 2024-09-30

38. Multi-Channel Man-in-the-Middle attacks against protected Wi-Fi networks: A state of the art review. URL: <https://shorturl.gg/nLe99UK> (Дата звернення: 01.03.2025).

39. MitM attack detection in BLE networks using reconstruction and classification machine learning techniques. URL: <https://inria.hal.science/hal-02948407/file/ml-ble-mlcs2020-CR.pdf> (Дата звернення: 1.03.2025).

40. Bluetooth Low Energy (BLE) Security and Privacy URL: https://drive.google.com/file/d/1baqfDYmEw1pMSBVy4pUKgKGWcOEhj_qe/view (Дата звернення: 10.03.2025).

41. Tahira Ali, Rashid Baloch, Mohsan Azeem, Dr. Muhammad Farhan, Sana Naseem, Bushra Mohsin A Systematic Review of Bluetooth Security Threats, Attacks & Analysis. International Journal of Computer Trends and Technology Volume 69 Issue 7, 1-18, July, 2021. (Дата звернення: 10.03.2025).

42. Security and privacy threats for Bluetooth low energy in IoT and wearable devices: a comprehensive survey. URL: <https://ieeexplore.ieee.org/stamp/stamp.jsp?arnumber=9706334> (Дата звернення: 28.03.2025).

43. Snort URL: <https://www.snort.org/> (Дата звернення: 05.05.2025).

44. What is Python? Executive summary URL: <https://surl.lu/fexgvh> (Дата звернення: 05.05.2025).

45. What is Scapy? URL: <https://scapy.net/> (Дата звернення: 05.05.2025).

ДОДАТКИ



Додаток А

Лістинг програмного коду

```
from scapy.all import ARP, sniff
import os

# Кеш таблиці ARP
arp_table = {}

def detect_arp_spoof(packet):
    if packet.haslayer(ARP) and packet.op == 2:
        src_ip = packet.psrc
        src_mac = packet.hwsrc

        if src_ip in arp_table:
            if arp_table[src_ip] != src_mac:
                print(f"[ALERT] Виявлено ARP Spoofing: {src_ip}
має два MAC: {arp_table[src_ip]} і {src_mac}")
                os.system(f'arp -d {src_ip}')
            else:
                arp_table[src_ip] = src_mac

print("Моніторинг ARP-пакетів...")
sniff(filter="arp", prn=detect_arp_spoof, store=0)
```

ДЕКЛАРАЦІЯ

про дотримання академічної доброчесності

Я, Ласкавчук Максим Андрійович, здобувач вищої освіти ОПП «Кібербезпека» автор кваліфікаційної (бакалаврської) роботи на тему: «Попередження комп'ютерних атак типу MAN IN THE MIDDLE, що здійснюються з використанням генеративного штучного інтелекту»

Повністю вказується ПІБ та статус (освітня (освітньо-наукова) програма – для здобувачів вищої освіти, назва кваліфікаційної роботи)

що нижче підписався, розуміючи та підтримуючи загальновизнані засади справедливості, доброчесності та законності,

ЗОБОВ'ЯЗУЮСЬ:

дотримуватися принципів та правил академічної доброчесності, що визначені законодавством України, локальними нормативними актами Донецького національного університету імені Василя Стуса, положеннями, правилами, умовами, визначеними іншими суб'єктами, та не допускати їх порушення.

ПІДТВЕРДЖУЮ:

що мені відомі положення статті 42 Закону України «Про освіту»;
що у даній роботі не представляв чийсь роботи повністю або частково як свої власні. Там, де я скористалася/скористався працею інших, я зробив відповідні посилання на джерела інформації;

що дана робота не передавалась іншим особам і подається вперше, не порушує авторських та суміжних прав закріплених статтями 21-25 Закону України «Про авторське право та суміжні права», а дані та інформація не отримувались в недозволений спосіб.

УСВІДОМЛЮЮ:

що ця робота може бути перевірена університетом на плагіат або інші порушення академічної доброчесності, в тому числі з використанням спеціалізованих сервісів;

що у разі порушення академічної доброчесності, до мене можуть бути застосовані процедури, передбачені законодавством України та Кодексом академічної доброчесності та корпоративної етики Донецького національного університету імені Василя Стуса, іншими локальними нормативними актами університету, та я можу бути притягнутий до академічної відповідальності.

