

МЕХАНІЗМИ ТА МЕТОДИ ФІШИНГУ ЯК ПЕРШОГО КРОКУ ДО ОТРИМАННЯ ДОСТУПУ

Анотація. Розглянуто фішинг – техніку надсилання фішингових повідомлень. Аналіз зроблено на підставі даних у відкритому доступі. Проаналізовано процес фішингової атаки, та досліджено технічні вектори того, як користувачі стають жертвами атаки. Також розглянуто існуючі параметри фішингових атак та відповідні підходи до запобігання.

Ключові слова: фішинг, кібербезпека, багатфакторна аутентифікація, соціальна інженерія.

1. Проблема фішингу. Найпоширенішим підходом для запуску фішингової атаки є надсилання фішингового електронного листа. Згідно зі звітом Verizon [1] про розслідування витоку даних за 2023 рік, соціальна інженерія — це тактика, використання якої збільшилось з 14% 2022 року до 21% у 2023 році (рисунок 1).

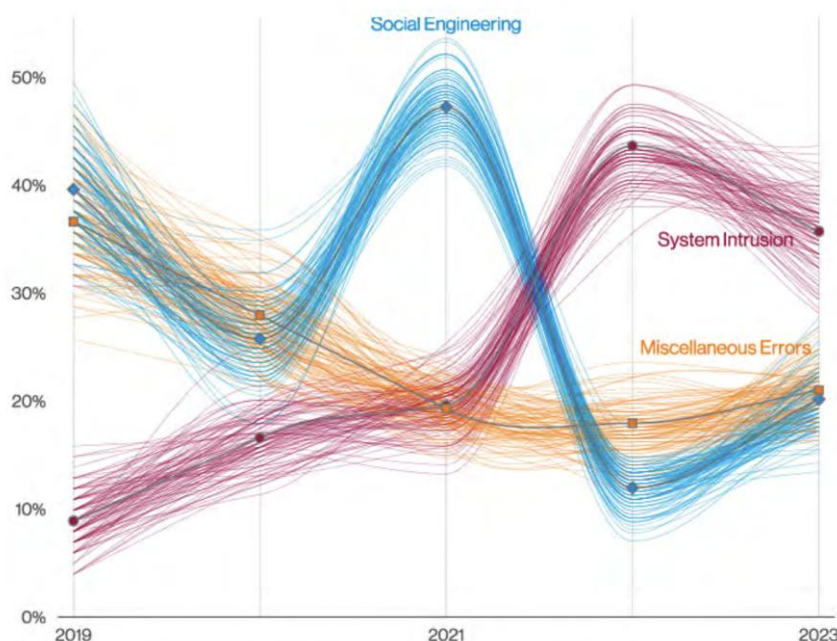


Рисунок 1 - Соціальна інженерія поміж інших атак [1]

Це зростання в першу чергу представлено фішинговими атаками, які виявилися в 18% зломів, і сценаріями претексту (4%). Половина випадків витоку

даних, у 2023 році була зафіксована з використанням життєвого циклу атаки соціальної інженерії: «розслідування», «побудова стосунків», «гра» та «вихід». У більшості випадків фішингові атаки соціальної інженерії відбувалися через електронну пошту, а саме 98% з зареєстрованих випадках, тоді як лише 2% за допомогою інших комунікацій, такі як телефон, соціальні мережі або внутрішній додаток для обміну повідомленнями, в якому деякі люди можуть розслабитися [1].

Атаки соціальної інженерії часто дуже ефективні та надзвичайно прибуткові для кіберзлочинців. Можливо, саме тому атаки компрометації бізнес-електронної пошти (BEC, Business Email Compromise) що, по суті, є атаками з використанням претекстів, майже подвоїлися з усього набору даних про інциденти, як можна побачити на малюнку нижче, і тепер становлять понад 50% інцидентів у шаблоні соціальної інженерії [1].

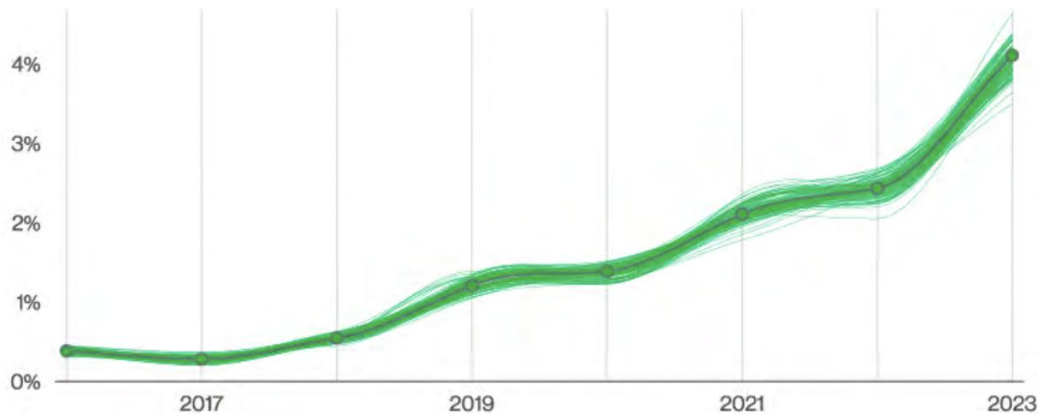


Рисунок 2 - Pretexting інциденти з часом [1]

