



DOI 10.28925/2663-4023.2025.30.990

УДК 004.056.53:351.853

Гарасимчук Олег Ігорович

кандидат технічних наук, доцент

доцент кафедри захисту інформації

Національний Університет "Львівська політехніка", Львів, Україна

ORCID: 0000-0002-8742-8872

oleh.i.harasymchuk@lpnu.ua

Оліярник Юлія Сергіївна

Студентка кафедри захисту інформації

Національний Університет "Львівська політехніка", Львів, Україна

ORCID: 0009-0005-1018-651X

yuliia.oliarnyk.kb.2023@lpnu.ua

Нестор Андрій Олександрович

Студент кафедри захисту інформації

Національний Університет "Львівська політехніка", Львів, Україна

ORCID: 0009-0004-1601-9598

andrii.nestor.kb.2023@lpnu.ua

Наконечний Тарас Ігорович

Аспірант кафедри захисту інформації

Національний Університет "Львівська політехніка", Львів, Україна

ORCID: 0009-0003-4487-9424

taras.i.nakonechnyi@lpnu.ua

ПСИХОЛОГІЧНІ МЕТОДИ ШАХРАЙСТВА В КІБЕРПРОСТОРІ ТА СПОСОБИ ЇМ ПРОТИДІЯТИ

Анотація. У статті досліджено методи соціальної інженерії, які використовуються зловмисниками для отримання несанкціонованого доступу до конфіденційної інформації та маніпулювання поведінкою жертв. Розглянуто основні види атак, такі як фішинг, вішинг, смішинг, прітекстинг, спір-фішинг і вейлінг, а також їхні особливості, механізми реалізації та способи обману користувачів. Особливу увагу приділено психологічним аспектам соціальної інженерії, включаючи вплив страху, довіри, терміновості, соціального доказу та когнітивних упереджень на процес прийняття рішень. Окреслено сучасні підходи до захисту від соціоінженерних атак, що включають поєднання технологічних та освітніх методів. Запропоновано заходи з підвищення цифрової грамотності користувачів, розробки політик інформаційної безпеки, застосування багатфакторної автентифікації, систем аналізу поведінки користувачів та штучного інтелекту для виявлення загроз. Особливу увагу приділено використанню великих мовних моделей для ідентифікації шахрайських схем та автоматизації кібербезпеки. Результати дослідження свідчать про необхідність комплексного підходу до захисту від атак соціальної інженерії, який передбачає синергію між технологічними інструментами та людським фактором. Запропоновані рекомендації спрямовані на мінімізацію ризиків та підвищення загального рівня безпеки в цифровому середовищі.

Ключові слова: соціальна інженерія, психологічні маніпуляції, кібербезпека, фішинг, вішинг, інформаційна безпека, штучний інтелект, захист від шахрайства.

**ВСТУП**

У сучасному цифровому світі інформаційна безпека стала однією з ключових проблем як для окремих користувачів, так і для великих організацій. З розвитком інформаційних технологій зростає і кількість загроз, пов'язаних із витоком даних, фінансовими шахрайствами та порушенням конфіденційності. Однак, попри вдосконалення технічних засобів захисту, людський фактор залишається найбільш вразливим елементом у системах безпеки. Саме тому соціальна інженерія, яка використовує методи психологічного маніпулювання, є одним із найефективніших способів отримання несанкціонованого доступу до інформації. Зловмисники, застосовуючи різноманітні методи впливу, змушують жертву самостійно передавати конфіденційні дані або виконувати певні дії, що призводить до значних втрат.

Соціальна інженерія охоплює широкий спектр методів, зокрема фішинг, вішинг, смішинг, байтинг, претекстинг та інші техніки, що базуються на використанні довіри, страху, цікавості або терміновості [1]. Ці атаки можуть бути спрямовані як на окремих осіб, так і на великі компанії, державні установи та фінансові організації. Найбільш небезпечним є те, що навіть найсучасніші засоби захисту не можуть повністю гарантувати безпеку, якщо користувач самостійно розкриває свої дані або відкриває доступ до системи. Саме тому важливим завданням є дослідження механізмів таких атак та пошук ефективних методів їхньої нейтралізації.