(дата)

(підпис)

ISSN 2617-0922 (Online)

ISSN 2617-0914 (Print)

ДОНЕЦЬКИЙ НАЦІОНАЛЬНИЙ УНІВЕРСИТЕТ ІМЕНІ ВАСИЛЯ СТУСА
СТУДЕНТСЬКЕ НАУКОВЕ ТОВАРИСТВО



ВІСНИК

СТУДЕНТСЬКОГО НАУКОВОГО ТОВАРИСТВА
ДОНЕЦЬКОГО НАЦІОНАЛЬНОГО УНІВЕРСИТЕТУ
ІМЕНІ ВАСИЛЯ СТУСА

ВИПУСК 17

ТОМ 1

Вінниця
2025

ЕКОНОМІЧНІ НАУКИ

<i>Бартош А. А.</i> Дивідендна політика корпорацій: фактори впливу та модель прийняття рішень.....	112
<i>Боцвін М. Р.</i> Мотивація персоналу в сучасних воєнних умовах	116
<i>Винідіктова В. С.</i> Податкове планування як інструмент оптимізації діяльності підприємств.....	122
<i>Винідіктова В. С.</i> Виклики впровадження цифрових технологій обліку на вітчизняних будівельних підприємствах	127
<i>Герцун З. С.</i> Характеристика понять ліквідності та платоспроможності й аналіз шляхів їх покращення в умовах воєнного часу у роздрібній торгівлі	131
<i>Максим'як А. Я.</i> Стратегічне формування облікової політики як ключовий механізм оптимізації податкових зобов'язань бізнесу	136
<i>Максим'як А. Я.</i> Облік витрат туристичних компаній у контексті реалізації цілей сталого розвитку.....	142
<i>Шевчук В. Р.</i> Плинність кадрів на підприємстві та шляхи її зменшення	146
<i>Юздепська А. А.</i> Напрями розвитку страхування кредитно-фінансових ризиків банків України в умовах кризи.....	150
<i>Янковська Н. В.</i> Облікове рішення сучасних викликів для вітчизняних бюджетних установ	154
<i>Янковська Н. В.</i> Проблеми аналізу діяльності вітчизняних бюджетних установ	158

ПРИРОДНИЧІ ТА ТЕХНІЧНІ НАУКИ

<i>Байраківська В. В.</i> GPT-4 та майбутнє генеративного штучного інтелекту.....	162
<i>Бояренко А. В.</i> Теоретичні основи цифрової грамотності та комунікаційної діяльності.....	165
<i>Бурківський О. С.</i> Оптимізація військової логістики в умовах бойових дій за допомогою алгоритму A*	168
<i>Срмак Д. М.</i> Технологія блокчейн: переваги, виклики та загрози у сфері кібербезпеки.....	172
<i>Коба Б. О.</i> Статистичне дослідження валютної пари «євро – долар» з використанням штучного інтелекту	176
<i>Ковальчук Л. П.</i> Шкідливі міждомени запити.....	181
<i>Круць Ю. О.</i> Роль Agile-методології у вдосконаленні цифрових комунікацій підприємства	184
<i>Ласкавчук М. А.</i> Попередження комп'ютерної атаки типу «Man in the middle», що здійснюється за допомогою генеративного штучного інтелекту.....	187
<i>Лисак В. О.</i> Застосування полінома Лагранжа для інтерполяції обсягу електронної торгівлі у сфері інформаційних послуг	192
<i>Оліх В. І.</i> Оптимізаційне моделювання рекламних витрат (витрат на рекламу) приватного підприємства	196
<i>Поліщук В. С.</i> Розробка системи локаторів та оптимізації логістичних маршрутів для військових з урахуванням уникнення дронів	199
<i>Родюк А. І.</i> Застосування диференціальних рівнянь із частинними похідними в математичному моделюванні.....	203
<i>Родюк К. О.</i> Принцип Діріхле та його застосування	207
<i>Труханська В. О.</i> Аналіз NLP-технологій для CRM-систем.....	210
<i>Турецька О. В.</i> Стан нормативно-правового регулювання штучного інтелекту в Україні та світі.....	215
<i>Філоненко М. В.</i> Адвокаційні технології в бібліотечній справі: зарубіжний досвід	220
<i>Флуд Д. В.</i> Цифровізація національної архівної спадщини України в умовах глобальних викликів.....	224
<i>Черніхевич О. В.</i> Безпека сайту: комплексний підхід до захисту від загроз.....	230
<i>Шпикуляк І. С.</i> Стратегії ефективної електронної комунікації для уникнення спаму	233
<i>Юкальчук А. І.</i> Структура нейронної мережі для попередження та виявлення кібератак типу user to root та remote to local на підприємства критичної інфраструктури	238
<i>Ярошенко Б. М.</i> Метод дихотомії у знаходженні коренів апроксимованої функції причинно-наслідкового зв'язку метрик інвестиційних процесів у виробництві електронної продукції.....	242