Фішинг залишається дуже успішним способом для зловмисників отримати доступ до приватних даних і комп'ютерних систем. Що відбувається після цього першого електронного листа, розвиток може бути різним. Атаки зазвичай здійснюються двома основними шляхами. Найчастіше, якщо зловмисники вимагають або виманюють облікові дані та отримують їх, далі використовуватимуть ці облікові дані для доступу до папки «Вхідні» користувача (в 32% випадків). Інший розвиток зловмисники можуть вигадати достовірну історію (хоча й вигадану). Ціллю є переконання когось виконати бажання зловмисника. Переконання когось змінити банківський рахунок для заявленого одержувача, зустрічається в 56% випадків. Звичайно, також можна використовувати комбінацію тактик. Зловмисники можуть використати отриманий доступ до папки "Вхідні" користувача, щоб знайти ланцюжок електронної пошти, який вони

можуть захопити, або шукати в адресній книзі жертви людей, які можуть стати ціллю для подальших дій.

Фішинговий електронний лист часто додає зовнішнє посилання, яке перенаправляє жертв на підроблений веб-сайт і вимагає від жертви ввести свою особисту інформацію. Робочий процес фішингової атаки проілюстрували Джайн і Гупта [2], як показано на рисунку нижче. У цьому випадку фішинговий веб-сайт є незаконним клоном законного веб-сайту та має високу візуальну схожість. Ці фішингові сайти завжди включають деякі поля введення, наприклад текстове поле, щоб вимагати від жертви введення конфіденційної інформації. Інформація надсилається фішеру, коли жертви надсилають свої дані.

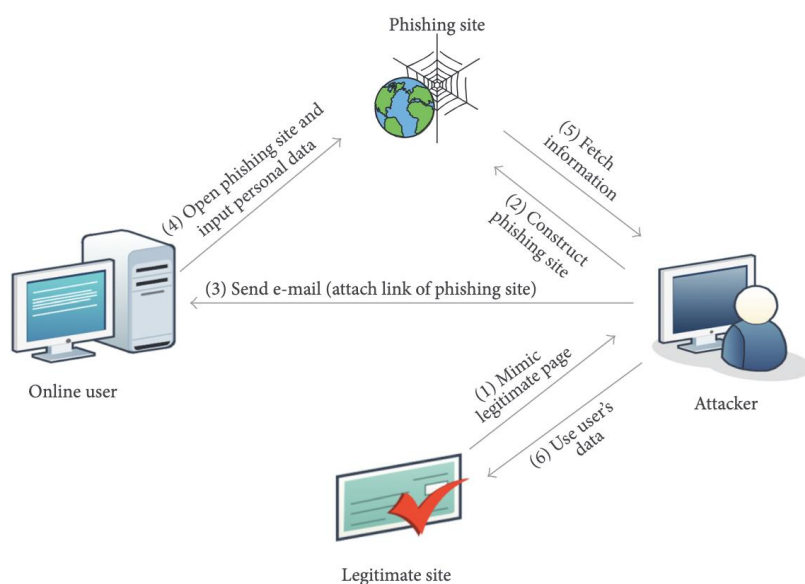


Рисунок 3 - Ілюстрація фішинг механізму [2]

Повна процедура того, як фішер запускає фішингову атаку, розділена на шість кроків.

- Крок 1. Фішери збирають і дублюють дані з добре відомого законного (цільового) веб-сайту.
- Крок 2. Фішери розробляють власний фішинговий веб-сайт відповідно до даних, зібраних на кроці 1.
- Крок 3. Фішери надсилають цей шкідливий веб-сайт кільком жертвам електронною поштою, зазвичай із посиланням у тексті електронного листа.
- Крок 4. Жертви клацають шкідливе посилання в електронному листі, переходять на фішинговий веб-сайт, вводять і надсилають свою особисту інформацію.
- Крок 5: дані жертв надсилаються фішеру з фішингового веб-сайту.

- Крок 6: Зрештою, фішери отримують доступ до цільового веб-сайту, використовуючи ідентифікаційну інформацію жертви.

