

<https://doi.org/10.31891/2219-9365-2025-81-55>

УДК 005.21:005.8:004.8.

БОХОНЬКО Олександр

Хмельницький національний університет

<https://orcid.org/0000-0002-7228-9195>

ЛИСЕНКО Сергій

Хмельницький національний університет

<https://orcid.org/0000-0001-7243-8747>

e-mail: sirogyk@ukr.net

МОДЕЛІ АТАК СОЦІАЛЬНОЇ ІНЖЕНЕРІЇ

З розвитком сучасних технологій Інтернет став ключовим для обміну різноманітною інформацією та комунікацій. Як наслідок, подібна еволюція принесла децентралізований доступ до даних та інформації через обмін файлами за допомогою платформ, зокрема таких як соціальні мережі, які як правило, не є досить захищеними.

Стаття присвячена дослідженню моделей атак соціальної інженерії, що використовують як психологічні особливості людської поведінки, так і технічні вразливості інформаційних систем. Автори інтегрують підходи з психології, соціології та кібербезпеки, що дозволяє глибше зрозуміти природу цих атак і розробити ефективніші засоби їх запобігання.

В статті описано ключові форми атак соціальної інженерії, включаючи фішинг, вішинг та інші методи, які маніпулюють довірою і викликають емоційні реакції у жертв. Значну увагу приділено аналізу основних векторів атак, специфіці цільових характеристик та психологічних тригерів, які зловмисники використовують для досягнення своїх цілей.

Запропоновані моделі атак охоплюють весь життєвий цикл атакуючих дій – від етапу початкової розвідки та збору інформації про жертву до впровадження методів маніпуляції та подальшого аналізу результатів після атаки. Окрім цього, моделі враховують поведінкові та психологічні аспекти, такі як ефект авторитету, страх, терміновість та інші фактори, що дозволяють зловмисникам успішно маніпулювати жертвами. У статті розглянуто вплив цих факторів на здатність організацій розпізнавати та нейтралізувати загрози на різних етапах їхнього життєвого циклу.

Результати дослідження мають на меті підвищити точність прогнозування атак, сприяти покращенню навчальних програм для співробітників, підвищити рівень обізнаності про ризики соціальної інженерії та посилити загальний рівень інформаційної безпеки.

Автори пропонують рекомендації щодо створення політик та інструментів захисту, які можуть значно зменшити вплив атак соціальної інженерії на організації та окремих осіб.

Ключові слова: соціальна інженерія, психологічні тригери, кібербезпека, моделі атак, інформаційна безпека.

BOKHONKO Oleksandr, LYSENKO Sergiy

Khmelnytskyi National University

SOCIAL ENGINEERING ATTACKS MODELS

With the development of modern technologies, the Internet has become the key to the exchange of various information and communications. As a result, such an evolution has brought decentralized access to data and information through file sharing through platforms, in particular such as social networks, which are generally not sufficiently secure.

The article is devoted to the study of social engineering attack models that use both psychological features of human behavior and technical vulnerabilities of information systems. The authors integrate approaches from psychology, sociology, and cyber security, which allows for a deeper understanding of the nature of these attacks and the development of more effective means of preventing them.

The article describes key forms of social engineering attacks, including phishing, vishing, and other methods that manipulate trust and elicit emotional responses from victims. Considerable attention is paid to the analysis of the main attack vectors, the specificity of target characteristics and psychological triggers that attackers use to achieve their goals.

The proposed attack models cover the entire life cycle of attacking actions - from the stage of initial reconnaissance and gathering information about the victim to the implementation of manipulation methods and further analysis of the results after the attack. In addition, the models take into account behavioral and psychological aspects such as the effect of authority, fear, urgency and other factors that allow attackers to successfully manipulate victims. The article examines the influence of these factors on the ability of organizations to recognize and neutralize threats at various stages of their life cycle.

The results of the study are intended to improve the accuracy of attack prediction, contribute to the improvement of training programs for employees, increase the level of awareness of the risks of social engineering and strengthen the overall level of information security.

The authors offer recommendations for creating policies and protection tools that can significantly reduce the impact of social engineering attacks on organizations and individuals.

Keywords: social engineering, psychological triggers, cyber security, attack models, information security.

ПОСТАНОВКА ПРОБЛЕМИ У ЗАГАЛЬНОМУ ВИГЛЯДІ ТА ЇЇ ЗВ'ЯЗОК ІЗ ВАЖЛИВИМИ НАУКОВИМИ ЧИ ПРАКТИЧНИМИ ЗАВДАННЯМИ

Створення моделей атак соціальної інженерії має вирішальне значення для розуміння, виявлення та пом'якшення різноманітних тактик, що розвиваються, які використовують кіберзлочинці для маніпулювання людьми та атак на організації. З розвитком технологій атаки соціальної інженерії стають більш витонченими та частими [1-3]. Зловмисники постійно вдосконалюють свої методи, тому виявляти

та запобігати таким загрозам стає дедалі складніше. Комплексні моделі допомагають зрозуміти нюанси та розвиваються стратегії, створюючи більш ефективний захист [4-6].

Атаки соціальної інженерії включають широкий спектр тактик, кожна з яких використовує різні аспекти людської психології та технологічні вразливості [7-9]. Від вішингу та фішингу до більш спеціалізованих форм, як “spear phishing” і “water-holing”, кожен вектор атаки вимагає певних стратегій виявлення та реагування. Моделювання цих атак допомагає точно класифікувати їх і розробити цільовий захист [10-12]. Такі атаки використовують людський елемент, часто націлюючись на такі емоції, як страх, цікавість, довіра чи терміновість [13-15]. Розуміння технік психологічної маніпуляції, які використовуються в цих атаках, має вирішальне значення для навчання окремих осіб та організацій [16-18]. Моделі допомагають у визначенні загальних психологічних тактик і розробці навчальних програм, які підвищують обізнаність і стійкість користувачів [19-21].

Економічні та репутаційні наслідки атак соціальної інженерії можуть бути руйнівними [22-24]. Для підприємств такі атаки можуть призвести до фінансових втрат, витоку даних і шкоди довірі клієнтів [25]. Моделі допомагають організаціям зрозуміти потенційні ризики та розробити проактивні заходи для пом'якшення цих впливів [26-27].

Зі збільшенням нормативних вимог щодо захисту даних і кібербезпеки організації зобов'язані впроваджувати надійні заходи безпеки [28]. Моделі атак соціальної інженерії допомагають узгодити політику організації з правовими вимогами, забезпечуючи відповідність і мінімізуючи ризик юридичних наслідків [29-30].

Розробка моделей атак соціальної інженерії допомагає створювати автоматизовані системи виявлення, які можуть ідентифікувати загрози та реагувати на них у реальному часі [31]. Це особливо важливо для великих організацій, де ручний моніторинг є недоцільним. Моделі забезпечують основу для алгоритмів, які можуть аналізувати шаблони, виявляти аномалії та позначати потенційні атаки [32-34].

Атаки соціальної інженерії впливають на різні сектори, включаючи фінанси, охорону здоров'я, освіту та уряд. Моделюючи ці атаки, організації в різних галузях можуть обмінюватися ідеями та найкращими практиками, сприяючи спільному підходу до кібербезпеки [35]. Ландшафт кібербезпеки є динамічним, у ньому регулярно з'являються нові тактики соціальної інженерії. Моделі дозволяють здійснювати безперервний моніторинг і адаптацію, забезпечуючи ефективність захисту від останніх загроз [34]. Моделі служать освітнім ресурсом як для фахівців з кібербезпеки, так і для широкої громадськості. Вони надають структуровані знання, які можна використовувати в навчальних програмах, семінарах та кампаніях з підвищення обізнаності та для підвищення загальної грамотності з питань кібербезпеки [35].

Таким чином, розуміння особливостей кожного типу атак соціальної інженерії дозволяє організаціям ефективніше розподіляти ресурси. Моделі допомагають визначити пріоритети зусиль на основі ймовірності та потенційного впливу різних векторів атак, оптимізуючи використання заходів безпеки та персоналу [36].

ФОРМУЛЮВАННЯ ЦІЛЕЙ СТАТТІ

Метою дослідження є новий підхід до створення моделі атаки соціальної інженерії. Розроблені моделі мають враховувати всі аспекти завивки атак для розробки ефективних методів виявлення в майбутньому.

ВИКЛАД ОСНОВНОГО МАТЕРІАЛУ

Відомі моделі виявлення атак соціальної інженерії

Різні дослідження досліджували моделі виявлення атак соціальної інженерії за допомогою різних підходів. Наприклад, одна запропонована модель [37] використовує методи сканування та класифікації для автоматичного визначення шкідливих посилань, навіть якщо сторінки вимагають часткового оновлення або введення облікових даних автентифікації. Ця модель продемонструвала вищу точність, досягнувши до 72% ефективності в користувальницьких додатках порівняно з існуючими методами.