Протидія атакам соціальної інженерії потребує комплексного підходу, що поєднує технологічні рішення, навчальні програми та формування цифрової обізнаності користувачів. Зокрема, використання багатофакторної аутентифікації, політики мінімальних привілеїв, сучасних систем виявлення загроз, а також навчання персоналу методам розпізнавання шахрайських схем може значно знизити ризики. Однак, з огляду на постійний розвиток тактик соціальної інженерії, необхідно не лише впроваджувати окремі методи захисту, а й досліджувати ефективність їхнього комбінування для створення максимально надійної системи кібербезпеки [2-3].

Аналіз останніх досліджень і публікацій. Проблеми загрози, що виникає внаслідок використання соціальної інженерії, детально розглядається у багатьох працях як вітчизняних так і закордонних дослідників, зокрема в роботах [1 – 3] акцентується увага на її небезпеці. У джерелах [1, 3, 6, 7] проведено глибокий аналіз типових атак, шаблонів та сценаріїв, що використовуються зловмисниками. Водночас, стратегії протидії соціоінженерним загрозам досліджено у працях [2–10]. Робота [5] присвячена виокремленню ключових суб'єктів та об'єктів впливу в соціальній інженерії. Тут суб'єкт визначається як атакуючий, тоді як об'єктом виступає особа чи система, що здійснює захист. Дії зловмисника здебільшого зводяться до застосування визначених шаблонів та сценаріїв впливу, спрямованих на досягнення переваги над захисником. Додатково розглянуто методики тестування на проникнення, зокрема, оцінку вразливостей людини [7]. Для цього рекомендується використання Social-Engineer Toolkit – як окремого програмного комплексу, так і вбудованого модуля в Kali Linux. У праці [8] пропонується підхід, що ґрунтується на залученні користувачів до процесу виявлення та сповіщення про потенційні загрози соціальної інженерії, сприймаючи людину як сенсор безпеки (CogniSense). Концептуальна структура захисту від соціальної інженерії (Social Engineering Defensive Framework) була розроблена у роботі [9]. Вона передбачає реалізацію протидії через чотири незалежні етапи. Окремий акцент зроблено на



підвищенні рівня обізнаності користувачів шляхом нагадувань та проведення спеціалізованих навчальних заходів [10].

Проблематика дослідження полягає у зростанні кількості та складності атак соціальної інженерії, які використовують психологічні вразливості користувачів замість технічних недоліків систем. Незважаючи на наявність різноманітних засобів захисту, більшість з них залишаються малоефективними через людський фактор та відсутність комплексного підходу до протидії. Тому актуальним є пошук оптимальних комбінацій технічних, організаційних і освітніх методів, здатних забезпечити підвищений рівень стійкості до подібних загроз.

Мета роботи: дослідити сучасні методи соціальної інженерії, визначити найбільш поширені техніки атак та оцінити ефективність існуючих засобів захисту з акцентом на можливість комбінованого використання різних методів протидії для підвищення рівня інформаційної безпеки.

Для досягнення поставленої мети необхідно вирішити такі завдання:

1. Провести класифікацію основних видів та технік атак соціальної інженерії.
2. Проаналізувати механізми впливу людського фактора на успішність соціально-інженерних атак.
3. Оцінити ефективність сучасних технічних, організаційних та освітніх заходів протидії.
4. Дослідити можливість комбінування різних методів захисту для зменшення ризиків атак.
5. Сформулювати практичні рекомендації щодо підвищення рівня захисту від загроз соціальної інженерії.

ОСНОВНА ЧАСТИНА

Соціальна інженерія – це сукупність методів та прийомів, спрямованих на маніпулювання людьми з метою отримання конфіденційної інформації, доступу до систем або виконання певних дій, вигідних зловмисникові. Основна особливість цього явища полягає у використанні людської психології як основного вектора атаки, а не технічних вразливостей інформаційних систем.