2. Методи, які мають вплив на успішність фішингові атаки. Фішинг вивчається з початку 2000-х років, але проблема не була повністю вирішена, оскільки фішери постійно вдосконалюють свій підхід до атак [3]. Фішинговий веб-сайт відносно легко ідентифікувати шляхом спостереження за URL-шляхом, і уважний і досвідчений користувач тепер знає про ці підозрілі веб-сайти [4].

Фішер вводить в оману жертву фішингової атаки, видаючи себе за візуально подібний веб-сайт, використовуючи таку тактику [2]:

1. Візуальний вигляд. Фішинговий веб-сайт має високу візуальну схожість із законним веб-сайтом. Це пов'язано з тим, що фішери дублюють HTML-код законного веб-сайту, щоб створити свій фішинговий веб-сайт.

2. Адресний рядок. Фішери приховують URL-адресу за допомогою сценарію або зображення, у результаті чого жертва вважає, що вона перебуває на правильному веб-сайті.

3. Вбудовані об'єкти. Фішери використовують вбудовані об'єкти, такі як зображення, сценарії тощо, щоб приховати свій шкідливий текстовий вміст і HTML-код.

4. Подібність Favicon. Favicon – це піктограма зображення, пов'язана з певним веб-сайтом. Фішери можуть дублювати фавікон цільового веб-сайту, у результаті чого жертва буде вірити, що вони перебувають на правильному веб-сайті.

Крім того, особа, яка не обізнана з кібер безпекою, є основною причиною того, що вона потрапляє під фішингову атаку. Відповідно до репорту KnowBe4 Phishing By Industry Benchmarking Report 2023 Edition [5]:

- Користувачам бракує детальних знань про (справжню) URL-адресу.
- Користувачі не знають, якій веб-сторінці можна довіряти.
- Користувачі не перевіряють повну URL-адресу через переспрямування сторінки або приховані URL-адреси.
- Користувачі не мають багато часу, щоб перевірити URL-адреси, або вони випадково заходять на веб-сторінку
- Користувачі не можуть відрізнити фішингові веб-сторінки від законних.

3. Сучасні загрози фішингу. Наразі фішингові атаки є більш складними, оскільки фішери розвивають свій підхід до атак за допомогою різних методів. Наприклад, загальну фішингову атаку можна відносно легко ідентифікувати шляхом спостереження за URL-шляхом, і багато користувачів зараз знають про такий вид атаки через підозрілі URL-адреси або очевидну інформацію попередження з браузерів. Однак такі шкідливі URL-адреси можна приховати за допомогою різноманітних передових технологій, таких як використання скороченої URL-адреси чи QR-коду.

Крім того, запобігання фішинговим атакам на мобільних платформах є більш складним, ніж очікувалося. Поряд із тими ж проблемами, що й настільні комп'ютери, мобільні пристрої все ще стикаються з додатковими проблемами. Згідно з нашим підсумком, більшість фішингових посилань надходять із фішингових електронних листів, а мобільні платформи не підтримують безпечну ідентифікацію. Мобільний користувач не може знати, чи URL-адреса, до якої він отримав доступ, є безпечною, і, зокрема, чи користувачеві бракує достатньої обізнаності щодо безпеки. Наприклад, Google Chrome забезпечує набагато кращі механізми захисту від фішингових атак, ніж інші веб-браузери [6]. Він друкує сторінку попередження, яка показує, що URL-адреса, до якої відкривається доступ, містить потенційний ризик для користувачів у їхніх веб-переглядачах, якщо ця URL-адреса є зловмисною. Також Google анонсував про додавання нового захисту стандартної функції безпечного перегляду Google Chrome, увімкнувши захист від фішингу в реальному часі для всіх користувачів [7]. Хоча функція розширеного безпечного перегляду залишається незмінною та пропонує найкращий захист у Chrome, Google тепер додає захист у реальному часі до стандартної функції безпечного перегляду для підвищення безпеки.

Розробник браузера каже, що робить це, оскільки локальний список безпечного перегляду оновлюється лише кожні 30–60 хвилин, але 60% усіх фішингових доменів залишаються активними лише 10 хвилин. Це створює значний часовий проміжок, який залишає людей незахищеними від нових шкідливих URL-адрес. «Щоб заблокувати ці небезпечні сайти в момент їх запуску, ми оновлюємо безпечний перегляд, щоб тепер він перевіряв сайти на відомі шкідливі сайти Google у режимі реального часу», — каже Google.

«Завдяки скороченню часу між ідентифікацією та запобіганням загрозам ми очікуємо на 25% покращеного захисту від шкідливих програм і фішингових загроз».

Google повідомила, що функція Enhanced Safe Browsing за вмиканням зв'язується безпосередньо з протоколом Safe Browsing і надсилає додаткові дані. Хоча конфіденційність трохи менша, вона забезпечує найкращий захист, оскільки може виявляти шкідливі URL-адреси ще до того, як Google їх побачить.