УДК 004.032.26:004.056

ПОПЕРЕДЖЕННЯ КОМП'ЮТЕРНОЇ АТАКИ ТИПУ «MAN IN THE MIDDLE», ЩО ЗДІЙСНЮЄТЬСЯ ЗА ДОПОМОГОЮ ГЕНЕРАТИВНОГО ШТУЧНОГО ІНТЕЛЕКТУ

М. А. Ласкавчук, Л. В. Загоруйко

Анотація. У дослідженні розглянуто комп'ютерну атаку типу «Man in the Middle», наведено її основні методи реалізації та потенційні загрози. Особливу увагу приділено ролі генеративного штучного інтелекту, який має змогу допомогти користувачам із мінімальними знаннями у сфері кібербезпеки здійснювати подібні атаки. Запропоновано підхід до виявлення та попередження атак за допомогою системи запобігання вторгненням, що підвищує рівень безпеки інформаційного середовища.

Ключові слова: комп'ютерна атака, людина посередній, кібератака, штучний інтелект, кібербезпека.

Вступ. За інформацією Державної служби спеціального зв'язку та захисту інформації України, урядова команда реагування на комп'ютерні надзвичайні події CERT-UA, яка діє при Держспецзв'язку, у 2024 р. опрацювала 4 315 кіберінцидентів [1]. Це на 69,8 % більше, ніж роком раніше, коли кіберзлочинці атакували український кіберпростір 2 541 раз. Найпоширенішими типами інцидентів є розповсюдження шкідливого програмного забезпечення, фішинг, шкідливе підключення, компрометація облікового запису або системи. Метою зловмисників є викрадення чутливої інформації, а також знищення даних та інформаційних систем [1]. Існують різні методи проведення комп'ютерних атак, одним із таких методів є атака типу «людина посередній» (англ. *man in the middle*, далі – MITM). Наразі неможливо навести статистику збитків, що відбуваються саме від комп'ютерних атак типу MITM, оскільки кожен випадок атаки має свої унікальні характеристики та наслідки. Поширеність кіберзагроз у всіх сферах суспільного та державного життя, а також удосконалення методів їх реалізації, зокрема із застосуванням штучного інтелекту, актуалізують необхідність глибокого аналізу механізмів та сценаріїв кібератак. Це вимагає перегляду наявних стратегій і тактик протидії, а також пошуку нових підходів для підвищення ефективності кіберзахисту.

Відповідно до статті 1 Закону України «Про основні засади забезпечення кібербезпеки України» [2], кібератака (комп'ютерна атака) – це спрямовані (навмисні) дії в кіберпросторі, які здійснюються за допомогою засобів електронних комунікацій (включно з інформаційно-комунікаційними технологіями, програмними, програмно-апаратними засобами, іншими технічними та технологічними засобами і обладнанням) та спрямовані на досягнення однієї або сукупності таких цілей: порушення конфіденційності, цілісності, доступності електронних інформаційних ресурсів, що обробляються (передаються, зберігаються) в комунікаційних та/або технологічних системах, отримання несанкціонованого доступу до таких ресурсів; порушення безпеки, сталого, надійного та штатного режиму функціонування комунікаційних та/або технологічних систем; використання комунікаційної системи, її ресурсів та засобів електронних комунікацій для здійснення кібератак на інші об'єкти кіберзахисту.

