

УДК 006.91, 681.2-5

DOI 10.31471/1993-9981-2024-2(53)-137-146

МЕТОДОЛОГІЧНА МАТРИЦЯ ЗАПОБІГАННЯ МІТМ-АТАКАМ У СТРАТЕГІЧНИХ ТА ЦИВІЛЬНИХ ІНФРАСТРУКТУРАХ

А. А. Шкітов¹, Н. Г. Бузоверя², О. А. Стадніченко³, К. М. Мажуга³, Т. В. Гуменюк²

¹*Відкритий міжнародний університет розвитку людини «Україна»;
вул. Львівська 23, м. Київ, 03115, Україна; e-mail: office@uiu.ua*

²*Івано-Франківський національний технічний університет нафти і газу;
вул. Карпатська, 15, м. Івано-Франківськ, 76019, Україна;
e-mail: nadiia.shyrmovska@nung.edu.ua*

³*Український інститут експертизи сортів рослин;
вул. Генерала Родимцева, 15, м. Київ, 03041, Україна; e-mail: sops@sops.gov.ua*

За сучасних умов кіберзагроза є однією з найбільших сучасних проблем для бізнесу, організацій і звичайних користувачів Інтернету. З кожним роком частота й складність кібератаки зростають, що підтверджують статистичні дані за останні роки. Таким чином, починаючи з 2020 року, відзначено значний сплеск кіберзлочинності, зокрема через атаки на розподілені інформаційні системи. Особливу загрозу становлять атаки типу «Man-in-the-Middle» (MITM), які дозволяють зловмисникам отримувати доступ до конфіденційної інформації як корпоративних систем, так і користувачів персональних даних. Методом дослідження є аналіз видів MITM-атак, їх методики проведення, а також розробка рекомендацій щодо запобігання таким загрозам. Дослідження базується на методах аналізу та опису. У ході роботи відзначаються основні типи MITM-атак, описані їх механізми дії та вплив на інформаційну безпеку. У результатах запропоновано низку методів і рекомендацій щодо протидії цим атакам, які охоплюють як технічні, так і організаційні аспекти. Запропоновані методи спрямовані на забезпечення високого рівня захисту комп'ютерних мереж, зокрема шляхом використання сучасних технологій шифрування, двофакторної автентифікації та системи виявлення аномалій у мережевому трафіку. Результати дослідження можуть бути корисними для адміністраторів безпеки, розробників програмного забезпечення та звичайних користувачів, які хочуть підвищити безпеку своїх даних, особливо в контексті використання публічної мережі Wi-Fi. Практичне значення отриманих результатів уможливорюється при можливості використання запропонованих рекомендацій для підвищення захисту комп'ютерних мереж, зниження ризику реалізації MITM-атак та розробки більш ефективного програмного забезпечення для забезпечення інформаційної безпеки.

Ключові слова: кібератака, MITM-атака, точка доступу, FakeAP, SSL, WPA, SSID, попередження MITM-атак.

Cyber threats are one of the biggest modern problems for businesses, organizations, and ordinary Internet users. Every year, the frequency and complexity of cyber attacks are increasing, as confirmed by statistical data in recent years. Thus, starting in 2020, a significant surge in cybercrime has been noted, in particular through attacks on distributed information systems. A particular threat is posed by Man-in-the-Middle (MITM) attacks, which allow attackers to gain access to confidential information, both corporate systems and users' personal data. The research method is to analyze the types of MITM attacks, their methods of implementation, and the development of recommendations for preventing such threats. The study is based on analysis and description methods. The work identifies the main types of MITM attacks, describes their mechanisms of action and impact on information security. The results suggest a number of methods and recommendations for countering these attacks, covering both technical and organizational aspects. The proposed methods are aimed at ensuring a high level of protection of computer networks, in particular by using modern encryption technologies, two-factor authentication and a system for detecting anomalies in network traffic. The results of the study may be useful for security administrators, software developers and ordinary users who want to increase the security of their data, especially in the context of using a public Wi-Fi network. The practical significance of the results obtained is made possible by the possibility of using the proposed recommendations to increase the security of computer networks, reduce the risk of implementing a MITM attack and develop more effective software to ensure information security.

Keywords: cyberattack, MITM attack, access point, FakeAP, SSL, WPA, SSID, MITM attack prevention.

Вступ

За сучасних умов цифрової епохи, MITM (Man-in-the-Middle) атаки — це одна з найнебезпечніших кіберзагроз, особливо для критичних інфраструктур, таких як енергетичні мережі та фінансові системи. Такі системи передбачають перехоплення або зміну переданих даних між двома сторонами зловмисником, що є актуальним у кіберсучасності. У цій статті розглядаються сучасні методи запобігання MITM-атакам та їх математичне моделювання з метою виявлення таких атак. Це полягає в наступному викладі.