У [38] пропонується нова модель, яка використовує текст підозрілих веб-сторінок замість їхніх URL-адрес для виявлення фішингових атак, використовуючи алгоритми обробки природної мови (NLP) і глибокого навчання (DL). Зокрема, для захоплення семантичних і синтаксичних особливостей тексту використовується шар вбудовування Keras із Global Vectors for Word Representation (GloVe).

У [39] автори розглядають проблему аналізу фішингових дій на веб-сторінках. Вони пропонують структуру машинного навчання, яка включає різні контрольовані алгоритми навчання, такі як випадкові ліси, опорні векторні машини та дерева рішень. Автори виконують оптимізацію гіперпараметрів для цих алгоритмів за допомогою фреймворків машинного навчання, таких як scikit-learn. Нарешті, вони обговорюють практичну реалізацію класифікації фішингових атак за допомогою реального набору даних, доступного на Kaggle.

Дослідження [40] моделює та оцінює принципи переконання в атаках соціальної інженерії (SET)

шляхом розрахунку ймовірності атаки (AOP) за допомогою моделі дерева атак та ймовірності успіху атаки (ASP) за допомогою моделі ланцюга Маркова. Він використовує дані про частоту SET, модальності та принципи переконання, щоб перевірити модель на різних прикладах SET. Результати свідчать про те, що модель ефективна для розуміння та ранжирування SET.

Дослідження [41] пропонує системний підхід до створення засобів протидії на основі типового процесу атак соціальної інженерії. Автори систематично аналізують кожен крок атаки, щоб запобігти, пом'якшити або усунути її, в результаті чого розроблено 62 контрзаходи. Вони також розробили набір шаблонів безпеки соціальної інженерії, які містять відповідні знання для надання практичної допомоги в аналізі захисту від таких атак.

Стаття [42] розглядає поширені типи кібершахрайства, включаючи соціальну інженерію, фішинг і DoS-атаки. Автор досліджує модель когнітивного обчислення для виявлення підозрілих банківських транзакцій, включаючи квантові обчислення для майбутніх застосувань. У статті оцінюються плюси, мінуси та результати моделі. Дослідник пропонує використовувати прогнозне моделювання з такими методами, як логістична регресія, дерева рішень і випадкові ліси для виявлення шахрайства в банківських операціях у реальному часі.

У статті [43] пропонується підхід, спрямований на захист бірж криптовалют та їхніх клієнтів від атак. Цей підхід формалізовано як синтез динамічної моделі протидії фішинговим атакам і моделі “персептрона”, представленій у вигляді базової штучної нейронної мережі. Динаміка протидії описується системою диференціальних рівнянь, які визначають зміни станів жертви фішингової атаки та зловмисника, який ці атаки організовує. Запропонований нейро-ігровий підхід ефективно прогнозує процес протидії фішингу, враховуючи витрати, пов'язані з різними стратегіями, які використовують залучені сторони.

[44] надає комплексний аналіз фішингових атак і методів виявлення фішингу на основі HTML і URL-адрес. Автор розглядає сучасні моделі глибокого навчання для виявлення фішингових атак на основі URL-адрес і гібридних атак. Дослідження порівнює кожен модель з точки зору попередньої обробки даних, вилучення функцій, дизайну моделі та продуктивності.

Стаття [45] містить таксономію алгоритмів глибокого навчання для виявлення фішингу через огляд 81 вибраних джерел. Дослідник представляє концепції фішингу та глибокого навчання, а потім класифікує існуючу літературу за категоріями, пов'язаними з виявленням фішингу та алгоритмами глибокого навчання. У статті оцінюються сучасні методи глибокого навчання, висвітлюються їхні сильні та слабкі сторони, а також обговорюються такі проблеми, як ручне налаштування параметрів, тривалий час навчання та низька точність розпізнавання.

Стаття [46] підкреслює потенціал великих мовних моделей (LLM) для посилення захисту від кібератак. Використовуючи свої глибокі лінгвістичні здібності, LLM може аналізувати розмови, виявляти ледве помітні ознаки вішинга та створювати адаптивні контрзаходи в режимі реального часу, пропонуючи багатообіцяючий підхід до покращення кібербезпеки в цій сфері.

У дослідженні [47] описується процес моделювання кібератак шляхом визначення його етапів за допомогою тематичного моделювання. Цей процес було застосовано для моделювання кібератак грумінгу та залякування, виявивши, що зловмисники використовують методи психологічної маніпуляції. Запропонований підхід дозволяє точно визначити комунікативні наміри зловмисників. Крім того, дослідження представляє функціональний прототип системи батьківського контролю, який підтримує запропонований процес моделювання.

Стаття [48] спрямована на покращення виявлення клонів для активістів соціальних медіа та жертв шляхом розробки методів моделювання, прогнозування та уникнення клонів. Запропонована модель використовує ансамблевий адаптивний випадковий підпростір з k -найближчими сусідами (k -NN) як базовий класифікатор і використовує теорію зважених графів для аналізу вузлів-клонів. Коли вага вузла перевищує певний поріг, довіра припиняється, а клони ранжуються та встановлюються пріоритети для видалення.

Стаття [49] зосереджена на тому, як навіть безпечний текст, згенерований великими мовними моделями (LLM), можна легко перетворити на потенційно шкідливий вміст за допомогою атак типу «приманки та перемикавання». Під час таких атак користувач спочатку ставить нешкідливі запитання, а потім використовує техніку «знайти та замінити», щоб перетворити відповіді на зловмисні розповіді. Тривожна ефективність цього підходу щодо створення токсичного вмісту підкреслює значну проблему в розробці надійних заходів безпеки для LLM. Дослідження також підкреслює, що зосередження виключно на безпеці результатів LLM є недостатнім і що слід також враховувати перетворення після обробки.

Стаття [50] окреслює основні причини кібератак і розглядає останні моделі атак і методи виявлення. У ньому обговорюються як технічні, так і нетехнічні рішення для раннього виявлення атак, наголошується на потенціалі нових технологій, таких як машинне навчання, глибоке навчання, хмарні платформи, великі дані та блокчейн у боротьбі з поточними та майбутніми кіберзагрозами.

Стаття [51] містить детальний аналіз впливу та еволюції програми-вимагача на різних етапах до її поточної форми. В ній автор також визначає найпоширеніші тактики, техніки та процедури (TTP) поширених сучасних варіантів програм-вимагачів на основі аналізу відповідно до структури MITRE ATT&CK.

Стаття [52] поєднує аналіз поточного ландшафту загроз кібербезпеці з розробкою емпіричної моделі загроз, наголошуючи на компоненті монетизації. Він демонструє результати обробки статистичних даних, які ілюструють розподіл конкретних загроз.

Дослідження [53] підкреслює, що різні провідні захисні механізми можна одночасно обійти, використовуючи прості адаптивні стратегії, навіть у рамках встановлених моделей загроз та з обмеженими можливостями супротивника (наприклад, використовуючи легко виявлені тригери, зберігаючи при цьому високий рівень успіху атаки). Зокрема, дослідження представляє нову методологію, яка покращує специфічність тригерів і коригує навчання в симбіотичний спосіб.

У статті [54] представлено метод глибокого тестування, призначений для виявлення троянів, які становлять значну загрозу для користувачів середнього та досвідченого рівня, зосереджуючись на незаконних інструкціях. Приховані інструкції можуть використовувати малоімовірну послідовність звичайних інструкцій як прелюдію, за якою слідує незаконна інструкція, яка активує троян. Такий підхід дозволяє трояну залишатися глибоко прихованим у процесорі.

Дослідження [55] являє собою першу спробу застосувати великі мовні моделі (LLM) для створення цільових бічних фішингових електронних листів, спрямованих на університет вищого рівня та його персонал (приблизно 9000 осіб) протягом 11-місячного періоду. Він також оцінює здатність інфраструктури фільтрації електронної пошти виявляти такі спроби фішингу, згенеровані LLM, надаючи розуміння їх ефективності та визначаючи потенційні області для вдосконалення.

Дослідження [56] порівнює ефективність фішингових електронних листів, створених автоматично за допомогою GPT-4, із тими, що створюються вручну за допомогою V-Triad. Воно також оцінює об'єднаний потенціал GPT-4 і V-Triad. Четверта група, яка отримувала типові фішингові листи, служила контрольною групою. Використовуючи підхід «червоної команди», дослідники змоделювали нападників і надіслали електронні листи 112 учасникам, відібраним для дослідження. Чотири популярні моделі великих мов (GPT, Claude, PaLM і LLaMA) використовувалися для виявлення фішингових намірів і порівняння результатів із виявленням людиною. Дослідження вивчало економічні аспекти фішингових атак, керованих штучним інтелектом, демонструючи, як великі мовні моделі підвищують стимули для фішингу та фішингу шляхом зниження вартості.