Комп'ютерна атака типу MITM – це ситуація, коли суб'єкт використовує різні методики та технічні рішення, спрямовані на отримання доступу до трафіка та його декодування. Ця комп'ютерна атака реалізується із застосуванням алгоритмів перехоплення трафіка, що передається між двома кінцевими пристроями [3].

Атаки MITM можуть починатися з [3–4]:

- «отруєння» кешу протоколу дозволу адрес (ARP);
- підробка DNS;
- підміна IP-адреси;
- викрадення сеансу;
- перехоплення SSL;
- атаки на Wi-Fi;
- атаки Evil Twin;
- перехоплення електронної пошти;
- перехоплення Bluetooth.

На рис. 1 наведено, як відбувається атака MITM.



Рис. 1. Процес атаки MITM

Один із випадків таких атак відомий як динамічне підслуховування, під час якого зломисники створюють незалежні асоціації з жертвами і впливають на комунікацію між користувачами, модифікуючи повідомлення між двома користувачами, щоб вплинути на довіру між ними. Зломисник має можливість перехоплювати повідомлення між ними з метою їх модифікації. Більшість криптографічних угод включають у себе певну перевірку кінцевих точок, в основному для того, щоб протистояти MITM-атакам [4].

Класична MITM-атака проходить у два етапи.

Перший – перехоплення даних, коли злочинець інтегрується в середовище передачі даних. Далі за допомогою спуфінгу реалізує підміну IP-адрес, ARP10-повідомлень, сервера доменних імен тощо. Атаки MITM найчастіше використовують ARP-кеш, який являє собою локальний кеш із призначеними IP-адресами і зіставленими фізичними унікальними ідентифікаторами

пристроїв у мережі (MAC-адресами). Внаслідок цього реалізуються завдання отримання відомостей про структуру досліджуваної мережі та з'ясування локальних ідентифікаторів і універсальних ідентифікаторів мережі (MAC-адрес).

Другий етап – дешифрація, отримання доступу до зашифрованих даних. Оскільки існує велике розмаїття методик проведення атаки та методик протидії, однією з яких є аналіз вихідного коду (як встановлюваних застосунків, так і дослідження переданих даних у межах «пісочниці» на віртуальній машині та без наявності виходу у відкриту мережу). Такий аналіз може здійснюватися вручну фахівцем або за допомогою рішень, створених на основі навченого за відповідним напрямом ШІ.

Ще один випадок – атака типу «людина в браузері». Це особливий підтип MITM-атаки, що відбувається на кінцевій точці зв'язку [5]. В основі «людина в браузері» лежить ідея розміщення шкідливого прозорого браузера між браузером жертви та вебсервером, до якого жертва звертається для отримання послуги (наприклад, банківської послуги, соціальної мережі тощо). Метою атаки є перехоплення, запис і маніпулювання будь-яким обміном даними між жертвою і постачальником послуг [18–19]. На рис. 2 наведено атаку «людина в браузері».

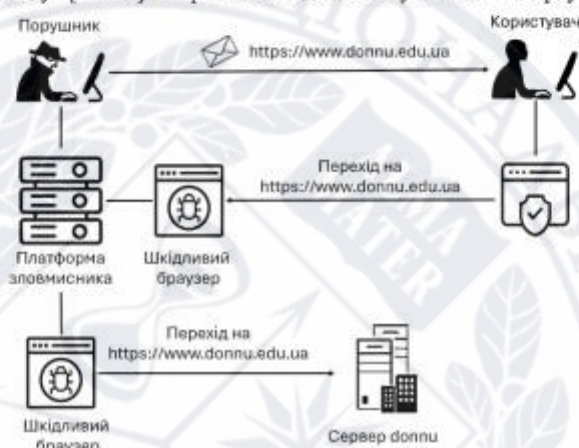


Рис. 2. Атака «людина в браузері»

– Порушник: комп'ютерний злочинець, метою якого є порушення приватності жертви, а також конфіденційності, цілісності та доступності даних, якими обмінюються дві кінцеві точки. Зловмисник встановлює і контролює шкідливий сервер, який слугує прозорим веббраузером для кінцевого користувача [6].

