



[DOI 10.28925/2663-4023.2026.34.1222](https://doi.org/10.28925/2663-4023.2026.34.1222)

УДК 004.056

Кривокульська Ольга Олексіївна

старший викладач кафедри кібербезпеки

Національний університет «Київський авіаційний інститут», Київ, Україна

ORCID:0009-0003-8518-6915

olha.kryvokulska@npp.kai.edu.ua

Яковенко Олеся Леонідівна

старший викладач кафедри кібербезпеки

Національний університет «Київський авіаційний інститут», Київ, Україна

ORCID:0000-0003-2998-9767

olesia.yakovenko@npp.kai.edu.ua

Марченко Ярослав Володимирович

асистент кафедри кібербезпеки

Національний університет «Київський авіаційний інститут», Київ, Україна

ORCID:0009-0006-9457-1235

yaroslav.marchenko@npp.kai.edu.ua

Якимчук Євгеній Анатолійович

асистент кафедри кібербезпеки

Національний університет «Київський авіаційний інститут», Київ, Україна

ORCID:0009-0009-9736-5260

yevhenii.yakymchuk@npp.kai.edu.ua

ФІШИНГ НОВОГО ПОКОЛІННЯ: ЕВОЛЮЦІЯ СОЦІАЛЬНОЇ ІНЖЕНЕРІЇ В УМОВАХ ЦИФРОВОЇ ТРАНСФОРМАЦІЇ

Анотація. У статті здійснено аналіз сучасних тенденцій розвитку фішингу, включаючи використання технологій штучного інтелекту для автоматизації створення фішингового контенту, підвищення рівня персоналізації та оптимізації сценаріїв атак. Розглянуто вплив людського фактора як ключового елемента вразливості, що визначає ефективність соціальної інженерії. Проведено кількісний аналіз динаміки фішингових атак за останні роки, який демонструє експоненційне зростання кількості інцидентів та свідчить про активне використання автоматизованих інструментів кіберзлочинцями. Окремо досліджено розподіл каналів реалізації фішингових атак, зокрема електронної пошти, SMS, голосових комунікацій та соціальних мереж, що дозволило визначити їхню відносну ефективність та рівень ризику. Запропоновано багатофакторну математичну модель оцінки ймовірності успішності фішингової атаки, яка враховує рівень захищеності інформаційної системи, ефективність застосованих методів соціальної інженерії, характеристики каналу поширення, а також поведінкові особливості користувачів. На відміну від існуючих підходів, модель розширено шляхом введення часової функції адаптації, що дозволяє враховувати зміну параметрів атаки в динаміці. Додатково використано елементи теорії ймовірностей для моделювання інтенсивності атак та підхід інформаційної ентропії для оцінки невизначеності у виборі каналів фішингу. Отримані результати підтверджують, що сучасний фішинг є складною багатовимірною загрозою, яка поєднує технічні, поведінкові та інформаційні аспекти. Запропонована модель дозволяє формалізувати процес оцінки ризику, здійснювати прогнозування ефективності атак та може бути використана для розробки адаптивних систем кіберзахисту. Практична значущість дослідження полягає у можливості застосування отриманих результатів у системах управління інформаційною безпекою, а також у підвищенні рівня обізнаності користувачів щодо сучасних методів соціальної інженерії.

Ключові слова: фішинг, соціальна інженерія, штучний інтелект, кібербезпека, математичне моделювання



ВСТУП

Фішинг трансформувався з простої техніки масових атак у складну адаптивну систему впливу, що поєднує технології штучного інтелекту, аналіз поведінки користувачів і багатоканальні комунікації. Сучасні атаки вже не обмежуються масовою розсилкою однакових повідомлень – вони персоналізуються під конкретну жертву з використанням даних із відкритих джерел та машинного навчання [7, 13].

Основна проблема полягає в дисбалансі між технічним захистом і людським фактором. Навіть найсучасніші системи фільтрації та двофакторна аутентифікація виявляються неефективними, коли атака експлуатує психологічні особливості користувача [11, 14].

Постановка проблеми. Ключова суперечність сучасної кібербезпеки полягає в зростаючому дисбалансі між рівнем технічного захисту та вразливістю «людського фактору» [9, 15]. Попри впровадження інтелектуальних систем фільтрації трафіку та протоколів MFA, зловмисники успішно експлуатують когнітивні упередження. Згідно зі звітом Verizon DBIR 2024 [1], людська помилка залишається першопричиною близько 74 % усіх інцидентів.