Стаття [57] представляє пізню модель злиття для гібридної системи фільтрації спаму з текстом і зображенням, порівнюючи її з традиційною ранньою моделлю злиття на основі оптичного розпізнавання символів (OCR). Згорточні нейронні мережі (CNN) і безперервний пакет слів були використані для вилучення ознак із текстових і графічних компонентів гібридного спаму відповідно. Результати підкреслюють переваги використання CNN над OCR для виявлення гібридного спаму.

Стаття [58] підкреслює серйозний вплив спаму на користувачів, зазначаючи, що оманливі електронні листи спаму можуть призвести до крадіжки особистої інформації або завантаження зловмисного програмного забезпечення, видаючи себе за законне. У документі обговорюються потенційні наслідки, такі як фінансові втрати та витоки даних, і розглядаються різні моделі машинного та глибокого навчання, запропоновані для боротьби з цими проблемами.

Стаття [59] досліджує використання вмісту електронної пошти, крім URL-адрес, для виявлення спаму. Зазначається, що хоча кілька методів існує, їм все ще бракує високої точності. Щоб вирішити цю проблему, у документі пропонується узагальнений дворівневий комплексний підхід для виявлення спаму, який не залежить від типу набору даних. Цей підхід порівнюється з різними методами машинного навчання, включаючи SVM, регресію та kNN, з прогнозами, зробленими за допомогою алгоритмів К-декомпозиції та голосування.

У статті [60] розглядається впровадження інтелектуальної системи моніторингу, призначеної для виявлення падінь у людей похилого віку за допомогою технології IoT. Система спрямована на підвищення безпеки та благополуччя людей похилого віку, забезпечуючи моніторинг у режимі реального часу та сповіщення у разі падіння, пропонуючи таким чином своєчасну допомогу.

Аспекти атак соціальної інженерії для створення моделі

Щоб підвищити ефективність виявлення атак соціальної інженерії, створення її моделей має включати кілька ключових кроків і міркувань. Оскільки атаки соціальної інженерії використовують як людську психологію, так і технічні вразливості, моделювання цих атак вимагає міждисциплінарного підходу, що поєднує знання психології, соціології та кібербезпеки.

Сьогодні атаки соціальної інженерії можуть приймати різні форми, включаючи фішинг, претекстинг, цькування, переслідування тощо. Кожен тип включає в себе різні техніки психологічної маніпуляції, щоб ввести в оману людей, щоб вони розголосили конфіденційну інформацію або вчинили певні дії [68-70].

Для побудови адекватних моделей атак необхідно враховувати:

1. *Вектори атак*, тобто шляхи, за допомогою яких зловмисники можуть досягти своїх цілей, як-от електронна пошта, телефонні дзвінки, соціальні мережі або особиста взаємодія.

2. *Цілі*, такі як особи, організації або системи. Ідентифікація цілі допомагає зрозуміти потенційний

вплив і ймовірні методи атаки.

Важливим ядром атак соціальної інженерії є поведінкові та психологічні аспекти:

1. *Психологічні тригери*. Ми повинні залучати до моделей психологічні тригери, які використовують соціальні інженери, такі як страх, жадібність, цікавість або бажання бути корисними, що має вирішальне значення. Ці тригери часто використовуються для створення відчуття терміновості або довіри.

2. *Людський фактор*. Модель повинна враховувати характеристики потенційних жертв, такі як рівень їхньої обізнаності, досвід і сприйнятливість до тактик соціальної інженерії.

Комплексна модель повинна охоплювати весь *життєвий цикл* атаки соціальної інженерії:

1. *Розвідка*. Це передбачає збір інформації про ціль.

2. *Взаємодія*. Взаємодія з ціллю для зміцнення довіри.

3. *Експлуатація*. Використання зібраної інформації та довіри для досягнення цілей зловмисника.

4. *Виконання*. Фактичне порушення або компрометація інформації.

5. *Пост аналіз атаки*. Розуміння наслідків і будь-яких подальших дій.

Щоб довести адекватність моделей атак, необхідно змоделювати різні сценарії атак, які мають допомогти зрозуміти потенційний вплив і ефективність різних засобів захисту, а також виконати перевірку моделі шляхом порівняння прогнозів моделі з даними реального світу, щоб забезпечити точність і надійність.

Ці особливості підкреслюють складність створення математичних моделей атак соціальної інженерії. Такі моделі є цінними для прогнозування потенційних атак, навчання окремих осіб і організацій і розробки більш ефективних засобів захисту.

Моделі атак соціальної інженерії

Визначимо множину S як множину атак соціальної інженерії, $S = \{\alpha, \beta, \gamma, \delta, \varepsilon, \zeta, \eta, \theta, \vartheta, \iota, \kappa, \lambda, \mu, \nu\}$, де α – атака вішингу (the vishing attack); β – фішингова атака (phishing attack); γ – клонування профілю (profile cloning); δ – grooming; ε – dumpster diving attacks; ϵ – tailgating; ζ – file masquerade; η – baiting; θ – scareware or pop-up windows; ϑ – water-holing; ι – trojan mail; κ – spear phishing; λ – spam mail; μ – interesting software; ν – hoaxing.

Давайте розглянемо моделі атак соціальної інженерії з використанням атак *вішинг* (vishing), *фішинг* (phishing) і *клонування профілю* (profile cloning), щоб представити покроковий процес атак від початкового вибору цілі до вилучення інформації та механізмів захисту. Цей структурований підхід допомагає проаналізувати та зрозуміти динаміку атак і взаємодію між різними компонентами, а також визначити вразливі місця для розробки стратегій пом'якшення таких атак.

Модель атаки типу “vishing”

Для розробки схеми функціонування моделі атаки вішингу необхідно врахувати ключові компоненти:

розробка приводу: створення правдоподібної причини для дзвінка, наприклад, видавання себе за довірену особу;

підробка ідентифікатора абонента: маніпулювання ідентифікатором абонента для показу надійного номера, підвищення довіри;

психологічне маніпулювання: використання терміновості, страху або авторитету, щоб змусити жертву надати інформацію;

збір інформації: вимагання особистої інформації під виглядом перевірки або обслуговування клієнтів;

використання експлойтів: використання отриманої інформації для шахрайських дій або крадіжки особистих даних;

подальші дії: додаткова маніпуляція або дії прикриття для подальшого використання або приховування шахрайства.

Отже, представимо модель атаки вішингу у вигляді кортежу:

$$M_{\alpha} = \langle A_{\alpha}, V_{\alpha}, P_{\alpha}, C_{\alpha}, T_{\alpha}, I_{\alpha}, D_{\alpha} \rangle, \quad (1)$$

де $A_{\alpha} = \{a_{\alpha 1}, a_{\alpha 2}, \dots, a_{\alpha N_{\alpha A}}\}$ є набором, який представляє осіб, які здійснюють атаки вішингом, $N_{\alpha A}$ – кількість осіб, які здійснюють атаки вішингом;

$V_{\alpha} = \{v_{\alpha 1}, v_{\alpha 2}, \dots, v_{\alpha N_{\alpha V}}\}$ – набір, що включає потенційних і реальних жертв, на яких нападають, $N_{\alpha V}$ – кількість жертв;

$P_{\alpha} = \{p_{\alpha 1}, p_{\alpha 2}, \dots, p_{\alpha N_{\alpha P}}\}$ це набір, який представляє різні приводи або сценарії, які використовують зловмисники для обману жертв; $N_{\alpha P}$ – кількість сценаріїв;

$C_{\alpha} = \{c_{\alpha 1}, c_{\alpha 2}, \dots, c_{\alpha N_{\alpha C}}\}$ – це набір, який включає різні ідентифікатори абонентів, які використовуються під час підробки, щоб зробити виклик легітимним; $N_{\alpha C}$ це кількість ідентифікаторів абонентів;

$T_{\alpha} = \{t_{\alpha 1}, t_{\alpha 2}, \dots, t_{\alpha N_{\alpha T}}\}$ – набір, що містить психологічні прийоми, застосовувані нападниками, $N_{\alpha T}$ –

кількість психологічних прийомів;

$I_\alpha = \{i_{\alpha 1}, i_{\alpha 2}, \dots, i_{\alpha N_{\alpha I}}\}$ – набір, що представляє конфіденційну інформацію, на яку спрямовані зловмисники, $N_{\alpha I}$ – номер конфіденційної інформації;

$D_\alpha = \{d_{\alpha 1}, d_{\alpha 2}, \dots, d_{\alpha N_{\alpha D}}\}$ – це набір, який складається з різноманітних захисних механізмів і засобів протидії, доступних потенційним жертвам, $N_{\alpha D}$ – це кількість захисних механізмів.

Визначимо функцію вибору претексту $f_{\alpha P}$, яка описує процес вибору претексту на основі профілю жертви, як:

$$f_{\alpha P}: V_\alpha \times A_\alpha \rightarrow P_\alpha, f_P(v_i, a_j) = p_k,$$

де p_k є відповідний привід для жертви $v_{\alpha i}$ з боку нападника $a_{\alpha j}$.