– Користувач: потерпілий, який за допомогою фішингових методів отримує доступ до шкідливого сервера (на якому розміщений прозорий веббраузер), налаштованого зловмисником, і починає користуватися вебдодатком [6].

– Ціль: вебдодаток, здатний надавати конфіденційні та/або цінні послуги (наприклад, банківські послуги, поштові сервіси тощо), доступ до яких здійснюється за допомогою механізму автентифікації [6].

Наступні кроки описують техніку атаки «людина в браузері» [6]:

а) порушник створює підроблене гіперпосилання, що вказує на шкідливий сервер. Посилання надсилається жертві, яку за допомогою фішингових методів заманюють натиснути на нього та пройти автентифікацію у вебдодатка;

б) натиснувши на шкідливе гіперпосилання у власному браузері, користувач мимоволі підключається до шкідливого сервера, де знаходиться прозорий веббраузер разом з низкою програм (наприклад, вебпрокси, сніффер, кейлоггер, надбудови тощо), які зловмисник використовує для перехоплення, запису та маніпулювання даними, якими обмінюються жертва та цільова вебпрограма;

eoss-conf.com



**ISSUE
№36**



COLLECTION OF SCIENTIFIC PAPERS



**2nd INTERNATIONAL
SCIENTIFIC
AND PRACTICAL
CONFERENCE**

**ACHIEVEMENTS OF
SCIENCE AND
APPLIED RESEARCH**

MAY 19-21, 2025. DUBLIN, IRELAND





Section: History and Cultural Studies

Кузьмінець Н. ПРОПАГАНДИСТСЬКО-АГІТАЦІЙНА ДІЯЛЬНІСТЬ РАДЯНСЬКОЇ ВЛАДИ В УКРАЇНІ У 1920-Х РОКАХ: МЕТОДИ ТА ІНСТРУМЕНТИ	69
--	----

Копансва В., Святненко Д. ЗАБЕЗПЕЧЕННЯ ГРАНТОВОЇ ДІЯЛЬНОСТІ В ІНФОРМАЦІЙНИХ УСТАНОВАХ УКРАЇНИ.....	73
---	----

Section: Information Technology, Cyber Security and Computer Engineering

Ксенченко Я.В., Мельник О.В. РОЗРОБКА ЗАСТОСУНКУ ДЛЯ ВИЗНАЧЕННЯ ТА УПРАВЛІННЯ ЕМОЦІЙНИМ СТАНОМ ЛЮДИНИ НА ОСНОВІ АЛГОРИТМІВ ШТУЧНОГО ІНТЕЛЕКТУ.....	77
--	----

Окунькова О. ЗАСОБИ МОВИ PYTHON ДЛЯ РОЗРОБКИ ЕКСПЕРТНИХ СИСТЕМ.....	80
--	----

Пурський В. КРИПТОГРАФІЯ ЯК ОСНОВА СУЧАСНОЇ КІБЕРБЕЗПЕКИ.....	84
---	----

Федотова О., Савік К. ТЕХНОЛОГІЇ АНАЛІТИЧНОГО ОПРАЦЮВАННЯ ДАНИХ У РОБОТІ НАУКОВИХ БІБЛІОТЕК УКРАЇНИ: СКЛАДНОЩІ ТА ПЕРСПЕКТИВНІ ШЛЯХИ ВИКОРИСТАННЯ.....	87
--	----

Ласкавчук М. СПЕЦИФІКА АТАК ТИПУ «MAN IN THE MIDDLE».....	91
---	----

Hladun O., Molchanova M., Zalutska O., Mazurets O. EFFECTIVENESS RESEARCH OF USING VIT NEURAL NETWORK ARCHITECTURE FOR CLASSIFYING THE DESTROYED BUILDINGS REMAINS.....	96
---	----

Мельник О.В., Сливка Р.М. ІНТЕГРАЦІЯ JETPACK COMPOSE, ROOM ТА KOTLIN COROUTINES У СТВОРЕННІ МОБІЛЬНОГО ЗАСТОСУНКУ.....	101
---	-----



Dílmezinárodní kolektivní monografie. Mezinárodní Ekonomický Institut s.r.o.. Česká republika: Mezinárodní Ekonomický Institut s.r.o. 2024. Str.323-340.