Метою статті є реалізація інтегрованої математичної моделі оцінювання ймовірності успішності фішингових атак, яка враховує взаємодію технічних параметрів захисту інформаційних систем, характеристик каналів комунікації та поведінкових особливостей користувачів в умовах застосування технологій штучного інтелекту.

Для досягнення поставленої мети у роботі використано сукупність загальнонаукових і спеціальних методів дослідження, зокрема: статистичний аналіз – для обробки емпіричних даних щодо динаміки фішингових атак у 2020-2024 рр. на основі звітів APWG [3] та ENISA [2]; методи математичного моделювання – для побудови стохастичної моделі оцінювання ризику з використанням апарату теорії ймовірностей і введенням детермінованих функцій адаптивності атак; методи теорії інформації – для кількісного оцінювання невизначеності та нерівномірності розподілу каналів фішингу на основі ентропії Шеннона; компаративний аналіз – для порівняння ефективності традиційних підходів до фішингу та атак нового покоління, зокрема Deepfake та AI-driven phishing.

Огляд літератури та аналіз попередніх досліджень сучасний науковий дискурс (Jain & Gupta [5], Khonji [6]) зосереджений на автоматизації виявлення загроз [8, 12]. Проте спостерігається фрагментарність: технічні аспекти часто розглядаються окремо від психологічних. Значний внесок у вивчення соціальної інженерії зробив К. Греднегі [4], довівши, що психологічне маніпулювання є ефективнішим за прямий технічний злам. Головний недолік наявних підходів – відсутність інтегрованої моделі «людина + канал + АІ», що враховує динаміку адаптації атаки в часі.

Математичне моделювання ризиків фішингу: розробка та аналіз базової моделі. Математичне моделювання є ключовим інструментом формалізації процесів кіберзагроз, зокрема фішингових атак, що поєднують технічні та поведінкові аспекти. Основною метою моделювання є кількісна оцінка ймовірності успішної атаки з урахуванням множини факторів, які впливають на її результативність.

У даному дослідженні запропоновано багатофакторний підхід, який дозволяє інтегрувати рівень захищеності системи, характеристики каналу доставки та складність атаки в єдину формалізовану модель.

Базова модель. Ймовірність успішного здійснення фішингової атаки визначається наступною залежністю:

$$P_{success} = (1 - S) \cdot E \cdot C \quad (1)$$

де:

S – рівень захищеності системи (0-1);

E – ефективність соціальної інженерії (0-1);

C – коефіцієнт каналу атаки (0-1).

Аналіз параметрів:

1. При високому рівні захисту ($S \rightarrow 1$) ризик знижується.

2. Параметр E залежить від якості підготовки атаки.

3. C враховує довіру до каналу (email, SMS, voice тощо).

Розширення моделі. Модель може бути доповнена часовим фактором та адаптивністю атак:

$$P = (1 - S) \cdot E \cdot C \cdot A(t) \quad (2)$$

де $A(t)$ – функція адаптації атаки в часі.

Найбільший вплив має параметр, який змінюється швидше; фактор часу стає критичним при активній адаптації атаки.

Аналіз даних. У межах дослідження було проведено візуалізацію ключових показників, що відображають динаміку та структуру фішингових загроз у період активної цифрової трансформації (2020-2024 рр.).

На Рисунку 1 представлена динаміка зростання кількості зафіксованих фішингових інцидентів у глобальному масштабі.

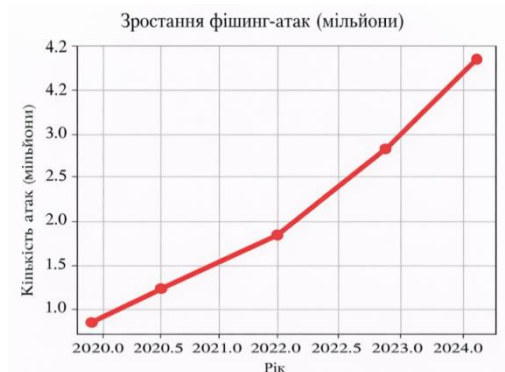


Рис. 1. Динаміка фішингових атак (млн).