Оскільки стандартна функція безпечного перегляду є опцією за замовчуванням, менеджер із продуктів Google Chrome Джасіка Бава повідомила BleepingComputer, що вони запроваджують захист у режимі реального часу з більшою мірою збереження конфіденційності через Fastly Oblivious HTTP Relays.

Протокол Oblivious передає частково хешовані URL-адреси користувачів системі безпечного перегляду Google, не розкриваючи особисту інформацію користувачів, таку як IP-адреси та заголовки запитів.

Однак ця стандартна функція безпечного перегляду в режимі реального часу, що забезпечує конфіденційність, має недолік. Оскільки він не надсилає стільки метаданих системі, він не зможе евристично визначити, чи є URL-адреса зловмисною, якщо її попередньо не помітить Google.

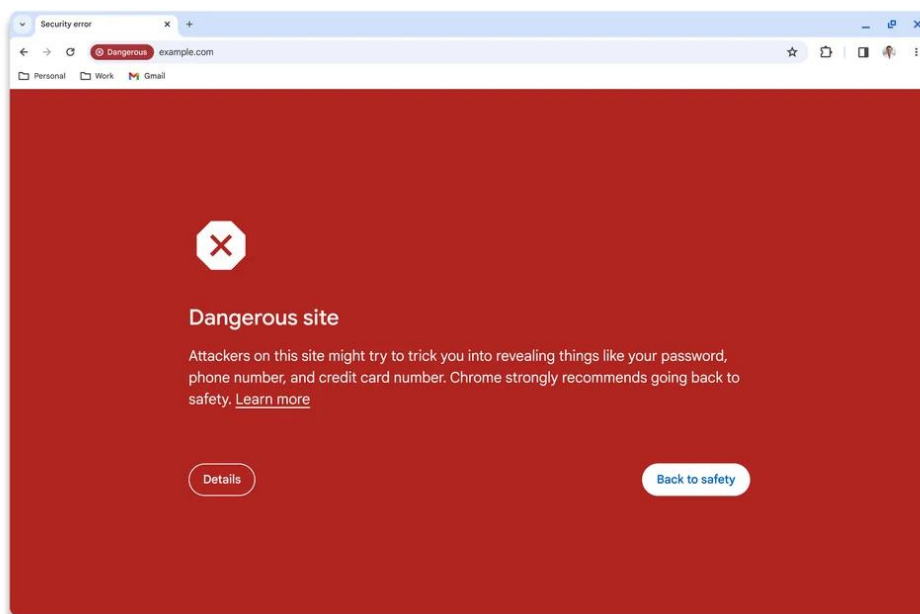


Рисунок 4 - Приклад роботи механізму захисту в Google Chrome [7]

4. Дослідження методів пом'якшення та захисту від фішингу. Виявлення фішингових атак є складною проблемою [5], оскільки ця атака використовує слабкі сторони людських характеристик, а не недоліки мережі. Для покращення захисту від фішингових атак можна виділити два напрями це або підвищення

ефективності технології виявлення фішингу, або розвиток освіти у користувачів. Загалом технологія виявлення фішингу зосереджена на ідентифікації зловмисних веб-сайтів, і існує два основні підходи, прийняті для пом'якшення фішингових атак: підхід на основі перевірки URL (за списком) та підхід на основі вмісту повідомлення. Схема на основі списків передбачає два типи списків: білий і чорний список. Ця схема є статичним підходом, у якому цільова URL-адреса порівнюється зі списками фішингу перед доступом до URL-адреси. У підході білого списку законні домени зберігаються в списку, і будь-який відповідний результат показує, що цільовий веб-сайт є законним веб-сайтом, і тому користувач може безпечно отримати доступ до цього веб-сайту. У підході до чорного списку шкідливі домени збираються в список, і будь-який відповідний результат показує, що цільовий веб-сайт, ймовірно, є фішинговим веб-сайтом, і користувачу буде надіслано сповіщення з попередженням у його браузері. Схема, заснована на вмісті, виявляє фішингові атаки шляхом вилучення певних функцій із цільової URL-адреси. Прикладом одного із методів є OpenPhish [8], що аналізує десятки мільйонів URL-адрес, щоб виявити фішинговий вміст. Цей звіт розбиває зміни в цільових брендах, галузях та фішинговій інфраструктурі. Дані нижче генеруються за допомогою їх бази даних фішингу.