8. Федотова О. О. Особливості використання інформаційних технологій на підприємствах в ході управлінської діяльності. Культурологія та соціальні комунікації: інноваційні стратегії розвитку: матер. міжнар. наук. конф., присвяч. 95-річчю ХДАК (21–22 листоп. 2024 р.) У 2 ч. Ч. 1 / Харків. держ. акад. культури; [під ред. доц. Н. Рябухи та ін.]. Харків: ХДАК, 2024. С.209–210.

СПЕЦИФІКА АТАК ТИПУ «MAN IN THE MIDDLE»

Ласкавчук Максим
 здобувач вищої освіти

Кафедра прикладної математики та кібербезпеки
 Донецький національний університет імені Василя Стуса, Україна

Анотація. У статті розглядається специфіка атак «Man In The Middle», наводяться її основні методи реалізації та потенційні загрози. Проаналізовано загрози для конфіденційності, цілісності та доступності даних під час таких атак, а також акцентується увага на актуальності проблеми в умовах розвитку цифрових комунікацій.

Ключові слова: комп'ютерна атака, людина посередині, кібератака, кібербезпека.

Введення. За інформацію Державної служби спеціального зв'язку та захисту інформації України, урядова команда реагування на комп'ютерні надзвичайні події CERT-UA, яка діє при Держспецзв'язку, в 2024 році опрацювала 4315 кіберінцидентів [1]. Це на 69,8% більше, ніж роком раніше, коли кіберзлочинці атакували український кіберпростір 2541 раз. Найпоширенішими типами інцидентів є розповсюдження шкідливого програмного забезпечення, фішинг, шкідливе підключення, компрометація облікового запису або системи. Існують різні методи проведення комп'ютерних атак, одним із таких методів є атака типу «людина посередині» (англ. man in the middle, далі - MITM)

Мета та задачі дослідження. Метою дослідження є вивчення специфіки комп'ютерних атак типу «Man In The Middle» (MITM), зосереджене на їх характерних особливостях, методах реалізації та загрозах для інформаційної безпеки. У процесі роботи було здійснено пошук, збір, систематизацію та аналіз інформації щодо MITM-атак, вивчено типові сценарії їх здійснення та потенційні наслідки.

Результати дослідження і їх обговорення. Комп'ютерна атака типу MITM - це атака, коли суб'єкт використовує різні методики та технічні рішення, спрямовані на отримання доступу до трафіку та його декодування. Дана



комп'ютерна атака реалізується із застосуванням алгоритмів перехоплення трафіку, що передається між двома кінцевими пристроями [2]. На рисунку 1 наведено, як відбувається типова атака MITM.



Рисунок 1 – Процес атаки «людина посередині»

*рисунок створений автором

Класична MITM-атака проходить у два етапи. Перший - перехоплення даних, коли злочинець інтегрується в середовище передачі даних. Далі за допомогою спуфінгу реалізує підміну IP-адрес, ARP-повідомлень, сервера доменних імен тощо. Другий етап – дешифрування, отримання доступу до зашифрованих даних.

Атака типу «людина в браузері», це особливий підтип MITM-атаки, яка здійснюється на стороні кінцевої точки зв'язку [3]. В основі «людина в браузері» лежить ідея розміщення шкідливого прозорого браузера між браузером жертви та веб-сервером, до якого жертва звертається для отримання послуги (наприклад, банківської послуги, соціальної мережі тощо). Метою атаки є перехоплення, запис і маніпулювання будь-яким обміном даними між жертвою і постачальником послуг [3-4]. На рисунку 2 наведено атаку «людина в браузері».