Графік демонструє чітко виражений висхідний тренд: від 1,0 млн атак у 2020 році до пікового значення 4,2 млн у 2024 році. Особливу увагу слід звернути на прискорення темпів зростання після 2022 року (коефіцієнт зростання $k > 1,5$), що корелює з широким впровадженням інструментів генеративного штучного інтелекту, які дозволили зловмисникам автоматизувати створення контенту. Середньорічний темп зростання становить 35-45 %, що свідчить про експоненційний характер процесу. Такі дані підтверджуються звітами APWG Phishing Report 2024 [3] та ENISA Threat Landscape 2023 [2].

Аналіз векторів та каналів поширення. Для деталізації механізмів доставки фішингових повідомлень було розроблено комплексну інфографіку (Рисунок 2), що поєднує кількісні дані та рівні ризику.



Рис. 2. Розподіл каналів фішингу (%).

Візуалізація розподілу каналів за 2024 рік підкреслює домінуючу роль електронної пошти (60%), проте вказує на критичну небезпеку нових векторів. SMS-фішинг (смішинг) та голосовий фішинг (вішинг) сукупно складають 40% атак, при цьому вони характеризуються вищим рівнем конверсії через меншу захищеність мобільних пристроїв порівняно з корпоративними поштовими серверами. Часова теплокарта атак на рисунку вказує на нерівномірність активності, що підтверджує ентропійну природу вибору каналу зловмисниками.

Таблиця 1

Порівняльна характеристика основних комунікаційних каналів фішингу в умовах цифрової трансформації

Канал	Частка (%)	Ризик (оцінка)
Email	60	Високий
SMS	20	Дуже високий
Voice	10	Критичний
Social Media	10	Середній



Email залишається домінуючим каналом, проте SMS- і voice-атаки демонструють найвищу конверсію через швидкість та відсутність фільтрів.

Для кількісної оцінки ступеня невизначеності та нерівномірності вибору зловмисниками каналів поширення атак у роботі застосовано апарат теорії інформації, зокрема ентропію Шеннона. Даний підхід дозволяє визначити, наскільки диверсифікованою є стратегія зловмисників або, навпаки, наскільки вона зосереджена на конкретних векторах.

Математично ентропія системи розподілу каналів H обчислюється за формулою:

$$H = - \sum_{i=1}^n p_i \log_2 p_i \quad (3)$$

де p_i – частка (ймовірність) використання i -го каналу зв'язку (згідно з даними Табл. 1).

Розрахунок:

Виходячи з отриманих емпіричних даних ($p_1 = 0.6$ для Email, $p_2 = 0.2$ для SMS, $p_3 = 0.1$ для Voice, $p_4 = 0.1$ для Social Media), маємо:

$$H = -(0.6 \cdot \log_2 0.6 + 0.2 \cdot \log_2 0.2 + 0.1 \cdot \log_2 0.1 + 0.1 \cdot \log_2 0.1) \approx 1.57 \text{біт}$$

Інтерпретація результатів: отримане значення $H \approx 1.57$ при максимально можливому для чотирьох каналів $H_{\max} = \log_2 4 = 2.0$ свідчить про помірну впорядкованість системи.

Зниження показника ентропії відносно максимуму вказує на наявність домінуючого вектора. У контексті кібербезпеки це дозволяє зробити такі висновки: прогнозованість загрози, оскільки розподіл не є рівномірним («білим шумом»), що дає можливість превентивно фокусувати ресурси захисту (до 60%) на фільтрації поштового трафіку; спеціалізація атак, адже попри появу нових методів (Deepfake, Voice AI) зловмисники зберігають консервативний підхід, використовуючи Email як найбільш економічно ефективний канал з низьким порогом входу; інформаційна надмірність, бо різниця між $H_{\max}$ та H свідчить про певну «надмірність» стратегії атак, що дозволяє системам моніторингу виявляти закономірності (паттерни) навіть в умовах великих масивів даних. Порівняльний аналіз технік фішингу. Для визначення найбільш небезпечних векторів атак в умовах сучасної цифрової трансформації було проведено компаративний аналіз основних типів фішингу. Порівняння базується на чотирьох метриках: рівень технічної складності, конверсія (ефективність), ступінь інтеграції штучного інтелекту та результируючий рівень ризику для інформаційної системи (Табл. 2).