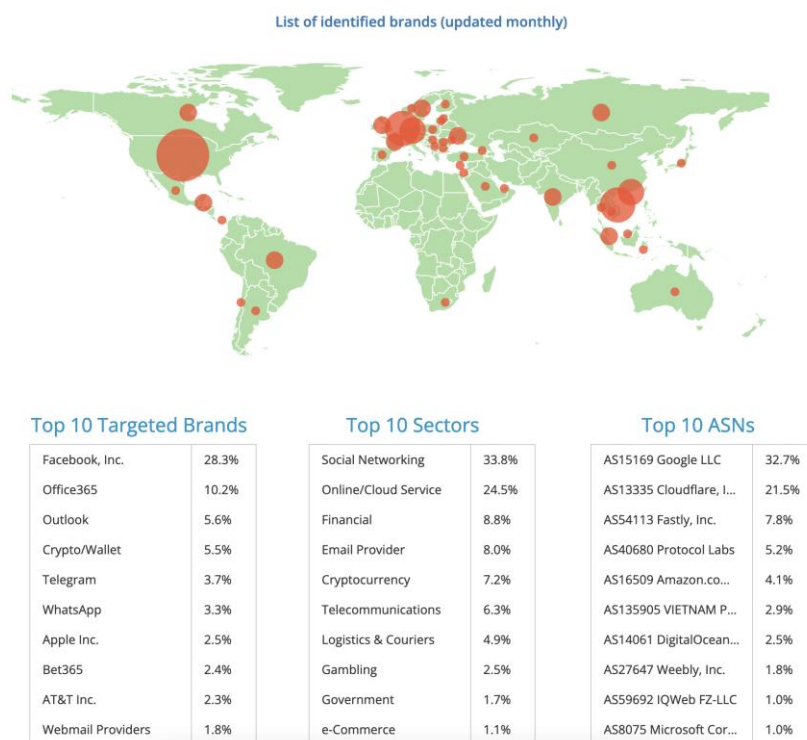


Рисунок 5 - Статистика брендів для impersonation [8]

5. Метод white black лістингу як захист від фішингу. В даний час різні дослідження базуються на цих двох підходах. Прихильники методи white black списків вважають, що перспективним рішенням є покращення продуктивності виявлення шляхом об'єднання чорних списків, перевірених вручну, з обчислювальними методами не лише для підвищення точності виявлення, але й для зменшення витрат часу на перевірку атак. Це рішення принесе користь сайтам, які підтримують служби фішингових чорних списків, наприклад PhishTank [9]. Комбінована техніка розглядалася у дослідженні «Smartening the Crowds: Computational Techniques for Improving Human Verification to Fight Phishing Scams» і складалась з чотирьох основних підходів; роблячи проблему непомітною для кінцевих користувачів, покращуючи дизайн інтерфейсів користувача, покращуючи навчання кінцевих користувачів [10]. Вони розробили систему, подібну до краудсорсингу, для виявлення фішингових веб-сайтів. Система вилучала неперевірені URL-адреси з PhishTank і фільтрувала URL-адреси, які не були в білому списку.

Вони розгорнули відповідні алгоритми (DBSCAN і shingling) для кластеризації схожих сторінок, щоб підвищити ефективність виявлення. Ці кластери були надіслані в Amazon Mechanical Turk для перевірки. Зрештою URL-адресу було визначено як фішингову або звичайною відповідно до ваги голосів від кожного учасника. Їхня система була швидшою за інші існуючі чорні списки. Крім того, їх рішення можна легко прийняти будь-яким існуючим чорним списком, перевіреним вручну, таким як ті, що пропонуються Google, Microsoft і PhishTank.

В дослідженні вони виявили, що велика кількість поточних фішингових веб-сайтів створюється за допомогою схожих інструментів, що призводить до створення схожих фішингових результатів через високу схожість вмісту. Таким чином, вони запропонували розширений метод ієрархічного чорного списку для виявлення фішингових атак шляхом використання існуючих зворотних списків. у їхньому підході для аналізу вмісту фішингових веб-сторінок було застосовано техніку n-grams на перевірених вручну чорних списках URL-адрес, а також одну майже дубльовану фішингову веб-сторінку було ідентифіковано ймовірнісним способом за допомогою методу shingling. Крім того, щоб зменшити кількість помилкових спрацьовувань, вони використали модуль фільтра, який додатково визначив легітимність потенційного фішингового веб-сайту за допомогою методів пошуку інформації в пошукових системах.

Підхід на основі списку не індексує лише цільову URL-адресу. Деякі дослідники [11] вважають, що більшість сучасних веб-сайтів складаються з кількох ресурсів, таких як CSS, JS та зображення. Однак ці законні (фірмові) веб-сайти зазвичай отримують вміст ресурсів з іншого домену через обмеження браузера щодо максимальної кількості одночасних підключень до того самого домену. Наприклад, законний веб-сайт PayPal отримує ресурси CSS, JS і файли зображень із paypalobjects.com. Таким чином, аналіз зв'язку запиту ресурсів є ефективним методом запобігання появі фішингових веб-сайтів. Geng та ін. запропонували новий підхід, який базується на зв'язках запитів ресурсів бренду майнінгу для виявлення фішингових атак. Їхній підхід не тільки ефективний і простий у застосуванні, але й є ефективним доповненням до існуючих методів. У їхньому рішенні цільова URL-адреса та всі запитувані домени перевірятимуться з чорного та білого списків.