Давайте визначимо функцію підробки ідентифікатора абонента $f_{\alpha C}$, яка моделює вибір ідентифікатора абонента для підвищення довіри, як:

$$f_{\alpha C}: P_\alpha \rightarrow C_\alpha, f_{\alpha C}(p_{\alpha k}) = c_{\alpha m},$$

де c_m ідентифікатор абонента, обраний як претекст $p_{\alpha k}$.

Позначимо психологічну маніпулятивну функцію $f_{\alpha T}$, яка фіксує використання психологічних тактик впливу на жертву, як:

$$f_{\alpha T}: P_\alpha \times A_\alpha \rightarrow T_\alpha, f_{\alpha T}(p_{\alpha k}, a_{\alpha j}) = t_{\alpha n},$$

де $t_{\alpha n}$ тактика, використана $p_{\alpha k}$ нападником як привід $a_{\alpha j}$.

Позначимо функцію вилучення інформації $f_{\alpha I}$, яка описує вилучення інформації від жертви, як:

$$f_{\alpha I}: V_\alpha \times P_\alpha \times T_\alpha \rightarrow I_\alpha, f_{\alpha I}(v_{\alpha i}, p_{\alpha k}, t_{\alpha n}) = i_{\alpha o},$$

де $i_{\alpha o}$ інформація, отримана від жертви $v_{\alpha i}$ з використанням приводу $p_{\alpha k}$ та тактики $t_{\alpha n}$.

Позначимо захисно-активаційну функцію $f_{\alpha D}$, яка являє собою активацію захисних механізмів жертвою, як:

$$f_{\alpha D}: V_\alpha \times I_\alpha \rightarrow D_\alpha, f_{\alpha D}(v_{\alpha i}, i_{\alpha o}) = d_{\alpha p},$$

де $d_{\alpha p}$ знаходиться захисний механізм, який активує жертва v_i після отримання інформації $i_{\alpha o}$.

Щоб описати вплив на елементи наборів, визначимо функцію впливу g_α що кількісно визначає успіх атаки та ефективність захисту. Функція впливу враховує ймовірність успіху різних етапів атаки та ймовірність ефективності механізмів захисту.

Позначимо функцію впливу g_α :

$$g_\alpha: V_\alpha \times P_\alpha \times T_\alpha \times I_\alpha \times D_\alpha \rightarrow \mathbb{R}, g_\alpha(v_{\alpha i}, p_{\alpha k}, t_{\alpha n}, i_{\alpha o}, d_{\alpha p}) = S_\alpha,$$

де S_α — дійсне число, що представляє рівень успіху атаки, причому більші значення вказують на більший успіх;

функція g_α може включати такі фактори, як

ймовірність того, що жертва $v_{\alpha i}$ піддається приводу $p_{\alpha k}$ (позначено $P_\alpha(v_{\alpha i}, p_{\alpha k})$).

ефективність психологічної тактики $t_{\alpha n}$ щодо потерпілого $v_{\alpha i}$ (познач $T_\alpha(v_{\alpha i}, t_{\alpha n})$).

ймовірність отримання інформації $i_{\alpha o}$ (позначається $I_\alpha(i_{\alpha o})$);

ефективність захисного механізму $d_{\alpha p}$ (позначається $D_\alpha(d_{\alpha p})$).

Таким чином, для конкретної жертви, $v_{\alpha i}$ на яку нападає $a_{\alpha 1}$:

1. Вибір претексту: $p_k = f_P(v_i, a_j)$.

2. Підробка ідентифікатора абонента: $c_{\alpha m} = f_{\alpha C}(p_{\alpha k})$.

3. Психологічна тактика: $t_{\alpha n} = f_{\alpha T}(p_{\alpha k}, a_{\alpha j})$.

4. Витяг інформації: $i_{\alpha o} = f_{\alpha I}(v_{\alpha i}, p_{\alpha k}, t_{\alpha n})$.

5. Активація захисту: $d_{\alpha p} = f_{\alpha D}(v_{\alpha i}, i_{\alpha o})$.

6. Оцінка впливу: $S_\alpha = g_\alpha(v_{\alpha i}, p_{\alpha k}, t_{\alpha n}, i_{\alpha o}, d_{\alpha p})$.

Модель атаки типу “phishing”

Для розробки схеми функціонування моделі фішингової атаки необхідно врахувати ключові складові: ідентифікація цілей: зловмисники ідентифікують і вибирають потенційних жертв, часто

використовуючи загальнодоступну інформацію, соціальні медіа або витоки даних. Цілі можуть варіюватися від окремих осіб до конкретних співробітників в організації (фішинг);

розробка претексту: створюється правдоподібний і відповідний контексту привід, щоб фішингове повідомлення виглядало законним. це може включати імітацію довіреної організації, такої як банк, державна установа або онлайн-сервіс.

створення повідомлень: фішингове повідомлення створюється з великою увагою до деталей, часто повторюючи дизайн, мову та тон законних повідомлень від імітованої особи;

механізм доставки: повідомлення доставляється до цілі через різні канали, включаючи електронну пошту, sms (smishing) або соціальні мережі;

виконання атаки: коли ціль взаємодіє з фішинговим повідомленням;

Подальші дії: зловмисники можуть використовувати отриману інформацію для здійснення додаткових атак проти тієї самої жертви чи інших (наприклад, використання скомпрометованих облікових записів для надсилання подальших фішингових повідомлень).

Отже, представимо модель атаки вішингу у вигляді кортежу:

$$M_{\beta} = \langle A_{\beta}, V_{\beta}, P_{\beta}, C_{\beta}, T_{\beta}, I_{\beta}, D_{\beta}, \rangle, \quad (2)$$

де $A_{\beta} = \{a_{\beta 1}, a_{\beta 2}, \dots, a_{\beta N_{\beta A}}\}$ є набором, який представляє осіб, які здійснюють фішингові атаки, $N_{\beta A}$ – кількість осіб, які здійснюють фішингові атаки;

$V_{\beta} = \{v_{\beta 1}, v_{\beta 2}, \dots, v_{\beta N_{\beta V}}\}$ – набір, що включає потенційних і реальних жертв, на яких нападають, $N_{\beta V}$ – кількість жертв;

$A_{\beta} = \{a_{\beta 1}, a_{\beta 2}, \dots, a_{\beta N_{\beta A}}\}$ це набір, який представляє окремих осіб або групи, які здійснюють фішингові атаки.

$T_{\beta} = \{t_{\beta 1}, t_{\beta 2}, \dots, t_{\beta N_{\beta T}}\}$ це набір, який містить потенційних жертв, націлених на фішингові повідомлення.

$P_{\beta} = \{p_{\beta 1}, p_{\beta 2}, \dots, p_{\beta N_{\beta P}}\}$ це набір приводів, які використовуються у фішингових повідомленнях для видавання за легітимні джерела.

$D_{\beta} = \{d_{\beta 1}, d_{\beta 2}, \dots, d_{\beta N_{\beta D}}\}$ це набір, який використовує канали для доставки фішингових повідомлень (наприклад, електронної пошти, SMS).

$M_{\beta} = \{m_{\beta 1}, m_{\beta 2}, \dots, m_{\beta N_{\beta M}}\}$ це набір, який містить компоненти фішингового повідомлення, включаючи посилання, вкладення та вміст.

$AC_{\beta} = \{ac_{\beta 1}, ac_{\beta 2}, \dots, ac_{\beta N_{\beta AC}}\}$ це набір, який включає в себе дії, до яких заохочується ціль (наприклад, натискання посилання, завантаження файлу).

$I_{\beta} = \{i_{\beta 1}, i_{\beta 2}, \dots, i_{\beta N_{\beta I}}\}$ це набір, який представляє інформацію або результати, які шукають зловмисники (наприклад, облікові дані, встановлення зловмисного програмного забезпечення).

$DF_{\beta} = \{df_{\beta 1}, df_{\beta 2}, \dots, df_{\beta N_{\beta DF}}\}$ це набір, який містить засоби протидії та захисту, що застосовуються цілями.

Визначимо функцію вибору претексту $f_{\beta P}$ який описує процес вибору приводу на основі профілю жертви як: $f_{\beta P}: T_{\beta} \times A_{\beta} \rightarrow P_{\beta}$.

Ця функція описує, як зловмисник $a_{\beta j}$ вибирає відповідний привід $p_{\beta k}$ для певної цілі $t_{\beta i}$.

Зловмисники обирають привід, який, на їхню думку, буде найбільш переконливим для цілі. Наприклад, якщо відомо, що мета є клієнтом певного банку, привід може включати видавання себе за цей банк $f_{\beta P}(t_{\beta i}, a_{\beta j},) = p_{\beta k}$, наприклад:

$$f_{\beta P}(t_{\beta John}, a_{\beta Hacker1}) = p_{\beta Bank\ Alert},$$

де зловмисник $Hacker1$ вибирає претекст «Банківське сповіщення» для цільового Джона.