Суть атаки man-in-the-middle доволі проста: злочинець таємно перехоплює трафік з одного комп'ютера та відправляє його кінцевому одержувачу, попередньо прочитавши та змінивши на свою користь.

Атаки MITM надають злочинцю можливість робити такі дії як підміна криптовалютного гаманця для викрадення коштів, перенаправлення браузера на шкідливий вебсайт або ж просто пасивний збір інформації з метою її подальшого злочинного використання.

Кожного разу, коли третя сторона перехоплює інтернет-трафік, це можна ідентифікувати як атаку MITM. Такі дії зовсім неважко зробити злочинцеві навіть без належної автентифікації. Наприклад, загальнодоступні мережі Wi-Fi є гарним джерелом для MITM, оскільки ні маршрутизатор, ні підключений комп'ютер не перевіряють її ідентичність.



Рисунок 1 – схематичний принцип атаки типу man-in-the-middle

У випадку публічної атаки через мережу Wi-Fi зловмисник повинен перебувати поблизу та під'єднатися до тієї

самої мережі, або ж просто мати комп'ютер в мережі, здатний перехоплювати трафік. Не всі атаки MITM вимагають, щоб зловмисник фізично знаходився поруч зі своєю жертвою, оскільки існує велика кількість штамів зловмисного програмного забезпечення, яке здатне викрасти трафік та підмінити інформацію в будь-якому місці, де є можливість інфікувати комп'ютер жертви.

Боротьба з атаками MITM вимагає використання певної форми автентифікації кінцевої точки, наприклад TLS або SSL, яка застосовує ключ ідентифікації, що в ідеалі не може бути підробленим. Слід зазначити, що методи автентифікації стають все сильнішими, що веде до наскрізного шифрування деяких систем.

Описові передумови у MITM-атак

У цьому контексті MITM-атаки можуть бути виконані через різні вразливості, зокрема мережеві протоколи та незахищені з'єднання. Основні цілі атак: перехоплення конфіденційної інформації (паролі, фінансові дані); модифікація переданих даних; використання здобутої інформації для інших атак (фішинг, соціальна інженерія) [2]. При цьому, нові методи атак структурують Spoofing SSL/TLS. Це, коли зловмисники можуть підмінити сертифікати для підключення до фальшивого сайту. Wi-Fi MITM-атаки: атаки в публічних Wi-Fi мережах, де перехоплюється незашифрований трафік.

Можливі загрози для інфраструктур

Такі загрози розкривають стратегічні та цивільні об'єкти. Наприклад, енергетичні мережі, мережі водопостачання та фінансові системи мають низку вразливостей: критичність переданих даних; залежність від мережевих систем управління; використання застарілих протоколів безпеки.

Методи запобігання MITM-атакам

Такі методи стверджують шифрування даних. Це протоколи TLS/SSL, що забезпечують захищене шифрування, яке ускладнює перехоплення та розшифровку даних зловмисниками. Адже TLS використовує асиметричне шифрування

для встановлення з'єднання та симетричне для передачі даних.

За особливих обставин двофакторна аутентифікація (2FA) логічно додає завдяки алгоритму пріоритетності потенційно-ефективний рівень захисту, вимагаючи не лише пароллю, але й одноразового коду чи апаратного токена, без якого ніяк неможливо увійти в систему користування.

При цьому віртуальні приватні мережі (VPN) виступають в якості драйвера, що шифрує трафік між пристроями, роблячи його надійно захищеним навіть, у публічних мережах. Особливо актуальним це постає під час воєнного стану в Україні, коли довкола є значна кількість внутрішніх ворогів (шахраїв, зловмисників, прихованих колаборантів). Адже російський ворог їх часто використовує «в сліпу», використовуючи такі вигідні ворогу умови, що здійснюють інформаційно-психологічний тиск. Серед використаних ворогом «в сліпу», можна віднести категорію тих осіб, що є за своїм станом «синдрому» поведінки на межі адекватності та не адекватності, а саме: вживання алкогольних напоїв та тютюнових виробів, що обумовлюють їх внутрішній авітаміноз. Саме такий метод спостереження потрібний під час воєнно-правового режиму.

В силу вище означеного, системи виявлення вторгнень (IDS) дозволяють виявляти підозрілу активність у реальному часі. Підписні IDS працюють на основі відомих сигнатур атак, тоді як аномальні IDS шукають відхилення від нормального трафіку.