6. Метод аналізу вмісту як захист від фішингу. З досліджень «Machine learning based phishing detection from URLs» [12] бачимо що було розроблено систему виявлення фішингу на основі навчання. Згідно з їхнім підходом, навчальні дані включають багато функцій, які стосуються як фішингу, так і законних класів веб-сайтів. Було вибрано сорок різних функцій на основі NLP, таких як підрахунок необроблених слів, перевірка домена, найбільша довжина слова, найкоротша довжина слова, TLD тощо. Крім того, сім різних алгоритмів були проведені в процесі навчання з метою вибору оптимального алгоритму, який має найвищу точність. Згідно з результатами цього підходу, алгоритм Random Forest показав найкращу продуктивність із найвищим рівнем точності виявлення серед цих семи алгоритмів для виявлення фішингових веб-сайтів. Niakanlahiji, Chi та Al-Shaer [13] запропонували масштабований багатофункціональний підхід машинного навчання для виявлення фішингових веб-сайтів під назвою «PhishMon». На відміну від інших підходів до виявлення фішингу на основі машинного навчання, вони вибрали п'ятнадцять нових функцій, які можна ефективно обчислити, використовуючи функції, які не вимагають взаємодії зі сторонніми службами, такими як пошукові системи та сервери WHOIS, таким чином зменшуючи час прийняття рішень. Функції були витягнуті з чотирьох аспектів: response HTTP, сертифікати SSL, документ HTML і файл JavaScript відповідно. Їхнє рішення забезпечує високу точність виявлення фішингових веб-сайтів, оскільки вибрані функції фіксують різні характеристики легітимних веб-додатків, а також їхні базові веб-інфраструктури.

7. MFA як додатковий рівень захисту. Поряд із вищезазначеними дослідженнями деякі дослідники вважали, що двофакторну автентифікацію можна застосувати для пом'якшення фішингових атак. Двофакторна автентифікація (також називається 2FA, MFA) — це механізм безпеки, який реалізує два вектори для автентифікації безпеки, і він вважається більш безпечним, ніж традиційна система автентифікації. Три загальновизнаних фактора автентифікації: те що ви знаєте або something that you know (наприклад, паролі), що у вас є або something that you have (наприклад, токени) і що ви є або what you are (біометрія) [14]. За останні кілька років 2FA реалізовано на більшості популярних веб-сайтів, таких як Google, Microsoft. 2FA допомагає користувачам захистити облікові записи; додатковий запит автентифікації надсилається користувачеві на іншому векторі, щоб перевірити та підтвердити додаткову ідентифікацію. Для пом'якшення фішингових атак 2FA може захистити обліковий запис користувача, навіть якщо фішери збирають ідентифікаційні дані користувача, фішер все одно не може отримати доступ до цільового веб-сайту, оскільки ідентифікацію користувача потрібно перевіряти за допомогою двох векторів. По суті, інші особи не можуть увійти в обліковий запис користувача без згоди користувача.

Однак лише 2FA не запобіжить успіху всіх фішингових атак. Існують різні підходи, які можуть обійти механізм 2FA. Наприклад, під час фішингової атаки, якщо жертва переходить на фішингову сторінку та вводить свої облікові дані, фішер може використовувати ці дані в режимі реального часу для входу на законний сайт [15]. Ілюстрацію цього сценарію показано на рисунку 6 нижче. Запит на код 2FA надсилається жертві (на кроці 4), потім ця жертва вводить код на фішинговому веб-сайті (на кроці 5). Згодом цей фішер використовує цей код для входу на справжній сайт як ідентифікований користувач. Зрештою, цей фішер може надіслати будь-що назад жертві, переконавши її, що веб-сайт, на який він отримав доступ, є законним, щоб зменшити її обізнаність про безпеку.