Визначимо функцію вибору каналу $f_{\beta D}$ як:

$$f_{\beta D}: P_{\beta} \times A_{\beta} \rightarrow D_{\beta}.$$

Ця функція показує, як зловмисник $a_{\beta j}$ вибирає канал доставки $d_{\beta m}$ для вибраного приводу $p_{\beta k}$.

Зловмисник вибирає найефективніший канал доставки для надсилання фішингового повідомлення, як-от електронна пошта, SMS або соціальні мережі, залежно від вибраного приводу $f_{\beta D}(p_{\beta k}, a_{\beta j},) = d_{\beta m}$, наприклад:

$$f_{\beta D}(p_{\beta Bank\ Alert}, a_{\beta Hacker1}) = d_{\beta Email},$$

де Hacker1 вирішує використати електронну пошту як канал доставки для приводу «Банківське сповіщення».

Давайте визначимо функцію створення повідомлення $f_{\beta M}$ як:

$$f_{\beta M}: P_{\beta} \times D_{\beta} \rightarrow M_{\beta}.$$

Функція $f_{\beta M}$ моделює створення фішингового повідомлення $M_{\beta n}$ на основі претексту $p_{\beta k}$ та каналу доставки $d_{\beta m}$.

Повідомлення створено відповідно до претексту та каналу доставки. Він включає візуальний дизайн, мову та певний зміст, які роблять його легітимним $f_{\beta M}(p_{\beta k}, d_{\beta m}) = m_{\beta n}$, наприклад:

$$f_{\beta M}(p_{\beta Bank\ Alert}, d_{\beta Email}) = m_{\beta Phishing\ Email}$$

де фішингове повідомлення електронної пошти містить фірмовий знак банку та повідомлення про сповіщення системи безпеки.

Позначимо функцію спонування до дії $f_{\beta AC}$ як:

$$f_{\beta AC}: M_{\beta} \times T_{\beta} \rightarrow AC_{\beta}.$$

Функція $f_{\beta AC}$ описує, як створене повідомлення $m_{\beta n}$ спонукає ціль $t_{\beta i}$ виконати певну дію $ac_{\beta o}$.

Фішингове повідомлення спрямоване на те, щоб переконати ціль виконати дію, наприклад натиснути посилання, завантажити вкладений файл або надати конфіденційну інформацію $f_{\beta AC}(m_{\beta n}, t_{\beta i}) = ac_{\beta o}$, наприклад:

$$f_{\beta AC}(m_{\beta Phishing\ Email}, t_{\beta John}) = ac_{\beta Click\ Link},$$

де фішинговий електронний лист переконує Джона натиснути посилання.

Позначимо функцію вилучення інформації $f_{\beta I}$ як:

$$f_{\beta I}: AC_{\beta} \times T_{\beta} \rightarrow I_{\beta}.$$

Ця функція представляє вилучення конфіденційної інформації $i_{\beta p}$ з цілі $t_{\beta i}$ на основі дії $ac_{\beta o}$.

Після того, як мета виконує індуковану дію, зловмисник захоплює потрібну інформацію, таку як облікові дані для входу або особисті дані $f_{\beta I}(ac_{\beta o}, t_{\beta i}) = i_{\beta p}$, наприклад:

$$f_{\beta I}(ac_{\beta Click\ Link}, t_{\beta John}) = i_{\beta Login\ Credentials},$$

де, натиснувши посилання, Джона обманом змушують надати свої облікові дані для входу.

Позначимо функцію активації захисту $f_{\beta DF}$ як:

$$f_{\beta DF}: T_{\beta} \times M_{\beta} \rightarrow DF_{\beta}.$$

Ця функція моделює активацію захисних механізмів $df_{\beta q}$ цілі $t_{\beta i}$ у відповідь на фішингове повідомлення $m_{\beta n}$. Цілі можуть використовувати різні засоби захисту, щоб розпізнати та пом'якшити спробу фішингу, як-от фільтри електронної пошти, навчання користувачів або багатофакторну автентифікацію $f_{\beta DF}(t_{\beta i}, m_{\beta n}) = df_{\beta q}$, наприклад:

$$f_{\beta DF}(t_{\beta John}, m_{\beta Phishing\ Email}) = df_{\beta Two-FactorAuth},$$

де Джон використовує двофакторну автентифікацію як механізм захисту від фішингових електронних листів.

Таким чином, для конкретної жертви, $v_{\beta i}$ на яку нападає $a_{\beta 1}$:

Вибір претексту для мети зловмисниками: $f_{\beta P}(t_{\beta i}, a_{\beta j}) = p_{\beta k}$.

Вибір зловмисниками каналу доставки для приводу: $f_{\beta D}(p_{\beta k}, a_{\beta j}) = d_{\beta m}$.

Створення фішингового повідомлення зловмисниками: $f_{\beta M}(p_{\beta k}, d_{\beta m}) = m_{\beta n}$.

Спонування мети повідомлення до дії: $f_{\beta A^c}(m_{\beta n}, t_{\beta i}) = ac_{\beta o}$.

Результатом отримання дії є вилучення інформації: $f_{\beta I}(ac_{\beta o}, t_{\beta i}) = i_{\beta p}$.

Активізація захисних механізмів мішенями: $f_{\beta D^F}(t_{\beta i}, m_{\beta n}) = df_{\beta q}$.

Позначимо функцію впливу g_{β} так:

$$g_{\beta} : T_{\beta} \times P_{\beta} \times D_{\beta} \times M_{\beta} \times AC_{\beta} \times I_{\beta} \times DF_{\beta} \rightarrow R_{\beta}, \\ g_{\beta}(t_{\beta i}, p_{\beta k}, d_{\beta m}, m_{\beta n}, ac_{\beta o}, i_{\beta p}, df_{\beta q}) = S_{\beta},$$

де: S_{β} - рівень успішності фішингової атаки.

Фактори, що впливають, S_{β} можуть включати ефективність приводу, достовірність каналу доставки, переконливу силу компонентів повідомлення, ймовірність того, що об'єкт виконає індуковану дію, і ефективність захисних механізмів об'єкта.

Модель атаки типу "profile cloning"

Для розробки схеми функціонування моделі атаки клонування профілю необхідно врахувати ключові компоненти:

- ✓ цільовий користувач: особа, чий профіль буде клоновано;
- ✓ збір інформації: зловмисники збирають дані про ціль;
- ✓ створення профілю: створення підробленого профілю, що імітує цільовий;
- ✓ встановлення довіри: зміцнення довіри з контактною особою;
- ✓ почати контакт: звернення до контактів жертви;
- ✓ експлуатація: отримання конфіденційної інформації від жертви;
- ✓ подальші дії: продовження взаємодії або поворот до нових цілей;
- ✓ і включати контрзаходи;
- ✓ налаштування конфіденційності: обмеження видимості публічної інформації.
- ✓ перевірка: перевірка неочікуваних запитів друзів.
- ✓ обізнаність та навчання: навчати про ризики та стратегії захисту.
- ✓ повідомляти та блокувати: повідомляти про підозрілі профілі та блокувати їх.
- ✓ двофакторна аутентифікація: безпечні облікові записи з додатковою перевіркою.

Отже, представимо модель атаки клонування профілю у вигляді кортежу:

$$M_{\gamma} = \langle A_{\gamma}, T_{\gamma}, I_{\gamma}, C_{\gamma}, V_{\gamma}, E_{\gamma} \rangle, (1)$$

де $A_{\gamma} = \{a_{\gamma 1}, a_{\gamma 2}, \dots, a_{\gamma N_{\gamma A}}\}$ — набір, який представляє осіб або групи, які здійснюють атаки клонування профілю.

$T_{\gamma} = \{t_{\gamma 1}, t_{\gamma 2}, \dots, t_{\gamma N_{\gamma T}}\}$ це набір, який представляє користувачів, профілі яких клоновано;

$I_{\gamma} = \{i_{\gamma 1}, i_{\gamma 2}, \dots, i_{\gamma N_{\gamma I}}\}$ набір, що включає загальнодоступну інформацію про ціль, що використовується для клонування;

$C_{\gamma} = \{c_{\gamma 1}, c_{\gamma 2}, \dots, c_{\gamma N_{\gamma C}}\}$ — набір, що містить підроблені профілі, створені зловмисниками;

$V_{\gamma} = \{v_{\gamma 1}, v_{\gamma 2}, \dots, v_{\gamma N_{\gamma V}}\}$ — набір, що містить контакти об'єкта, які є потенційними жертвами подальшої експлуатації;

$E_{\gamma} = \{e_{\gamma 1}, e_{\gamma 2}, \dots, e_{\gamma N_{\gamma E}}\}$ це набір, який представляє методи соціальної інженерії, що використовуються для отримання інформації з контактів цільової групи.