Таким чином, мережева сегментація є однією з ефективно дієвих методів запобігання MITM-атакам в якості сегментації мережі. За цих умов розбиття мережі на ізольовані сегменти певною мірою обмежує можливість зловмисника отримати доступ до всього середовища, навіть якщо він успішно виконає атаку на одну ділянку. Адже сегментовані мережі вводять ворога в оману, а саме цим знижують ймовірність масового

перехоплення даних або впливу на всі системи одночасно.

Тому, у цьому змістовному балансі ірраціональної сенсорики та раціональної комбінаторики штучного інтелекту і машинного навчання, що є предметом означеної теми дослідження, запроваджуються алгоритми AI в ієрархічній її побудові, що здатні аналізувати мегавеликі обсяги трафіку та системно виявляти аномалії, які можуть сигналізувати про MITM-атаки. Відтак, технології UEBA (User and Entity Behavior Analytics) відстежують відхилення у поведінці користувачів.

З огляду на це, фізична безпека та апаратні методи захисту дають змогу запобігати MITM-атакам є не лише для кібербезпеки, а й для фізичного захисту апаратного забезпечення, що містить конфіденційну інформацію в умовах повоєнної України, особливо у стратегічно-критичних інфраструктурах, як прерогативи спецслужб. Саме за таких умов використання апаратних модулів безпеки (HSM) для шифрування, контролю фізичного доступу до мережевих пристроїв та впровадження апаратних токенів як частини двофакторної аутентифікації є ефективними рішеннями для зменшення ризику атак.

Математичне моделювання MITM-атак у воєнній стеганографії

Для захисту мережевих інфраструктур можна застосовувати математичні моделі, такі як орієнтовані граfi. Нехай система представляється графом $G=(V, E)$ де V — це вузли (користувачі, пристрої), а E — зв'язки між ними. MITM-атака моделюється додаванням нових ребер зловмисником для перехоплення або зміни даних.

В означеній проблематиці варто застосовувати як п'ятизначну, десятизначну, так і 16-ти значну кодову інформацію на бітовій основі. Саме така ієрархія кодування та декодування є надто важливою в орієнтованому граfi мережі. Адже в кожній ієрархічно обумовленій категорії орієнтованого графа мережі має

бути поруч із ймовірно хибними числовими знаками кодового програмування відповідна кількість вірогідних чисел в означеному програмуванні.

00000	00000000000	000000000000000000
10100	0100101001	1100010010000100

При цьому варто використовувати математичний метод індукції, що виражається формулою: $n!(a_n+1)$

$v_i \in V$ – вузли, які представляють пристрої або сервери.

$e_{ij} \in E$ – орієнтовані ребра, що представляють передачу даних від вузла v_i до вузла v_j .

Таким чином, у вірогідному алгоритмі пріоритетності, якщо зломисник перехоплює канал зв'язку, ми можемо нейромоделювати в системі генеративної розробки програмування - це додаванням нового ребра e_{zm} , де z – вузол зломисника, а m – цільовий вузол. Ймовірність успішної MITM-атаки залежить від таких факторів:

$$P(\text{MITM}) = P(\text{Intercept}) \times P(\text{Decrypt}) \times P(\text{Modify})$$

де:

- $P(\text{Intercept})$ – ймовірність перехоплення даних;
- $P(\text{Decrypt})$ – ймовірність успішного дешифрування;
- $P(\text{Modify})$ – ймовірність зміни даних без виявлення.

Застосування двофакторної аутентифікації знижує ймовірність атаки:

$$P(\text{MITM2FA}) = P(\text{MITM}) \times (1 - p_{2FA}).$$

Для моделювання можна використовувати криптографічні протоколи. Наприклад, при шифруванні симетричним та асиметричним методами:

$$c = \text{Easym}(k) + \text{Esym}(m)$$

де: c – зашифроване повідомлення, k – симетричний ключ, m – повідомлення.

Тому для контейнерів фіксованого розміру більш доцільним є використання методу псевдовипадкової перестановки (вибору) [8], зміст якого полягає в тому,

що генератор псевдовипадкових чисел (генератор ПВЧ) утворить послідовність індексів j_1, j_2, j_3 і зберігає – k -й бітів повідомлення в пікселі з індексом j_k [10].

Нехай N – загальна кількість бітів (наймолодших) у наявному контейнері; P^N – перестановка чисел $\{1, 2, \dots, N\}$. Тоді, якщо в нас є для приховання конфіденційне повідомлення довжиною n бітів, то ці біти можна просто вмонтувати замість бітів контейнера $P^N(1), P^N(2), \dots, P^N(n)$.