У контексті кібербезпеки соціальна інженерія відіграє ключову роль, оскільки навіть найнадійніші системи захисту можуть виявитися безсилими перед людським фактором. Досвідчені зловмисники використовують методи переконання, викликають довіру або паніку, аби змусити людину розкрити паролі, передати важливі дані або відкрити доступ до критично важливих ресурсів. Багато гучних випадків зламу інформаційних систем починалися саме з атаки соціальної інженерії, коли зловмисники маніпулювали співробітниками компаній або окремими користувачами [11-13].

З розвитком цифрових технологій та широким розповсюдженням соціальних мереж, соціальна інженерія отримала нові можливості для атак. Відкритий доступ до персональних даних, публічні фотографії, інформація про професійну діяльність – усе це створює сприятливий ґрунт для шахрайських схем. Саме тому усвідомлення загроз, навчання персоналу та впровадження заходів протидії є критично важливими складовими кібербезпеки будь-якої організації.

Успішність атак соціальної інженерії значною мірою залежить від використання психологічних методів впливу. Зловмисники застосовують різні техніки маніпуляції, зокрема створення відчуття терміновості, виклик страху, авторитету або довіри, аби



спонукати людину до дій, які суперечать її власним інтересам або політиці безпеки організації.

Основою всіх методів соціальної інженерії є здатність впливати на когнітивні процеси людини, використовуючи її емоції, несвідомі упередження, довіру до авторитетів і природну схильність до соціальної взаємодії. Людина – істота емоційна, і саме емоції часто беруть гору над логічним мисленням. Вплив через страх, довіру, поспіх або навіть цікавість дозволяє зловмисникам легко маніпулювати жертвою.

Одним із найефективніших методів є використання страху, оскільки в стані стресу людина приймає менш обдумані рішення. Наприклад, повідомлення про заблокований банківський рахунок або загрозу витоку конфіденційної інформації викликає паніку, і жертва може миттєво передати дані, не аналізуючи ситуацію. Це пов'язано з активацією механізму "бий або тікай", який еволюційно закладений у людську природу. У такому стані мозок працює у спрощеному режимі, орієнтуючись на негайне усунення загрози.

Ще одним ключовим чинником є довіра до авторитетів. Людина схильна підкорятися тим, кого вона вважає носіями влади або експертами у певній сфері. Якщо зловмисник видає себе за представника технічної підтримки, юриста чи навіть керівника, жертва може автоматично виконати його вимоги без додаткових перевірок. Це явище пояснюється ефектом соціального впливу: людина змалку привчається слідувати інструкціям авторитетних осіб і вважає їхні слова істинними без сумніву.

Маніпуляція через почуття взаємності є ще одним потужним інструментом. Люди відчують себе зобов'язаними віддати щось у відповідь, якщо їм надали послугу або виявили доброзичливість. Зловмисник може спершу створити враження допомоги – наприклад, запропонувати вирішити уявну проблему жертви, а потім попросити про послугу у відповідь, яка насправді стане компрометацією даних чи системи безпеки.

Також широко використовується принцип соціального доказу. Люди схильні довіряти тому, що роблять інші, особливо якщо інформація підтверджена групою або суспільством. Наприклад, у шахрайських схемах часто використовують підроблені відгуки або заяви про те, що "багато ваших колег вже надали ці дані". Це знижує критичне мислення жертви і робить її більш схильною до прийняття потрібного зловмиснику рішення.

Зловмисники також експлуатують людську цікавість, яка є природною рисою мислення. Якщо жертва отримує повідомлення, що містить інтригуючу або скандальну інформацію, вона може натиснути на посилання, яке веде на шкідливий сайт або програму. Цікавість є сильною мотиваційною силою, особливо коли інформація виглядає унікальною чи ексклюзивною.