Позначимо функцію збору інформації $f_{\gamma I}$, яка описує, як зловмисники збирають інформацію про ціль, як:

$$f_{\gamma I} : T_{\gamma} \times A_{\gamma} \rightarrow I_{\gamma}, f_{\gamma I}(t_{\gamma i}, a_{\gamma j}) = i_{\gamma k}.$$

Позначимо функцію створення профілю $f_{\gamma C}$, яка моделює створення клонованого профілю з використанням зібраної інформації:

$$f_{\gamma C} : I_{\gamma} \times A_{\gamma} \rightarrow C_{\gamma}, f_{\gamma C}(t_{\gamma i}, a_{\gamma j}) = c_{\gamma m}.$$

Давайте позначимо функцію встановлення довіри $f_{\gamma EC}$, яка представляє процес встановлення довіри до клонованого профілю шляхом взаємодії з контактами цілі як:

$$f_{\gamma EC} : C_{\gamma} \times V_{\gamma} \rightarrow C_{\gamma}, f_{\gamma EC}(c_{\gamma m}, v_{\gamma n}) = c'_{\gamma m}$$

Давайте визначимо функцію ініціації контакту $f_{\gamma C_I}$, яка описує, як зломисники використовують клонований профіль, щоб ініціювати контакт із з'єднаннями цілі, як:

$$f_{\gamma C_I}: C_\gamma \times V_\gamma \rightarrow V_\gamma, f_{\gamma C_I}(c'_{\gamma m}, v_{\gamma n}) = v'_{\gamma n}.$$

Позначимо функцію експлуатації $f_{\gamma E}$, яка моделює використання методів соціальної інженерії для експлуатації контактів об'єкта, як:

$$f_{\gamma E}: V_\gamma \times C_\gamma \times E_\gamma \rightarrow I_\gamma, f_{\gamma E}(v'_{\gamma n}, c'_{\gamma m}, e_{\gamma p}) = i_{\gamma q}$$

Позначимо активаційну функцію захисту $f_{\gamma D^F}$, що представляє активацію механізмів захисту мішенню або її контактами, як:

$$f_{\gamma D^F}: T_\gamma \times C_\gamma \times V_\gamma \rightarrow D_\gamma, f_{\gamma D^F}(t_{\gamma i}, c_{\gamma m}, v_{\gamma n}) = d_{\gamma r}$$

Щоб описати загальний вплив атаки клонування профілю, ми визначаємо функцію впливу, яка кількісно визначає успіх атаки та ефективність механізмів захисту.

Позначимо функцію впливу g_γ як:

$$g_\gamma: T_\gamma \times I_\gamma \times C_\gamma \times V_\gamma \times E_\gamma \times D_\gamma \rightarrow R_\gamma, \\ g_\gamma: (t_{\gamma i}, i_{\gamma k}, c_{\gamma m}, v_{\gamma n}, e_{\gamma p}, d_{\gamma r}) = S_\gamma,$$

де S_γ представляє рівень успіху атаки клонування профілю.

Фактори, що впливають, S_γ можуть включати: кількість і якість інформації, $i_{\gamma k}$ зібраної про ціль; достовірність клонованого профілю $c_{\gamma m}$; ефективність методів соціальної інженерії $e_{\gamma p}$; реагування та захисні механізми, $d_{\gamma r}$ активовані ціллю або її контактами.

Таким чином, для конкретної жертви, $v_{\gamma i}$ на яку нападає $a_{\gamma 1}$:

$f_{\gamma I}(t_{\gamma i}, a_{\gamma j}) = i_{\gamma k}$: зломисники збирають інформацію про ціль.

$f_{\gamma C}(t_{\gamma k}, a_{\gamma j}) = c_{\gamma m}$: зломисники створюють клонований профіль.

$f_{\gamma E^C}(c_{\gamma m}, v_{\gamma n}) = c'_{\gamma m}$: зломисники встановлюють достовірність клонованого профілю.

$f_{\gamma C^I}(c'_{\gamma m}, v_{\gamma n}) = v'_{\gamma n}$: зломисники ініціюють контакт із з'єднаннями цілі.

$f_{\gamma E}(v'_{\gamma n}, c'_{\gamma m}, e_{\gamma p}) = i_{\gamma q}$: зломисники використовують зв'язки цілі за допомогою методів соціальної інженерії.

$f_{\gamma D^F}(t_{\gamma i}, c_{\gamma m}, v_{\gamma n}) = d_{\gamma r}$: цілі або їхні зв'язки активують захисні механізми.

ВИСНОВКИ З ДАНОГО ДОСЛІДЖЕННЯ І ПЕРСПЕКТИВИ ПОДАЛЬШИХ РОЗВІДОК У ДАНОМУ НАПРЯМІ

У статті представлені моделі атак соціальної інженерії. В якості прикладів були використані атаки на вішинг, фішинг і клонування профілю.

Запропоновані моделі представляють різноманітні форми атак соціальної інженерії з урахуванням різних векторів атак, характеристик цілей і методів психологічної маніпуляції, які використовують зломисники.

Запропоновані моделі охоплюють весь життєвий цикл атаки, від розвідки до аналізу після атаки, забезпечуючи всебічне розуміння загроз.

Зрештою, ці моделі відіграють важливу роль у прогнозуванні потенційних атак, покращенні навчання та розробці надійного захисту від загроз соціальної інженерії, а також є основою для нових методів виявлення атак соціальної інженерії.

Література

- [1] P. Zambrano, J. Torres, L. Tello-Oquendo, Á. Yáñez, and L. Velásquez, "On the modeling of cyber-attacks associated with social engineering: A parental control prototype," Journal of Information Security and Applications, vol. 75, p. 103501, 2023/06/01/ 2023, doi: <https://doi.org/10.1016/j.jisa.2023.103501>.
- [2] W. Syafitri, Z. Shukur, U. A. Mokhtar, R. Sulaiman, and M. A. Ibrahim, "Social Engineering Attacks Prevention: A Systematic Literature Review," IEEE Access, vol. 10, pp. 39325-39343, 2022, doi: <https://doi.org/10.1109/ACCESS.2022.3162594>.
- [3] R. F. Abu Hweidi and D. Eleyan, "Social Engineering Attack Concepts, Frameworks, and Awareness: A Systematic Literature Review," International Journal of Computing and Digital Systems, 2023, doi: <http://dx.doi.org/10.12785/ijcds/130155>.