Рисунок 2 - Атака «людина в браузері»

*рисунок створений автором



Наступні кроки описують техніку атаки «людина в браузері» [4]:

а) порушник створює підроблене гіперпосилання, що вказує на шкідливий сервер. Посилання надсилається жертві, яку за допомогою фішингових методів заманюють натиснути на нього та пройти аутентифікацію у веб-додатку;

б) натиснувши на шкідливе гіперпосилання у власному браузері, користувач підключається до шкідливого сервера, де знаходиться прозорий веб-браузер разом з низкою програм (наприклад, веб-проксі, сніффер, кейлоггер, надбудови тощо), які зловмисник використовує для перехоплення, запису та маніпулювання даними, якими обмінюються жертва та цільова веб-програма;

в) шкідливий сервер отримує доступ до цільового веб-додатку (наприклад, банківського додатку, соціальної мережі тощо) через шкідливий веб-браузер абсолютно прозорим способом. Під час атаки користувач може використовувати всі функціональні можливості цільового додатка (плюс ті, які тепер може запропонувати зловмисний сервер). Всі операції жертви тепер можуть бути перехоплені, а отже, записані і використані зловмисником для маніпуляцій. Наприклад, порушник зможе перехопити облікові дані, надані жертвою в якійсь формі автентифікації, незалежно від того, які протоколи безпеки налаштовані у веб-додатку.

Атаки MITM можуть починатися з [5]:

а) отруєння кешу протоколу дозволу адрес;

б) підробки DNS;

в) підміни IP-адреси;

г) викрадення сеансу;

д) перехоплення SSL.

Отруєння кешу протоколу дозволу адрес. ARP є надійною та життєво важливою процедурою для інфраструктури локальних мереж. Зловмисники вносять зміни в кеш-таблицю локального ARP і підпорядковують MAC-адресу хоста цільовій IP-адресі. Атака здійснюється з метою отримання доступ до конфіденційної інформації користувача.

Підробка DNS. Ця атака є найнебезпечнішою атакою, що виконується з використанням кешування для підвищення продуктивності. Вона може складатися з трьох основних типів атак. Перший, підслуховування пакетів даних під час процесу відповіді. Другий, використання атаки на день народження для розміщення кешу. Третій, злом несанкціонованого DNS. Щоб завершити підміну DNS, зловмисник контролює локальний доступ до DNS, змушуючи ціль використовувати неавторизований сервер.

Підміна IP-адреси. Атака, при якій зловмисник зловмисник маніпулює вихідною IP-адресою, щоб видати себе за іншу особу або пристрій у мережі. Оскільки IP-адреса визначає відправника та отримувача пакета, її підробка дозволяє зловмиснику втручатися у зв'язок, отримувати або змінювати інформацію без відома справжніх сторін. Це дає змогу контролювати потік даних і приховувати реальне джерело атаки.



Викрадення сеансу. Зловмисник викрадає або видає себе за токен сеансу користувача, щоб отримати несанкціонований доступ або виконати дії від його імені;

Сніфінг - це один із способів атаки, за допомогою якого зловмисник перехоплює мережу і викрадає ідентифікатор сеансу. Зловмисник продовжує стежити за мережним трафіком жертви і намагається виявити будь-який пакет, що проходить в незашифрованому вигляді, якщо зловмисник знаходить пакет без шифрування, то він намагається з'ясувати, чи містить він ідентифікатор сеансу чи ні. Якщо зловмисник отримує ідентифікатор сеансу в пакеті, то, використовуючи цей ідентифікатор, він перехоплює вже створений сеанс між жертвою і сервером.

Перехоплення SSL. Зловмисник перехоплює та розшифровує зашифровані SSL/TLS-повідомлення, надаючи підроблені сертифікати. Оскільки під час атаки на перехоплення SSL веб-браузери зазвичай не видають жодних попереджень кінцевим користувачам. Виконавши успішну атаку на перехоплення SSL, порушники можуть легко отримати облікові дані жертви і використовувати їх не компрометуючи себе. Це можуть бути банківські рахунки, дані кредитних карток, та інша конфіденційна інформація. На рис. 3 продемонстровано базову атаку на зняття SSL.