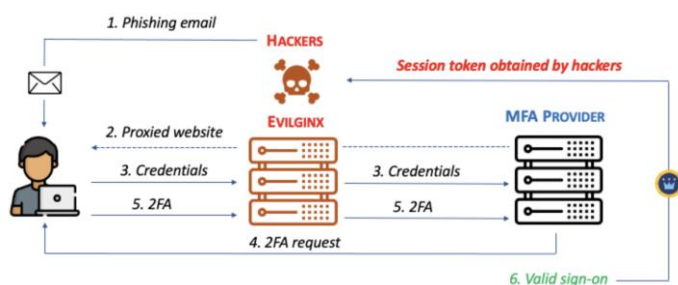


Рисунок 5 - Man-in-the-middle attack with evilginx [16]

Висновки. Пом'якшення фішингових атак є важливою темою дослідження, яку варто вивчити. Хоча було проведено багато досліджень, ця загроза все ще існує в реальному світі, поширеність якої постійно зростає. Згідно з результатами досліджень виявлення фішингових атак є складною проблемою. Для пом'якшення фішингових атак використовуються дві основні стратегії; або підвищення продуктивності технології виявлення фішингу, або покращення обізнаності людей про ці атаки. Розвиток людської досвідченості є основним способом подолання фішингових атак, оскільки фішингові атаки використовують слабкі сторони людських якостей, а не недоліки мережі. Крім того, люди завжди є найслабшою ланкою в атаках соціальної інженерії. Належне навчання є важливим, щоб гарантувати, що користувачі розуміють, як розпізнавати фішингові атаки, а якість навчання може вплинути на успіх запобігання фішингу. Таким чином, експеримент з тренінгами та навчанням вимагає від користувачів безпосередньої участі, оскільки якість навчання залежить від порівняння показників успішної ідентифікації фішингу в учасників до та після відповідного тренінгу. Однак для покращення продуктивності технології виявлення фішингу дослідники зосереджуються на двох аспектах пом'якшення фішингових атак; один базується на виявленні фішингових електронних листів, оскільки електронна пошта є найбільш вразливим засобом для запуску фішингової атаки. Таким чином, фішингові атаки будуть заблоковані з джерела, якщо буде виявлено фішинговий електронний лист. Інший спосіб зосереджений на виявленні фішингових веб-сайтів, оскільки більшість фішингових атак використовують фішинговий веб-сайт для незаконного збору даних жертви. Порівняно з виявленням фішингових веб-сайтів, виявлення фішингової електронної пошти може потребувати участі користувачів, щоб отримати кращий результат виявлення. Тому що успіх фішингової електронної пошти залежить від її контексту. Зокрема, коли передумова фішингового електронного листа узгоджується з робочим контекстом користувача (або поточною ситуацією). Таким чином, у цьому випадку досвід користувача та поточна ситуація є релевантними характеристиками, і ці змінні матимуть високу вагу, якщо дослідники запровадять машинне навчання для навчання відповідних даних. Крім того, для оцінювання потрібні учасники з різним досвідом для охоплення різноманітних ситуацій і тестів. Більшість антифішингових рішень впроваджено для пом'якшення загальних фішингових атак, але вони ігнорують деякі конкретні ситуації, наприклад розширені фішингові атаки. Щоб запобігти

розширеним фішинговим атакам, фішингові веб-сайти важко виявити, якщо жертва піддається атаці з використання викрадених DNS даних, оскільки вміст URL-адреси та вміст веб-сайту є такими самими, як і законні веб-сайти. Більшість підходів, заснованих на вмісті, можуть не спрацювати, оскільки вміст URL-адреси, до якої здійснюється доступ, є важливим фактор у прийнятті рішення. Однак дійсний і надійний сертифікат SSL неможливо підробити. Таким чином, ми розглядаємо визначення узгодженості сертифіката SSL між веб-сайтом, до якого ви отримали доступ, і законним веб-сайтом, щоб ідентифікувати фішинговий веб-сайт, який перебуває під атакою викрадення DNS. Крім того, під час атаки викрадення Інтернет-провайдера (ISP hijacking) фішери можуть не фішингувати жертву безпосередньо з законного веб-сайту, але вони можуть заманити жертв, скеровуючи їх на фішинговий веб-сайт, вставляючи переконливий контекст на законному веб-сайті. Щоб запобігти атакам захоплення субдоменів, фішинговий веб-сайт важко виявити, якщо фішери розмістили цей веб-сайт у субдоміні, взятому з законного веб-сайту. Незалежно від веб-вмісту, URL-адреси та інформації сертифіката SSL, усі вони будуть такими самими, як і законний веб-сайт. Більше того, підхід до перебору субдоменів потребує вдосконалення, оскільки більшість поточних інструментів базується на грубому переборі, існуючі словники можуть не охоплювати всі випадки субдоменів, оскільки деякі субдомени можуть бути безглуздими. Таким чином, ми бачимо що ще є місце для покращення та вдосконалення захисту від фішингових атак.