- [4] A. Raval, S. Chakrabarty, H. Jasoliya and D. Swain, "Understanding People's awareness towards social engineering with survey," 2022 IEEE 2nd International Symposium on Sustainable Energy, Signal Processing and Cyber Security (iSSSC), Gunupur, Odisha, India, 2022, pp. 1-5, doi: <https://doi.org/10.1109/iSSSC56467.2022.10051531>.
- [5] A. Solomon et al., "Contextual security awareness: A context-based approach for assessing the security awareness of users," Knowledge-Based Systems, vol. 246, p. 108709, 2022/06/21/ 2022, doi: <https://doi.org/10.1016/j.knosys.2022.108709>.
- [6] W. Fuertes, D. Ar'evalo, J. D. Castro, M. Ron, C. A. Estrada, R. Andrade, F. F. Pe'na, and E. Benavides, "Impact of social engineering attacks: A literature review," Smart Innovation, Systems and Technologies, vol. 255, pp. 25–35, 2022. [Online]. Available: <https://link.springer.com/chapter/10.1007/978-981-16-4884-73>.
- [7] L. Karadsheh, H. Alryalat, J. Alqatawna, S. F. Alhawari, and M. A. A. Jarrah, "The impact of social engineer attack phases on improved security countermeasures: Social engineer involvement as mediating variable," International Journal of Digital Crime and Forensics, vol. 14, pp. 1–26, 1 2022.
- [8] T. B. G. Herath, P. Khanna, and M. Ahmed, "Cybersecurity practices for social media users: A systematic literature review," Journal of Cybersecurity and Privacy 2022, Vol. 2, Pages 1-18, vol. 2, pp. 1–18, 1 2022. [Online]. Available: <https://www.mdpi.com/2624-800X/2/1/1/htmlhttps://www.mdpi.com/2624-800X/2/1/1>.
- [9] L. Karadsheh, H. Alryalat, J. Alqatawna, S. F. Alhawari, and M. A. A. Jarrah, "The impact of social engineer attack phases on improved security countermeasures: Social engineer involvement as mediating variable," International Journal of Digital Crime and Forensics, vol. 14, pp. 1–26, 1 2022.
- [10] P. Harr, "Defending Against AI-Based Phishing Attacks," 2023. <https://www.forbes.com/sites/forbestechcouncil/2023/08/04/defending-against-ai-based-phishing-attacks/?sh=1d4d61b83da6> (accessed Sep. 24, 2023).
- [11] E. Ogala , R. Ogala, M. U. Agbane , E. R. Adeiza, B. Tenuche, S. Sunday, Detecting Telecoms Fraud in a Cloud-Base Environment by Analyzing the Content of a Phone Conversation, Asian Journal of Research in Computer Science. 4(3) (2022) 115-131.
- [12] M. Nandakumar, R. Nachiappan, Sunil, A.K., Neves, J.C., Proença, H.P., M. Sathiyarayanan, ScamBlk: A Voice Recognition-Based Natural Language Processing Approach for the Detection of Telecommunication Fraud, in: Bashir, A.K., Fortino, G., Khanna, A., Gupta, D. (eds) Proceedings of International Conference on Computing and Telecommunication Networks, Lecture Notes in Networks and Systems, volume 394, Springer, Singapore, 2022, doi.org/10.1007/978-981-19-0604-6_47.
- [13] S. Malhotra, G. Arora, R. Bathla, Detection and Analysis of Fraud Phone Calls using Artificial Intelligence, in: 2023 International Conference on Recent Advances in Electrical, Electronics & Digital Healthcare Technologies (REEDCON), 01-03 May, 2023, IEEE Xplore ISBN:978-1-6654-9382-6, doi:10.1109/REEDCON57544.2023.10150631.
- [14] C. J. Shalke, R. Achary, Social Engineering Attack and Scam Detection using Advanced Natural Language Processing Algorithm, in: 2022 6th International Conference on Trends in Electronics and Informatics (ICOEI), Tirunelveli, India, 2022, pp. 1749-1754, doi: 10.1109/ICOEI53556.2022.9776697.
- [15] N. Hasan, N. Ahmed, S. M. Ali, Improving sporadic demand forecasting using a modified k-nearest neighbor framework. Engineering Applications of Artificial Intelligence. 2024. Vol. 129, 107633.
- [16] D. Chicco, G. Jurman, The Matthews correlation coefficient (MCC) should replace the ROC AUC as the standard metric for assessing binary classification. BioData Mining. 2023. Vol. 16 (1). Pp.1-23.
- [17] M. F. Ansari, P. K. Sharma, B. Dash, Prevention of Phishing Attacks Using AI-Based Cybersecurity Awareness Training. International Journal of Smart Sensor and Adhoc Network., 2022, 61–72. <https://doi.org/10.47893/ijssan.2022.1221>.
- [18] M. A. Siddiqi, W. Pak, & M. A. Siddiqi, A Study on the Psychology of Social Engineering-Based Cyberattacks and Existing Countermeasures. In Applied Sciences, 13 (Switzerland), 2022, Vol. 12, Issue 12. MDPI. <https://doi.org/10.3390/app12126042>.
- [19] M. A. Siddiqi, W. Pak, M. A. Siddiqi, A Study on the Psychology of Social Engineering-Based Cyberattacks and Existing Countermeasures. In Applied Sciences 13 (Switzerland) (2022) (Vol. 12, Issue 12). MDPI. <https://doi.org/10.3390/app12126042>
- [20] R. Kaur, "Artificial intelligence for cybersecurity : Literature review and future research directions," vol. 97, no. January, 2023, doi: 10.1016/j.inffus.2023.101804.
- [21] G. Lawton, "What is generative AI? Everything you need to know," 2023. <https://www.techtarget.com/searchenterpriseai/definition/generative-AI> (accessed Sep. 29, 2023).
- [22] B. Violino, "Phishing attacks are increasing and getting more sophisticated. Here's how to avoid them," 2023. <https://www.cnn.com/2023/01/07/phishing-attacks-are-increasing-and-getting-more-sophisticated.html> (accessed Sep. 29, 2023).
- [23] P. Mwiinga, "Investigating the Far-Reaching Consequences of Cybercrime A Case Study on the Impact in Lusaka," no. July, 2023.
- [24] A. Haleem, M. Javaid, and R. Pratap, "BenchCouncil Transactions on Benchmarks , Standards and Evaluations An era of ChatGPT as a significant futuristic support tool : A study on features , abilities , and challenges," BenchCouncil Trans. Benchmarks, Stand. Eval., vol. 2, no. 4, p. 100089, 2023, doi: 10.1016/j.bench.2023.100089.
- [25] S. Sjouwerman, "How AI Is Changing Social Engineering Forever," 2023. <https://www.forbes.com/sites/forbestechcouncil/2023/05/26/how-ai-is-changing-social-engineering-forever/?sh=123037f5321b> (accessed Sep. 26, 2023).
- [26] A. Rudra, "Cybersecurity Risks of Generative AI," 2023. <https://securityboulevard.com/2023/07/cybersecurity-risks-of-generative-ai/> (accessed Sep. 26, 2023).
- [27] D. Gupta, "The Road Ahead: Adapting to the Generative AI Cybersecurity Landscape," 2023. <https://securityboulevard.com/2023/08/the-road-ahead-adapting-to-the-generative-ai-cybersecurity-landscape/> (accessed Sep. 24, 2023).
- [28] D. Riley, "Cybercriminals are using custom 'WormGPT' for business email compromise attacks," 2023. <https://siliconangle.com/2023/07/13/slashnext-warns-cybercriminals-using-custom-wormgpt-business-email-compromise-attacks/> (accessed Sep. 26, 2023).
- [29] S. Rushin, "The Dark Side of Generative AI: Unveiling the Cybersecurity Risk," 2023. <https://www.digit.fyi/comment-the-dark-side-of-generative-ai-unveiling-the-cybersecurity-risk/> (accessed Sep. 24, 2023).
- [30] Darktrace, "Major Upgrade to Darktrace/EmailTM Product Defends Organizations Against Evolving Cyber Threat Landscape, Including Generative AI Business Email Compromises and Novel Social Engineering Attacks," 2023. <https://darktrace.com/news/darktrace-email-defends-organizations-against-evolving-cyber-threat-landscape> (accessed Sep. 26, 2023).
- [31] S. Ortiz, "What is ChatGPT and why does it matter? Here's what you need to know," 2023. <https://www.zdnet.com/article/what-is-chatgpt-and-why-does-it-matter-heres-everything-you-need-to-know/> (accessed Sep. 29, 2023).
- [32] M. Vizard, "SlashNext Report Shows How Cybercriminals Use Generative AI," 2023. <https://securityboulevard.com/2023/07/slashnext-report-shows-how-cybercriminals-use-generative-ai/> (accessed Sep. 26, 2023).
- [33] E. K. Sing, "With generative AI, businesses need to rewrite the phishing rulebook," 2023. <https://identityweek.net/with-generative-ai-businesses-need-to-rewrite-the-phishing-rulebook/> (accessed Sep. 22, 2023).