Функція перестановки повинна бути псевдовипадковою, іншими словами, вона повинна забезпечувати вибір бітів контейнера приблизно випадковим чином. Таким чином, секретні біти будуть рівномірно розподілені у всьому бітовому просторі контейнера. Однак, при цьому індекс визначеного біта контейнера може з'явитися в послідовності більше одного разу і в цьому випадку може відбутися «перетинання» – перекручування вже вбудованого біта. Якщо кількість бітів повідомлення набагато менша кількості молодших бітів зображення, то ймовірність перетинання є незначною, і перекручені біти надалі можуть бути відновлені за допомогою коригувальних кодів.

На наш погляд, ймовірність одного перетинання оцінюється завдяки математичному моделюванню таким чином:

$$p \approx 1 - \exp\left[-\frac{I_m \cdot (I_m - 1)}{2 \cdot I_c}\right], I_m \ll I_c. \quad (1)$$

При збільшенні I_m і $I_c = \text{const}$ дана ймовірність прагне до одиниці.

Для запобігання перетинань можна запам'ятовувати всі індекси використаних елементів B , і перед модифікацією нового пікселя виконувати попередню його перевірку на повторюваність. Також можна застосовувати генератори ПВЧ без повторюваності чисел. Останній випадок розглянемо детальніше.

Функція перестановки також залежить від секретного ключа K . При цьому генератор псевдовипадкової перестановки P_N – це функція, що для кожного значення K виробляє різні псевдовипадкові перестановки чисел $\{1, 2, \dots, N\}$.

Позначимо через P_K^N генератор перестановок з відповідним ключем K .

Якщо перестановка P_K^N захищена за обчисленням в силу її можливої складності використання, (тобто злам алгоритму вимагає не виправдано більших великих витрат обчислювальних ресурсів зломисника), то можливість розкриття змісту або припущення самого тільки виду перестановок без володіння інформацією про секретний ключ K практично прирівнюється до нуля [47, с. 65].

Секретний генератор псевдовипадкової перестановки (генератор ПВП) може бути ефективно реалізований на основі генератора псевдовипадкової функції (генератор ПВФ), котрий, як і генератор ПВП, виробляє різні, так, що не підлягають передбаченню, і прогнозуванню алгоритми функції при кожному окремому значенні ключа; однак множина значень функцій не повинне прирівнюватися за множини її визначення. Генератор ПВФ легко реалізується із секретної хеш-функції H шляхом об'єднання аргументу i із секретним ключем K та взяття від результативного бітового рядка функції H .

$$F_k(i) = H(K_i),$$

де K_i – об'єднання (конкатенація) бітових рядків K_i ; $f_k(i)$ – результативна псевдовипадкова функція від i , що залежить від параметра K .

У цьому контексті генератор Лубі (Luby) і Рекоффа (Reckoff) побудований у такий спосіб, а саме: запис виду $a \oplus b$ має на увазі побітове додавання за модулем 2 аргументи a з аргументом b , причому результат додавання має ту ж розмірність, що і a .

Нехай i – рядок двійкових даних довжиною $2I$. Розділимо i на дві частини: x та y довжини I кожна, а ключ K на чотири частини: K_1, K_2, K_3, K_4 . Тоді,

$$y = y \oplus f_{k_1}(x) = y \oplus H(K_1 \circ x);$$

$$x = x \oplus f_{k_2}(y) = x \oplus H(K_2 \circ y);$$

$$y = y \oplus f_{k_3}(x) = y \oplus H(K_3 \circ x);$$

$$x = x \oplus f_{k_4}(y) = x \oplus H(K_4 \circ y);$$

повернення в $\circ x$.

Для кожного значення ключа K алгоритм повертає псевдовипадкову перестановку із чисел $\{1, \dots, 2^{2I}\}$. Лубі й Рекофф показали, що пропонується перестановка є настільки ж секретною, наскільки й генератор ПВФ. Вони також запропонували простий алгоритм перестановки з $\{1, \dots, 2^{2I+1}\}$. Якщо значення функції f_k становить досить довгі бітові послідовності, той же ефект можна одержати, прийнявши, що y – перші I біти рядка i , а x – останні $(I+1)$ біти.

Представлена вище конструкція дозволяє одержати перестановку $P_K^{2^k}$ з $\{1, \dots, 2^k\}$ для довільного K . Однак у випадку, коли кількість бітів контейнера становить N , виникає необхідність у перестановці P_K^N з $\{1, \dots, N\}$.