Таблиця 2

Компаративний аналіз фішингових атак із урахуванням складності реалізації, ефективності та впливу штучного інтелекту

Тип атаки	Рівень складності	Ефективність	Вплив AI	Ризик
Масовий фішинг	Низький	Середній	Низький	Середній
Spear phishing	Високий	Високий	Середній	Високий
AI phishing	Дуже високий	Дуже високий	Високий	Критичний
Deepfake phishing	Критичний	Максимальний	Дуже високий	Критичний

Результати компаративного аналізу технік фішингу, наведені у Таблиці 2, у поєднанні зі статистичними даними щодо динаміки зростання кількості інцидентів (Рисунок 1), переконливо підтверджують гіпотезу про зміну парадигми сучасних кіберзагроз. Експоненційне зростання кількості атак з 1,0 до 4,2 млн за останні п'ять років корелює не лише з технічною складністю, а й зі швидкістю адаптації зловмисників до захисних механізмів інформаційних систем.

Дані Таблиці 2 підтверджують необхідність включення коефіцієнта адаптивності $A(t)$ до загальної формули ризику, оскільки саме він відображає якісний перехід від «статистичного» масового фішингу до критично складних, AI-керованих сценаріїв.

Для візуалізації взаємозв'язку між еволюцією технік фішингу, зростанням ризику та запропонованими факторами математичної моделі розроблено структурно-логічну схему (Рисунок 3).

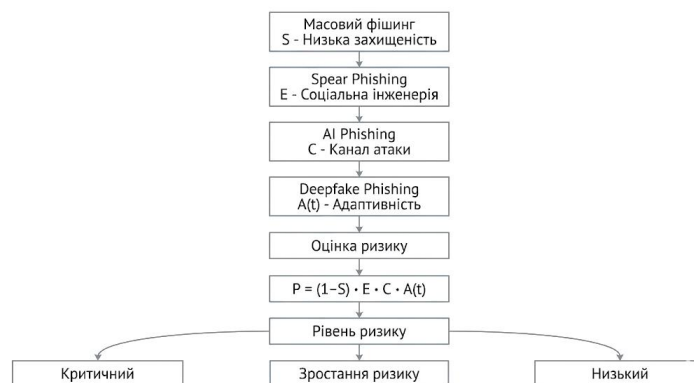


Рис. 3. Концептуальна модель еволюції фішингових атак: від масових розсилок S до критичних Deepfake-сценаріїв $A(t)$

Як видно з Рисунок 3, на початковому етапі «Масового фішингу» ризик (P) мінімізується високим рівнем системного захисту (S) та низькою ефективністю соціальної інженерії (E). У той час як на етапі «Deepfake Phishing», AI виступає мультиплікатором $A(t)$, що посилює складність атаки (C) та робить ризик (P) критичним, навіть за умови значних інвестицій у технічні засоби захисту інформації.

Інтеграція технологій штучного інтелекту (AI) до арсеналу кіберзлочинців спричинила докорінну зміну ландшафту соціальної інженерії[10], перетворивши фішинг на високотехнологічну адаптивну загрозу. Сучасний «фішинг нового покоління» базується на синергії кількох критичних напрямів використання неймереж. По-перше, застосування великих мовних моделей (LLM) дозволяє генерувати персоналізований контент з бездоганною граматичною точністю та ідеальним дотриманням офіційно-ділової стилістики. Це фактично нівелює один із традиційних маркерів виявлення загрози – мовну невідповідність та помилки, які раніше дозволяли користувачам ідентифікувати шкідливі повідомлення.

По-друге, стрімкий розвиток технологій синтезу медіаданих відкрив шлях до використання Deepfake-інструментів. Можливість імітації голосу та відеозображення керівництва компаній або довірених осіб у реальному часі під час відеоконференцій робить людський фактор довіри (H) критично вразливим, оскільки жертва втрачає здатність до візуальної та аудіальної верифікації співрозмовника. Водночас інтелектуальна автоматизація дозволяє зловмисникам масштабувати складні цільові атаки (Spear Phishing), трансформуючи їх у масові кампанії, де кожне окреме повідомлення адаптується під контекст конкретної жертви без залучення значних людських ресурсів з боку атакуючих.