- [34] L. Columbus, "How FraudGPT presages the future of weaponized AI," 2023. <https://venturebeat.com/security/how-fraudgpt-presages-the-future-of-weaponized-ai/> (accessed Sep. 24, 2023).
- [35] R. Bathgate, "Mandiant says generative AI will empower new breed of information operations, social engineering," 2023. <https://www.itpro.com/technology/artificial-intelligence/mandiant-says-generative-ai-will-empower-new-breed-of-information-operations-social-engineering> (accessed Sep. 26, 2023).
- [36] S. Das, "Back 'Voice scams hit 47% web users,'" 2023. <https://www.livemint.com/companies/start-ups/indias-internet-users-vulnerable-to-ai-powered-voice-scams-mcafee-reports-47-of-indian-users-encounter-or-know-victims-11683032655430.html> (accessed Sep. 26, 2023).
- [37] M. Al-Khateeb, M.R. Al-Mousa, A.S. Al-Sherideh, D. Almajali, M. Asassfeha, H. Khafajeh, "Awareness Model for Minimizing the Effects of Social Engineering Attacks in Web Applications," in *International Journal of Data and Network Science*, 2023, vol. 7, no. 2, pp. 791-800, ISSN 2561-8148, DOI: 10.5267/j.ijdns.2023.1.010.
- [38] E. Benavides-Astudillo, W. Fuertes, S. Sanchez-Gordon, D. Nuñez-Agurto, G. Rodríguez-Galán, "A Phishing-Attack-Detection Model Using Natural Language Processing and Deep Learning," in *Applied Sciences*, 2023, vol. 13, <https://doi.org/10.3390/app13095275>.
- [39] J.Soni, S.Sirigineedi, K.S.Vutukuru, S.S.C. Sirigineedi, N. Prabakar, H. Upadhyay, "Learning-Based Model for Phishing Attack Detection", in *Springer, Cham.*, 2023. vol. 240. https://doi.org/10.1007/978-3-031-28581-3_11
- [40] M. Aijaz, M. Nazir," Modelling and analysis of social engineering threats using the attack tree and the Markov model", in *International Journal of Information Technology*, 2023, Volume 16, pp. 1231–1238, <https://doi.org/10.1007/s41870-023-01540-z>.
- [41] L.Tong , S.Chuanyong , P.Qinyu , "Defending against social engineering attacks: A security pattern-based analysis framework" in *IET Information Security*, 2023, <https://doi.org/10.1049/ise2.12125>.
- [42] O.Kuzmenko, H.Yarovenko, L. Skrynka, "Analysis of Mathematical Models for Countering Banking Cyberfraud" in *Bulletin of Sumy State University. Series: Economics*, 2022, 111-120. <https://doi.org/10.21272/1817-9215.2022.2-13>.
- [43] V. Lakhno, V. Malyukov, I. Malyukova, O. Atkeldi, O. Kryvoruchko, A. Desiatko, K. Stepashkina, "Model for Analyzing Strategies in Dynamic Interaction of Phishing Attack Participants", in *Cybersecurity: Education, Science, Technology*, 2023, Volume 4, <https://doi.org/10.28925/2663-4023.2023.20.124141>.
- [44] S. Asiri, Y. Xiao, S. Alzahrani, S. Li, T. Li "A Survey of Intelligent Detection Designs of HTML URL Phishing Attacks", in: *IEEE Access* , 2023, Volume: 11, page(s) 6421 – 6443, ISSN: 2169-3536, DOI: 10.1109/ACCESS.2023.3237798.
- [45] N. Do, A. Selamat, O. Krejcar, E. Herrera-Viedma, H. Fujita, "Deep Learning for Phishing Detection: Taxonomy, Current Challenges and Future Directions", in *IEEE Access*, 2022, (Volume: 10) page(s) 36429 – 36463, ISSN: 2169-3536, DOI: 10.1109/ACCESS.2022.3151903.
- [46] G. Cimino, V. Deufemia, "Towards Enhanced Human Mitigation of Vishing Attacks: Leveraging Large Language Models for Real-Time User Guidance" in *eur-ws.org Workshop co-located with AVI 2024, University of Salerno, via Giovanni Paolo II, Fisciano (SA), Arenzano, Genoa, Italy*, 2024 pp 6.
- [47] P. Zambrano, J. Torres, L. Tello-Oquendo, Á. Yáñez, L. Velásquez, "On the modeling of cyber-attacks associated with social engineering: A parental control prototype", in *Journal of Information Security and Applications*, June 2023, Volume 75, <https://doi.org/10.1016/j.jisa.2023.103501>.
- [48] S. Sadhasivam, P. Valarmathie, K. Dinakaran, "Malicious Activities Prediction Over Online Social Networking Using Ensemble Model" in *Intelligent Automation & Soft Computing* 2023, 36(1), 461-479, <https://doi.org/10.32604/iasc.2023.028650>.
- [49] F.Bianchi, J. Zou, "Large Language Models are Vulnerable to Bait-and-Switch Attacks for Generating Harmful Content" in *Cite as arXiv:2402.13926*, <https://doi.org/10.48550/arXiv.2402.13926>.
- [50] Ö. Aslan, S. S. Aktuğ, M. Ozkan-Okay, A. A. Yilmaz, E. Akin, "A Comprehensive Review of Cyber Security Vulnerabilities, Threats, Attacks, and Solutions", in *Electronics*, 2023, 12(6), 1333; <https://doi.org/10.3390/electronics12061333>.
- [51] A. Raj, V. Narayan, V. Muskan, A. Sani, P. Sharma, S. S. Sarma, "Modern ransomware: Evolution, methodology, attack model, prevention and mitigation using multi-tiered approach", in *Security & Privacy*, June, 2024, <https://doi.org/10.1002/spy2.436>.
- [52] C. Lehman, "Generative AI in Cybersecurity: The Battlefield, The Threat, & Now The Defense," 2023. <https://www.unite.ai/generative-ai-in-cybersecurity-the-battlefield-the-threat-now-the-defense/> (accessed Sep. 24, 2023).
- [53] H. Peng, H. Qiu, H. Ma, S. Wang, A. Fu, S. F. Al-Sarawi, "On Model Outsourcing Adaptive Attacks to Deep Learning Backdoor Defenses" in *IEEE Transactions on Information Forensics and Security*, 2024, Vol. 19, pp. 2356 – 2369, DOI: 10.1109/TIFS.2024.3349869.
- [54] Y. Zhang, A. He, J. Li, A. Rezine, Z. Peng, E. Larsson, T. Yang, J. Jiang, H. Li, "On Modeling and Detecting Trojans in Instruction Sets" in: *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems*, 2024, DOI: 10.1109/TCAD.2024.3389558.
- [55] M. Bethany, A. Galiopoulos, E. Bethany, M. B. Karkevandi, N. Vishwamitra, P. Najafirad, "Large Language Model Lateral Spear Phishing: A Comparative Study in Large-Scale Organizational Settings", in *cite as arXiv:2401.09727*.
- [56] F. Heiding, B. Schneier, A. Vishwanath, J. Bernstein, P. S. Park "Devising and Detecting Phishing Emails Using Large Language Models" in: *IEEE Access*, 2024, Vol. 12, pp. 42131 – 42146, ISSN: 2169-3536, DOI: 10.1109/ACCESS.2024.3375882.
- [57] Z. Zhang, E. Damiani, H. Hamadi, C. Yeun, F. Taher, "A Late Multi-modal Fusion Model for Detecting Hybrid Spam E-mail", in *International Journal of Computer Theory and Engineering*, 2023, Vol. 15, No. 2, DOI:10.48550/arXiv.2210.14616.
- [58] G. Mohan, R. Kumar, S. N., P. Ayswariya, P. Yagnitha, "Comparative Analysis of Neural Network Models for Spam E-mail Detection", in: *2024 Fourth International Conference on Advances in Electrical, Computing, Communication and Sustainable Technologies (ICAECT)*, 2024, DOI: 10.1109/ICAECT60202.2024.10469042.
- [59] B. Aruna Kumari, C. Nagaraju, "A Generalized Two-Level Ensemble Method for Spam Mail Detection" in *J. Electrical Systems*, 2024, pp. 1570-1579, DOI:10.52783/jes.1462.
- [60] N. Harum, Z. Z. Abidin, W. M. Shah, A. Hassan, (2018). Implementation of smart monitoring system with fall detector for elderly using iot technology. *International Journal of Computing*, 17(4), 243-249. <https://doi.org/10.47839/ijc.17.4.1146>
- [61] Savenko O. Botnet detection technique for corporate area network / O. Savenko, S. Lysenko, A. Kryshchuk, Y. Klots / *Proceedings of the 7-th IEEE International Conference on Intelligent Data Acquisition and Advanced Computing Systems: Technology and Applications*, Berlin (Germany), September 12–14, 2013. – Berlin, 2013. – Pp. 363–368. ISBN 978-1-4799-1426-5
- [62] Pomorova O. A Technique for the Botnet Detection Based on DNS-Traffic Analysis [Text] / O. Pomorova, O. Savenko, S. Lysenko, A. Kryshchuk, K. Bobrovnikova // *Communications in Computer and Information Science*. – 2015. – Vol. 522. – Pp.127-138, ISSN: 1865-0929
- [63] Sergii Lysenko, Kira Bobrovnikova and Oleg Savenko. A Botnet Detection Approach Based on The Clonal Selection Algorithm // *Proceedings of 2018 IEEE 9th International Conference on Dependable Systems, Services and Technologies (DeSSerT-2018, Kyiv, Ukraine, May 24-27, 2018)* – Pp. 424-428
- [64] S. Lysenko, O. Savenko, K. Bobrovnikova, A. Kryshchuk, B. Savenko, "Information technology for botnets detection based on their behaviour in the corporate area network," in *Communications in Computer and Information Science*, 2017, Vol. 718, pp. 166–181.

-
- [65] O. Savenko, S. Lysenko, A. Kryschuk, "Multi-agent based approach of botnet detection in computer systems," in Communications in Computer and Information Science, 2012, Vol. 291. pp.171-180.
- [66] O. Pomorova, O. Savenko, S. Lysenko, A. Kryshchuk, K. Bobrovnikova, "Anti-evasion Technique for the Botnets Detection Based on the Passive DNS Monitoring and Active DNS Probing," in Communications in Computer and Information Science, 2016, Vol. 608, pp.83-95.
- [67] K. Bobrovnikova, S. Lysenko, B. Savenko, P. Gaj, O. Savenko, "Technique for IoT malware detection based on control flow graph analysis," in Radioelectronic and Computer Systems. 2022, pp. 141–153.
- [68] Bart Lenaerts-Bergmans. CrowdStrike. 10 Types of Social Engineering Attacks and how to prevent them. URL: <https://www.crowdstrike.com/cybersecurity-101/types-of-social-engineering-attacks/>
- [69] O. Savenko, A. Sachenko, S. Lysenko, G. Markowsky, N. Vasylykiv, "Botnet detection approach based on the distributed systems", in International Journal of Computing 2020, vol. 19(2), pp. 190-198. <https://doi.org/10.47839/ijc.19.2.1761A>.
- [70] M.Karpinski, A. Naglik, G.Shangytbayeva, I. Romanets, "Using graphic network simulator 3 for DDoS attacks simulation", International Journal of Computing, 2017, vol.16(4), pp. 219-2