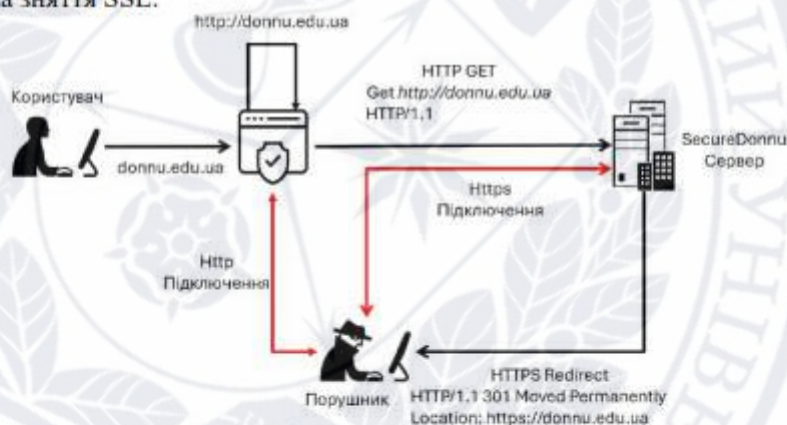


Рисунок 3 - Базова атака на зняття SSL

*рисунок створений автором

Пояснюється атака наступним чином:

- користувач заходить на сайт (наприклад, університетський) за його URL-адресою, наприклад: www.donnu.edu.ua;
- веб-браузер автоматично додає <http://> перед www.donnu.edu.ua, таким чином роблячи <http://www.donnu.edu.ua>;
- тепер браузер звертається до сервера SecureDonnu методом HTTP GET;
- сервер SecureDonnu отримує запит і перенаправляє його на <https://www.donnu.edu.ua>;



- оскільки перенаправлення HTTPS проходить через, незахищену мережу, можливі наступні варіанти:

- зловмисник може відключити HTTPS для користувача;
- зловмисник може взаємодіяти з SecureDonnu через HTTPS і робити все від імені реального користувача;
- зазвичай не відображається жодних попереджень у веб-браузері;
- користувач, як правило, не знає про цю атаку;
- користувач все ще може підключатися до сервера SecureDonnu через HTTP через зловмисника, навіть якщо університет фактично забезпечив HTTPS для своїх користувачів.

Висновки. У результаті проведеного дослідження встановлено, що атаки типу «Man In The Middle» мають різноманітні варіації та методи реалізації, що робить їх особливо небезпечними в умовах сучасних мережових середовищ. Показано, що навіть незначні помилки або неправильне налаштування мережових адаптерів можуть створити передумови для успішного здійснення таких атак, що, у свою чергу, призводить до порушення конфіденційності, цілісності та доступності інформації. Отже, забезпечення правильного конфігурування мережового обладнання та впровадження засобів захисту є критично важливими елементами інформаційної безпеки.

Список використаних джерел

1. CERT-UA. 2025. CERT-UA минулого року опрацювала 4315 кіберінцидентів. Державна служба спеціального зв'язку та захисту інформації України. URL: <https://cip.gov.ua/ua/news/cert-ua-minulogo-roku-opracyuvala-4315-kiberincidentiv>
2. Danish Javeed, Umar MohammedBadamasi, Cosmas Obiora Ndubuisi, Faiza Soomro and Muhammad Asif. Man in the Middle Attacks: Analysis, Motivation and Prevention. International Journal of Computer Networks and Communications Security VOL. 8, NO. 7, July 2020, 52–58
3. Sampsa Rauti, Samuli Laato, Tinja Pitkämäki. 2021. Man-in-the-browser attacks against IoT devices: a study of smart homes. URL: https://www.researchgate.net/profile/Samuli-Laato/publication/350907813_Man-in-the-Browser_Attacks_Against_IoT_Devices_A_Study_of_Smart_Homes/links/607d30b3881fa114b4111176/Man-in-the-Browser-Attacks-Against-IoT-Devices-A-Study-of-Smart-Homes.pdf
4. Franco Tommasi, Christian Catalano, Ivan Taurino. 2021. Browser-in-the-Middle (BitM) attack URL: <https://link.springer.com/content/pdf/10.1007/s10207-021-00548-5.pdf>
5. Zoran Cekerevac, Petar Cekerevac, Lyudmila Prigoda, Fawzi Al-Naima. 2025. Security risks from the modern man-in-the-middle attacks. URL: https://www.meste.org/mest/MEST_Najava/XXV_Cekerevac.pdf