Ефективність таких операцій суттєво підсилюється автоматизованим OSINT-аналізом, у межах якого алгоритми машинного навчання миттєво збирають та систематизують дані про об'єкт атаки з відкритих джерел, соціальних мереж та публічних реєстрів, забезпечуючи максимальну релевантність сценарію впливу. У контексті запропонованої математичної моделі штучний інтелект виступає ключовим мультиплікатором адаптивності $A(t)$. Це дозволяє атаці динамічно обходити статичні фільтри безпеки та системи проактивного захисту шляхом безперервної зміни сигнатур і семантичних структур контенту, що робить прогнозування та нейтралізацію таких загроз пріоритетним завданням для сучасних систем кіберзахисту.

Для верифікації запропонованої математичної моделі $P = (1 - S) \cdot E \cdot C \cdot H \cdot A(t)$ було проведено чисельний експеримент, результати якого наведено в табл. 3

Таблиця 3

Результати моделювання ймовірності успіху фішингових атак для різних сценаріїв стану інформаційної безпеки

Параметри сценарію	S (Захист)	E (Канал)	C (Складність)	H (Людина)	A(t) (AI)	P (Ймовірність)
Сценарій А (Secure)	0,8	0,7	0,6	0,5	1,0	0,042
Сценарій Б (Vulnerable)	0,3	0,9	0,8	0,9	1,5	0,756

У Сценарії Б враховано підвищену адаптивність AI $A(t) = 1.5$, що відображає здатність атаки еволюціонувати в часі.



Детальний аналіз проведеної симуляції дозволяє виявити критичні закономірності у зміні рівня ризику залежно від вхідних параметрів моделі. Порівняльний аналіз представлених сценаріїв свідчить про те, що одночасне зниження показника технічного захисту (S) та підвищення вразливості людського фактора (H) у поєднанні з активним застосуванням інструментів штучного інтелекту призводить до експоненційного зростання результуючого ризику (P), який збільшується більш ніж у 18 разів порівняно з базовим захищеним станом.

Особливої уваги заслуговує чинник критичної адаптивності загроз. Сценарій «Vulnerable» (Б) наочно демонструє, що навіть за умови збереження певних базових засобів кібербезпеки, високий показник адаптивності атаки $A(t)$ у синергії з психологічною невідповідністю користувача роблять інформаційну систему фактично незахищеною. У таких умовах ймовірність успішної реалізації фішингового сценарію перевищує 75 %, що вказує на неспроможність статичних методів захисту протистояти динамічним AI-загрозам.

Сформульовані висновки для стратегії захисту підкреслюють концептуальну обмеженість виключно технічного підходу. Результати симуляції доводять, що інвестиції в технологічний сегмент (S) мають спадну граничну ефективність, якщо вони не супроводжуються заходами з мінімізації людського фактора (H) через спеціалізовані тренінги та підвищення цифрової грамотності. Таким чином, сучасна архітектура кібербезпеки повинна базуватися на динамічних механізмах захисту, які здатні враховувати мінливу природу адаптивних алгоритмів штучного інтелекту та психологічні особливості об'єктів атаки.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Проведене дослідження дозволяє констатувати, що в умовах сучасної цифрової трансформації фішинг перестав бути примітивною технікою масових розсилок і перетворився на складну, багатофакторну систему адаптивного впливу. Аналіз динаміки загроз продемонстрував стійкий висхідний тренд: зростання кількості інцидентів з 1,0 млн у 2020 році до 4,2 млн у 2024 році свідчить про перехід кіберзлочинності на рейки інтелектуальної автоматизації. Отримані дані підтверджують, що середньорічний темп зростання кількості атак безпосередньо корелює з доступністю інструментів генеративного штучного інтелекту, які дозволяють зловмисникам масштабувати персоналізовані сценарії впливу.

Структурна декомпозиція каналів поширення за допомогою ентропійного аналізу Шеннона виявила помірну впорядкованість фішингових стратегій. Показник ентропії $H \approx 1.57$ біт вказує на те, що попри диверсифікацію методів, основний вектор загрози залишається прогнозованим і зосередженим у сегменті електронної пошти. Водночас зростаюча частка SMS та голосових комунікацій, підсилені технологіями Deepfake, створює нові виклики для систем автентифікації, оскільки ці канали характеризуються значно вищим рівнем міжособистісної довіри та меншим ступенем технічної фільтрації.