Перевага методу у методологічній матриці запобігання МІТМ атакам щодо стратегічних та цивільних інфраструктур полягає в тому, що існує можливість обмежитися тільки наявними для P_K^N аргументами. Нехай $k = \lceil \log_2(N) \rceil$ (квадратні дужки означають округлення до найменшого цілого, що більше або дорівнює аргументу). Тоді $2^{k-1} < N \leq 2^k$. При цьому підраховуються значення $P_K^{2^k}(1), P_K^{2^k}(2), \dots$ і з послідовності віддаляються будь-які числа, що перевищують N . Таким чином, одержують

значення $P_K^{2^k}(1), P_K^{2^k}(2), \dots$ і з послідовності віддаються будь-які числа, що перевищують N . Таким чином, одержують значення $P_K^N(1), P_K^N(2), \dots$. Помітимо, що це стає можливим, коли функція перестановки обчислена для значень аргументів які зростають, починаючи з одиниці. Отже, генератор ПВП P^N для довільного N може бути побудований на основі алгоритму Лубі й Рекоффа. Якщо ж N є складовим (як у випадку зображення), існує більше зручний спосіб побудови генератора ПВП. Наведений нижче алгоритм заснований на блоковому кодуванні з довільним розміром блоку [8]. Що знаходить конфіденційне застосування у цій методологічній матриці запобігання МІТМ-атакам, виходячи із воєнної стеганографії з критеріального огляду квантових симуляцій та космічної гіперспектроскопічності на експертній основі.

Кількість бітів контейнера повинна становити складене число із двох співмножників приблизно однакового порядку, тобто $N = X \cdot Y$ для деяких X і Y .

У випадку, коли дані ховаються в НЗБ пікселей цифрового зображення, параметри X і Y є розмірами даного зображення. З метою будь-якого ворожого введення в оману, а також для одержання координат i -го пікселя зображення при прихованні біта повідомлення ($i \in \{1, \dots, N\}$) необхідно виконати такі обчислення:

$$\begin{aligned}x &= (i, Y) + 1; \\y &= \text{mod}(i, Y) + 1; \\x &= \text{mod}(x + f_{k_1}(y), X) + 1; \\y &= (y + f_{k_2}(x), Y) + 1; \\x &= \text{mod}(x + f_{k_2}(y), X) + 1; \\i &= (x, y) \text{ або } i = (x - 1) \cdot Y + y,\end{aligned}$$

де (i, X) і $\text{mod}(i, X)$ – функції, які повертають, відповідно, ціле й залишок від розподілу i на X . Інший варіант формули застосуємо у випадку, якщо масив зображення попередньо був розгорнутий у вектор (по рядках). Додаток одиниці необхідний при індексації елементів масиву зображення, починаючи з одиниці.

Перші два раунди алгоритму необхідні для того, щоб комбінаторно «розсіяти» біти приховуваного повідомлення серед найменш значущих бітів контейнера. При цьому, перший раунд надає випадковий характер J_c – координатам пікселя-контейнера, а другий Y – координатам. Третій раунд необхідний для протидії атаці на відкритий (незашифрований) текст [47, с. 67].

У випадку використання тільки двох раундів, нехай $i = (b - 1) \cdot Y + a$, а $P_N^K(i)$ – значення перестановки. Якщо криптоаналітик здатний припустити значення a й може одержати пару «відкритий текст – кодований текст» ($i' = (z - 1) \cdot Y + a, P_N^K(i')$) для деякого z , то він здатний установити b .

Крок 1. Повідомлення, яке необхідно приховати: $M := \ll @ \text{Колос} \gg$

Контейнер C – підмасив B синьої кольорової складової зображення. При цьому кількість бітів у повідомленні: $L_M := \text{strlen}(M), L_M := 200$ бітів; геометричні розміри контейнера: $X := \text{rows}(C), X = 128$ пікселей; $Y := \text{cols}(C), Y = 128$ пікселей; $N := XY, N := 16384$.

Крок 2. Для формування ключа використовуємо модуль, що дозволяє на підставі первинного ключа $K_0 > 2$ сформувати вектор, що буде містити R пар ключів (кожна пара ключів буде використовуватися у відповідному раунді обчислення координат x і y).

Крок 3. Вбудовування бітів повідомлення в псевдовипадкові пікселі контейнера виконаємо за допомогою модуля. На початку модуля масиву S привласнюються значення вихідного масиву C . Також виконується конвертування повідомлення зі строкового формату у вектор двійкових даних $Mvec_{bin}$. При кіберобчисленні прямокутних координат x і y використовується операція векторизації, що дозволяє поелементно складати за модулем 2 двійкові вектори K і Y (або x). При цьому розмірність зазначених векторів повинна бути однаковою, для чого використовується функція *submatrix*.