ЛІТЕРАТУРА / REFERENCES

1. Verizon, “2023 Data Breach Investigations Report”. 2023. [Online]. Access: <https://www.verizon.com/business/en-gb/resources/reports/dbir/>
2. A. K. Jain and B. B. Gupta, “Phishing detection: Analysis of visual similarity based approaches”. In Journal of Security and Communication Networks, vol. 2017, Article ID. 5421046, pp. 1-20, Jan 2017. DOI: 10.1155/2017/5421046.
3. Zainab Alkhalil, Chaminda Hewage, Liqaa Nawaf and Imtiaz Khan, “Phishing Attacks: A Recent Comprehensive Study and a New Anatomy” 2021. DOI: 10.3389/fcomp.2021.563060
4. F. Castaño et al., “Evaluation of state-of-art phishing detection strategies based on machine learning”. 2021. DOI: 10.18239/jornadas_2021.34.06.
5. Report Phishing by industry benchmarking report 2023. Access: <https://info.knowbe4.com/en-us/phishing-by-industry-benchmarking-report>

6. HKCERT, Browser's Anti-phishing feature: What is it and how it helps to block phishing attack? Access: <https://www.hkcert.org/blog/browser-s-anti-phishing-feature-what-is-it-and-how-it-helps-to-block-phishing-attack>
7. Google is enabling Chrome real-time phishing protection for everyone. Доступ: <https://www.cnet.com/news/privacy/google-chrome-can-now-warn-you-in-real-time-if-youre-getting-phished/>
8. Phishing feeds. Access: <https://openphish.com/>
9. PhishTank. Access: <https://phishtank.org/>
10. G. Liu, G. Xiang, B. A. Pendleton, J. I. Hong, and W. Liu, "Smartening the crowds: Computational techniques for improving human verification to fight phishing scams". 2011. DOI: 10.1145/2078827.2078838.
11. G. G. Geng, Z. W. Yan, Y. Zeng, and X. B. Jin, "RRPhish: Anti-phishing via mining brand resources request". 2018. DOI: 10.1109/ICCE.2018.8326085.
12. O. K. Sahingoz, E. Buber, O. Demir, and B. Diri, "Machine learning based phishing detection from URLs." 2019. DOI: 10.1016/j.eswa.2018.09.029
13. A. Niakanlahiji, B. T. Chu, and E. Al-haer, "PhishMon: A machine learning framework for detecting phishing webpages". 2018. DOI: 10.1109/ISI.2018.8587410.
14. M. Papathanasaki, L. Maglaras, N. Ayres, "Modern Authentication Methods: A Comprehensive Survey" 2022 DOI:10.5772/acrt.08
15. Evilginx - Bypassing MFA, phishing is back on the menu. Access: <https://bleekseeks.com/blog/evilginx-bypassing-mfa-phishing-is-back-on-the-menu>
16. How hackers beat MFA at-scale. Access: <https://www.mantra.ms/blog/beating-mfa>

Received 11.05.2023.

Accepted 16.05.2023.

Phishing like the first step to gaining access

Phishing as a term that means the technique of sending phishing messages will be researched based on findings in public access and using the listed links. The process of a phishing attack will be analyzed, and then we will pay attention to the technical vectors of how users become victims of the attack. Finally, existing research on phishing attacks and related prevention approaches will be reviewed.

Mitigating phishing attacks is an important research topic worth exploring. Although a lot of research has been done, this threat still exists in the real world, and its prevalence is constantly increasing. According to research results, detecting phishing attacks is a difficult problem. There are two main strategies used to mitigate phishing attacks; or improving the performance of phishing detection technology or improving peo-

ple's awareness of these attacks. Developing human expertise is a key way to defeat phishing attacks, as phishing attacks exploit human weaknesses rather than network weaknesses. Also, humans are always the weakest link in social engineering attacks.

Compared to phishing website detection, phishing email detection may require user involvement to achieve better detection results. Because the success of a phishing email depends on its context. Specifically, when the premise of the phishing email is consistent with the user's work context (or current situation).

Most anti-phishing solutions are implemented to mitigate general phishing attacks, but they ignore some specific situations, such as advanced phishing attacks. To prevent advanced phishing attacks, phishing websites are difficult to detect if a victim is attacked using stolen DNS data because the URL content and website content are the same as legitimate websites. Most content-based approaches may not work because the content of the accessed URL is an important factor in the decision.

To prevent subdomain hijacking attacks, it is difficult to detect a phishing website if the phishers have hosted the website on a subdomain taken from a legitimate website. Regardless of the web content, URL, and SSL certificate information, they will all be the same as the legitimate website. Moreover, the approach to enumeration of subdomains needs improvement, as most current tools are based on rough enumeration, existing dictionaries may not cover all instances of subdomains, as some subdomains may be meaningless.

Key words: phishing, cyber security, multifactor authentication, social engineering.

Гуда Антон Ігорович – д.т.н, проф., професор кафедри ІТС, ІПБТ УДУНТ.

Кліщ Сергій Михайлович – асистент кафедри ІТС, ІПБТ УДУНТ.

Guda Anton – doctor of engineering's sciences, professor, IIBT.

Klishch Sergey – post-graduate student, assistant, IIBT.