Розроблена та верифікована у роботі багатофакторна математична модель $P = (1 - S) \cdot E \cdot C \cdot H \cdot F(t)$ доводить, що ефективність кіберзахисту сьогодні залежить не лише від технологічних параметрів системи, а й від здатності враховувати динамічну адаптивність атак. Результати симуляції різних сценаріїв безпеки продемонстрували, що синергія низького рівня захисту та високої адаптивності штучного інтелекту здатна підвищити ймовірність успішного зламу у 18 разів. Це підкреслює критичну роль людського фактора як найбільш вразливої ланки: навіть за умови високої технологічної стійкості інфраструктури, психологічна невідповідність користувача в поєднанні з AI-персоналізацією сценаріїв залишає ризик реалізації загрози на критично високому рівні.

Таким чином, запропонований у статті підхід до оцінки ризиків дозволяє перейти від статичного аналізу загроз до динамічного прогнозування, що є необхідною умовою для побудови адаптивних систем кіберзахисту нового покоління. Подальші дослідження у цьому напрямі мають бути зосереджені на розробці алгоритмів автоматизованого моніторингу параметра адаптивності в режимі реального часу.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Verizon Business. (2024). *2024 data breach investigations report* (100 pp.). <https://www.verizon.com/business/resources/reports/dbir/>
2. European Union Agency for Cybersecurity. (2023). *ENISA threat landscape 2023* (142 pp.).
3. Anti-Phishing Working Group. (2024). *Phishing activity trends report: 4th quarter 2023*. <https://apwg.org/trendsreports/>
4. Hadnagy, C. (2018). *Social engineering: The science of human hacking* (2nd ed.). Wiley.



5. Jain, A. K., & Gupta, B. B. (2022). Phishing detection using machine learning techniques: A comprehensive survey. *Computers & Security*, 114, 102577.
6. Khonji, M., Iraqi, Y., & Jones, A. (2013). Phishing detection: A literature survey. *IEEE Communications Surveys & Tutorials*, 15(4), 2091–2121.
7. Lastdrager, E. E. (2014). Achieving a consensual definition of phishing. *IEEE Security & Privacy*, 12(6), 92–95.
8. Gupta, B. B., & Tewari, A. (2020). A survey of machine learning techniques for detection of phishing. *Journal of Ambient Intelligence and Humanized Computing*, 11, 1561–1573.
9. Ferreira, A., & Lenzini, G. (2021). Socio-technical study on phishing. *Journal of Computer Security*, 29(2), 165–191.
10. Opara, C., Chen, Z., & Goldsmith, J. (2024). *Phishing detection with generative AI: Challenges and opportunities*. arXiv. <https://arxiv.org/abs/2401.00001>
11. Williams, E. J., Hinds, J., & Joinson, A. N. (2018). Exploring susceptible individuals' responses to personal and impersonal phishing. *Computers in Human Behavior*, 87, 227–237.
12. Alkhalil, A., et al. (2021). Phishing attacks: Recent comprehensive study and a new model for automated detection. *Papers in Computer Science*, 3(2), 12–25.
13. Chiew, K. L., et al. (2018). A survey of phishing attacks: Their types, vectors and technical approaches. *Expert Systems with Applications*, 106, 1–20.
14. Vishwanath, A., et al. (2011). Why do people click on phishing links? *Communications of the ACM*, 54(12), 52–59.
15. Abroshan, H., et al. (2022). Phishing: The role of human error and psychological traits. *Frontiers in Psychology*, 13, 1–15.

**Kryvokulska Olha**

Senior Lecturer of the Department of Cybersecurity
National University "Kyiv Aviation Institute", Kyiv, Ukraine
ORCID:0009-0003-8518-6915
olha.kryvokulska@npp.kai.edu.ua

Yakovenko Olesia

Senior Lecturer of the Department of Cybersecurity
National University "Kyiv Aviation Institute", Kyiv, Ukraine
ORCID:0000-0003-2998-9767
olesia.yakovenko@npp.kai.edu.ua

Marchenko Yaroslav

Assistant of the Department of Cybersecurity
National University "Kyiv Aviation Institute", Kyiv, Ukraine
ORCID:0009-0006-9457-1235
yaroslav.marchenko@npp.kai.edu.ua

Yakymchuk Yevhenii

Assistant of the Department of Cybersecurity
National University "Kyiv Aviation Institute", Kyiv, Ukraine
ORCID:0009-0009-9736-5260
yevhenii.yakymchuk@npp.kai.edu.ua