Крок 4. На стороні що приймає у програмному забезпеченні повинні бути відомі первинний ключ K_0 і масив кольоровості, у який виконувалося вбудовування (S). З останнього виходять значення X, Y, N .

За таких умов алгоритми, описувані в даному пункті, впроваджують ЦВДЗн в множини вихідного зображення. Їх перевагою є те, що для впровадження ЦВДЗн нема потреби виконувати обчислювально-громіздкі лінійні перетворення зображень. Тому, ЦВДЗн впроваджується за рахунок квантових маніпуляцій (операцій як надшвидкого здійснення відповідного результату, в жодному разі непомітного ворогові) яскравість $I(x, y) \in \{1, \dots, L\}$ або багатоколірними складовими $(r(x, y), b(x, y), g(x, y))$, а також алгоритму програмного забезпечення щодо виявлення, передбачення, запобігання та протидії можливих корупційних схем.

Дослідження MITM-атак

Поточні дослідження атак MITM зосереджені на дедалі більшій складності та доступності цих загроз, а також на заходах протидії їм. Основні напрямки дослідження включають:

Фішинг як послуга (PhaaS): ці набори інструментів демократизують атаки MITM,

роблячи їх доступними для менш досвідчених зловмисників. Імітуючи законні служби через проксі-сервери, ці атаки можуть викрасти облікові дані автентифікації, файли cookie сеансу та інші конфіденційні дані. Це створює значні ризики для підприємств, які покладаються на традиційні засоби безпеки.

Архітектури нульової довіри: дослідження підкреслюють застосування моделей нульової довіри для протидії загрозам MITM. Це передбачає впровадження повної перевірки TLS (Transport Layer Security) і надійних механізмів багатофакторної автентифікації, таких як FIDO2. Перехід від традиційної безпеки на основі VPN до архітектур на основі проксі має вирішальне значення для зменшення вразливостей.

III в атаці та захисті: Генеративні інструменти III все частіше використовуються для автоматизації та покращення виконання атак MITM, наприклад, шляхом кращого націлювання та імітації систем. І навпаки, штучний інтелект і машинне навчання також використовуються для виявлення аномалій і запобігання таким атакам.

Нові вразливості протоколів. Дослідження виявили слабкі місця в нових протоколах і системах зв'язку, зокрема тих, що використовуються в мережах IoT (Інтернет речей) і 5G. Для цього необхідно пом'якшити ці вразливості. Тому варто розробити більш безпечні методи шифрування та автентифікації.

З огляду на це, саме кібербезпека в критичній інфраструктурі: акцентує увагу на захисті основних систем, таких як фінансові мережі та хмарні сервіси, від перехоплення та злочинно-гібридних намірів маніпулювання даними на основі MITM.

Дослідження сторонніх компаній MITM атак.

Для цього необхідний такий алгоритм дій, як покрокова процедура, що полягає в наступному. Взяти характеристики URL-адрес фішингу облікових даних MITM