NEXT GENERATION PHISHING: THE EVOLUTION OF SOCIAL ENGINEERING IN THE CONTEXT OF DIGITAL TRANSFORMATION

Abstract. This article analyzes current trends in phishing, including the use of artificial intelligence technologies to automate the creation of phishing content, increase the level of personalization, and optimize attack scenarios. It examines the influence of the human factor as a key vulnerability that determines the effectiveness of social engineering. A quantitative analysis of the dynamics of phishing attacks in recent years is conducted, demonstrating an exponential increase in the number of incidents and indicating the active use of automated tools by cybercriminals. The distribution of channels used to carry out phishing attacks -specifically email, SMS, voice communications, and social networks – was examined separately, allowing for the determination of their relative effectiveness and risk level.

A multifactorial mathematical model for assessing the probability of a phishing attack's success is proposed, which takes into account the level of security of the information system, the effectiveness of the social engineering methods used, the characteristics of the distribution channel, as well as user behavioral patterns. Unlike existing approaches, the model has been extended by introducing a time-dependent adaptation function, which allows for changes in attack parameters over time. Additionally, elements of probability theory are used to model attack intensity, and the information entropy approach is employed to assess uncertainty in the selection of phishing channels.

The results confirm that modern phishing is a complex, multidimensional threat that combines technical, behavioral, and informational aspects. The proposed model allows for the formalization of the risk assessment process, enables the prediction of attack effectiveness, and can be used to develop adaptive cybersecurity systems. The practical significance of the study lies in the applicability of the results to information security management systems, as well as in raising user awareness of modern social engineering methods.

Keywords: phishing, social engineering, artificial intelligence, cybersecurity, mathematical modeling.

REFERENCES (TRANSLATED AND TRANSLITERATED)

1. Verizon Business. (2024). *2024 data breach investigations report* (100 pp.). <https://www.verizon.com/business/resources/reports/dbir/>
2. European Union Agency for Cybersecurity. (2023). *ENISA threat landscape 2023* (142 pp.).



3. Anti-Phishing Working Group. (2024). *Phishing activity trends report: 4th quarter 2023*. <https://apwg.org/trendsreports/>
4. Hadnagy, C. (2018). *Social engineering: The science of human hacking* (2nd ed.). Wiley.
5. Jain, A. K., & Gupta, B. B. (2022). Phishing detection using machine learning techniques: A comprehensive survey. *Computers & Security*, 114, 102577.
6. Khonji, M., Iraqi, Y., & Jones, A. (2013). Phishing detection: A literature survey. *IEEE Communications Surveys & Tutorials*, 15(4), 2091–2121.
7. Lastdrager, E. E. (2014). Achieving a consensual definition of phishing. *IEEE Security & Privacy*, 12(6), 92–95.
8. Gupta, B. B., & Tewari, A. (2020). A survey of machine learning techniques for detection of phishing. *Journal of Ambient Intelligence and Humanized Computing*, 11, 1561–1573.
9. Ferreira, A., & Lenzini, G. (2021). Socio-technical study on phishing. *Journal of Computer Security*, 29(2), 165–191.
10. Opara, C., Chen, Z., & Goldsmith, J. (2024). *Phishing detection with generative AI: Challenges and opportunities*. arXiv. <https://arxiv.org/abs/2401.00001>
11. Williams, E. J., Hinds, J., & Joinson, A. N. (2018). Exploring susceptible individuals' responses to personal and impersonal phishing. *Computers in Human Behavior*, 87, 227–237.
12. Alkhalil, A., et al. (2021). Phishing attacks: Recent comprehensive study and a new model for automated detection. *Papers in Computer Science*, 3(2), 12–25.
13. Chiew, K. L., et al. (2018). A survey of phishing attacks: Their types, vectors and technical approaches. *Expert Systems with Applications*, 106, 1–20.
14. Vishwanath, A., et al. (2011). Why do people click on phishing links? *Communications of the ACM*, 54(12), 52–59.
15. Abroshan, H., et al. (2022). Phishing: The role of human error and psychological traits. *Frontiers in Psychology*, 13, 1–15.

Отримано редакцією журналу / Received: 24.02.26

Прорецензовано / Revised: 03.03.26

Схвалено до друку / Accepted: 24.09.26



This work is licensed under Creative Commons Attribution-noncommercial-sharealike 4.0 International License.