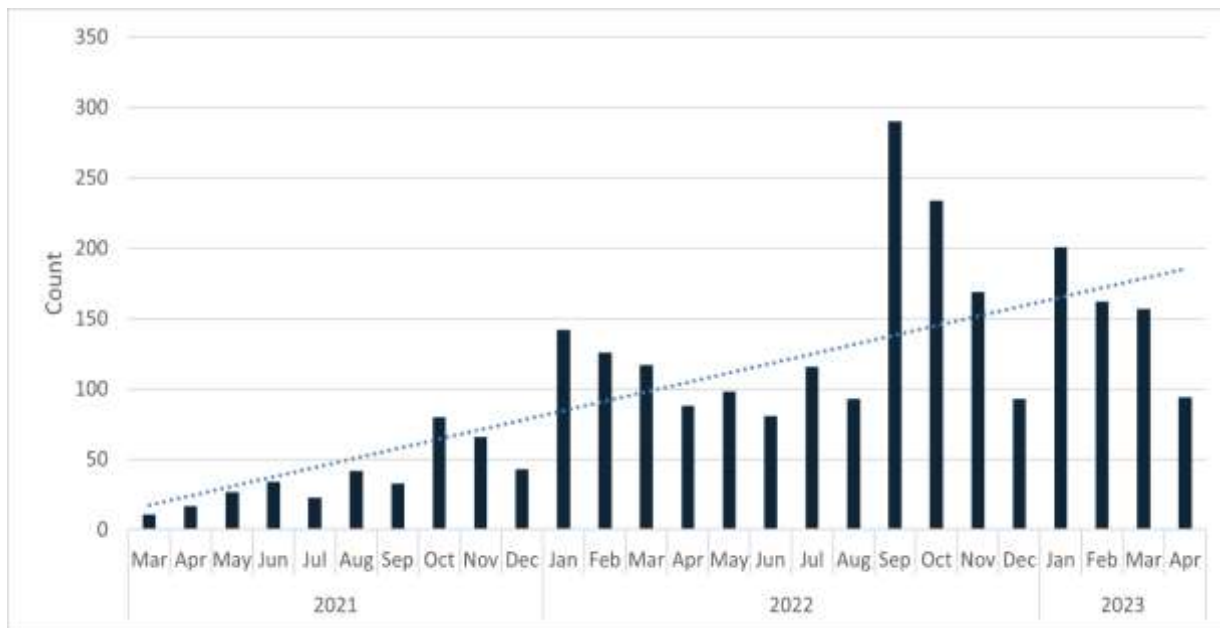


Рисунок 2 – Кількість фішингових кампаній типу «man-in-the-middle» на місяць

(описані далі в цій статті). Для цього варто окреслити порівняльний аналіз, а саме: зіставити їх із фішинговими кампаніями, виявленими в поштових скриньках клієнтів. Так, наприклад, Cofense Intelligence здійснювала метод спостережень. Проте, на наш погляд, варто звернути увагу на постійне зростання об'єму атак типу «людина посередині»[3] з метою збирання облікових даних початку 2021 року.

З огляду на це, на рисунку 2 показані основні сплески в січні та вересні 2022 року, за якими були необхідно важливі виправлення полідискретного (багатоточкового) характеру, що зберігають вихідну довгострокову траєкторію [1]. Крім того, сплески та кореляції (виправлення) можуть бути ритмодинамічно пов'язані – як циклічна обумовленість. Саме в них суб'єкти загроз розробляють, тестують та вдосконалюють можливості своїх наборів для фішингу MITM. Також існують численні інструменти з відкритим вихідним кодом (такі як evilginx2, CIPHERGinx та Muraena), які можуть надати суб'єктам загроз можливості MITM.

Отже, виявлено значну кількість цільових сторінок типу «зловмисник

посередині», що намагаються перехопити облікові дані Office 365 (рисунок 3). А, зі свого боку, Outlook і Amazon здійснюють подальшу покрокову процедуру із великим відривом. Адже більшість компаній, що спостерігаються, мають шкідливий URL, вбудований в тіло листа, а не у вкладення. При цьому, мало хто з вбудованих URL адресується безпосередньо до сервера man-in-the-middle. Оскільки значна кількість адрес проходить через одне або кілька перенаправлень URL, перш ніж досягти кінцевого URL, який фактично проводить атаку man-in-the-middle.

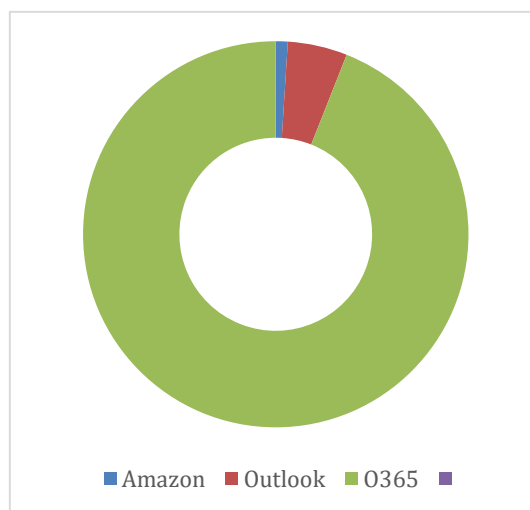


Рисунок 3 – Тип сторінок входу, що перехоплюються серверами типу man-in-the-middle

Висновки

Обґрунтовано, що MITM-атаки залишаються значною загрозою для стратегічних і цивільних інфраструктур. З'ясовано, що впровадження сучасних методів захисту, таких як шифрування, 2FA, VPN та IDS на основі математичного моделювання та штучного інтелекту може значно зменшити ризики атак. Передбачено та прогнозовано, що подальші дослідження, а також впровадження інноваційних рішень забезпечать ефективний рівень кібербезпеки в майбутньому. При цьому варто окреслити системний аналіз щодо кіберстійкості на рівні інтелектуальної особи, правової держави та інформаційного суспільства знань. Адже в цьому напрямку в кіберсучасності повоєнної України вже здійснюється академічний рейтинг глобально-регіонального характеру наукової західноєвропейської традиції, оскільки така інтелектуальна традиція не одне вже покоління мобілізує найкращий інтелект з усього світу (AD - рейтинг).

Список використаних джерел

1. Smith J. Man-in-the-Middle Attacks: An Overview and Prevention Strategies. *Cybersecurity Journal*. 2018.
2. Jones A. The Role of Encryption in Preventing MITM Attacks. *International Journal of Cryptography*. 2019.
3. Brown M. Critical Infrastructure Security: Lessons from Recent Attacks. Security in Practice. 2020.
4. Лісовський П. М., Лісовська Ю. П. Дискретна математика війни: кодери та декодер: навчальний посібник. Київ: Видавничий дім "Кондор", 2024. 112 с.
5. Лісовський П.М., Лісовська Ю.П. Воєнна стеганографія: квантова симуляція та космічна гіперспектроскопія: навч. посіб. Київ: Ліра-К, 2024. 134 с.
6. Лісовський П.М., Лісовська Ю.П. Кібервійська як квантове програмне забезпечення інформаційного капіталу: монографія. Київ: Ліра-К, 2024. 210 с.

7. Gray T. Man-in-the-Middle Attacks in IoT Systems. *Springer Advances in Computing*. 2022.

8. White P. and Green L. The Role of Artificial Intelligence in Detecting MITM Attacks. *IEEE Transactions on Cybernetics*, 2020. DOI: [10.1109/TCYB.2020.1234567](https://doi.org/10.1109/TCYB.2020.1234567)

9. Fernandez R. Advanced Threat Modeling Techniques for Mitigating MITM Attacks. Elsevier Computer Security Series. 2022.

10. Lee A., & Walker J. Ethical Hacking and Penetration Testing Guide. Auerbach Publications. 2021.

11. Johnson L. Network Defense and Countermeasures: Principles and Practices. Pearson. 2020.

References

1. Smith J. Man-in-the-Middle Attacks: An Overview and Prevention Strategies. *Cybersecurity Journal*. 2018.
2. Jones A. The Role of Encryption in Preventing MITM Attacks. *International Journal of Cryptography*. 2019.
3. Brown M. Critical Infrastructure Security: Lessons from Recent Attacks. Security in Practice. 2020.
4. Lisovskyi P. M., Lisovska Yu. P. Дискретна математика війни: кодери та декодер. Kyiv: Vydavnychiy dim "Kondor", 2024. 112 p. [in Ukrainian]
5. Lisovskyi P.M., Lisovska Yu.P. Voienna steganohrafiia: kvantova symuliatsiia ta kosmichna hiperspektroskopiiia. Kyiv: Lira-K, 2024. 134 p. [in Ukrainian]
6. Lisovskyi P.M., Lisovska Yu.P. Kiberviiska yak kvantove prohranne zabezpechennia informatsiinoho kapitalu: monohrafiia. Kyiv : Lira-K, 2024. 210 p.
7. Gray T. Man-in-the-Middle Attacks in IoT Systems. *Springer Advances in Computing*. 2022.
8. White P. and Green L. The Role of Artificial Intelligence in Detecting MITM Attacks. *IEEE Transactions on Cybernetics*, 2020. DOI: [10.1109/TCYB.2020.1234567](https://doi.org/10.1109/TCYB.2020.1234567)

9. Fernandez R. Advanced Threat Modeling Techniques for Mitigating MITM Attacks. Elsevier Computer Security Series. 2022.

10. Lee A., & Walker J. Ethical Hacking and Penetration Testing Guide. Auerbach Publications. 2021.

11. Johnson L. Network Defense and Countermeasures: Principles and Practices. Pearson. 2020.

METHODOLOGICAL MATRIX FOR PREVENTING MITM ATTACKS IN STRATEGIC AND CIVIL INFRASTRUCTURES

A. A. Shkitov^{1*}, N. G. Buzoveria²,
O. A. Stadnichenko³, K. M. Mazhuga³,
T. V. Humeniuk²

¹Open International University of Human Development
“Ukraine”; 23 Lvivska St., Kyiv, 03115, Ukraine;
e-mail: *office@uu.ua*

²Ivano-Frankivsk National Technical
University of Oil and Gas;
15 Karpatska St., Ivano-Frankivsk, 76019, Ukraine;
e-mail: *nadiia.shyrmovska@nung.edu.ua*

³Ukrainian Institute of Plant Variety Examination;
15 Henerala Rodymtseva Str, 15, Kyiv, 03041,
Ukraine; e-mail: *sops@sops.gov.ua*