Ще одним важливим аспектом соціальної інженерії є ефект поспіху. Люди значно частіше припускаються помилок, коли перебувають під тиском часу. Атаки часто здійснюються під приводом терміновості: "Терміново змініть пароль, інакше ваш обліковий запис буде видалено", або "Ваш колега вже надіслав документи, вам залишилось лише підтвердити". В умовах дефіциту часу жертва менш схильна до перевірки деталей і приймає рішення на автоматизованому рівні.

Соціальні інженери майстерно комбінують ці психологічні методи, адаптуючи підхід залежно від жертви. Використання цих маніпуляційних технік дозволяє їм отримувати потрібну інформацію, обходити механізми безпеки та здійснювати атаки без необхідності втручання в технічні аспекти системи. Саме тому підвищення психологічної грамотності та розвиток критичного мислення є найкращим способом захисту від подібних загроз.

У межах соціоінженерного підходу людські вразливості розглядаються як сукупність слабких місць, базових потреб, схильностей та особистих пристрастей. Використання цих характеристик у маніпулятивних цілях дозволяє зловмисникам здобути несанкціонований доступ до конфіденційної інформації. У результаті формується нова модель поведінки жертви, що створює сприятливе середовище для реалізації загроз інформаційній безпеці. Це, своєю чергою, знижує ефективність систем захисту та їх здатність протидіяти таким впливам (див., наприклад, [2], рис. 1).



Рис. 1. Використання соціоінженерного підходу

Маніпулятивний вплив може проявлятися у різних формах, зокрема, таких як шахрайство, обман, афери, інтриги, містифікації та провокації [14]. Перед їх застосуванням проводиться ґрунтовний аналіз, що включає визначення змісту маніпулятивної тактики, її детальне планування, організацію та контроль виконання [2].

Як впливає з рис. 1, застосування соціоінженерного підходу передбачає цілеспрямований вплив на свідомість або підсвідомість людини всупереч її справжній волі, але за формальної згоди. Такий вплив дозволяє маніпулювати поведінкою осіб, що займають ключові ролі, зокрема керівництва, адміністраторів і звичайних користувачів, використовуючи їхні слабкі місця, інтереси, базові потреби, схильності, переконання, звички, а також психоемоційний стан. Саме експлуатація цих людських вразливостей і знаходить відображення у таких проявах, як шахрайство, обман, афери, інтриги, містифікації та провокації. Водночас кожний із зазначених форм маніпулятивного впливу передусє ретельне опрацювання, що включає визначення її сутності, детальне планування, організацію та контроль реалізації. Характер і особливості застосування цих методів змінюються залежно від специфіки атаки соціальної інженерії, зокрема [14]:

– *Формінг* (Forming) є одним із методів соціальної інженерії, який передбачає створення або модифікацію веб-форм для отримання конфіденційних даних користувачів. Зловмисники використовують спеціальні технічні засоби для підміни легітимних форм введення даних на фальшиві, що дозволяє їм перехоплювати інформацію ще до її потрапляння до законного отримувача. Основною метою такого методу є отримання облікових даних, паролів, фінансової інформації або інших критично важливих відомостей. Особливо поширеним є використання формінгу у поєднанні з фішинговими атаками, коли жертва отримує електронний лист із посиланням на підроблений сайт. Використання цього методу ускладнюється розширеним застосуванням двофакторної автентифікації, однак деякі його варіанти дозволяють навіть обходити такі захисні механізми.

– *Фішинг* (Phishing) є одним із найбільш поширених методів соціальної інженерії, який полягає у спробах шахраїв отримати конфіденційну інформацію користувачів



шляхом масового розсилання фальшивих повідомлень, що імітують офіційні звернення від банків, державних установ або великих компаній. Основна особливість фішингових атак полягає у використанні соціального фактору, коли жертва через довіру до джерела інформації сама надає свої особисті дані. Зловмисники можуть змушувати користувачів перейти за шкідливими посиланнями, завантажити шкідливі вкладення або ввести свої паролі на підроблених веб-ресурсах. Така атака може використовуватися як на індивідуальному рівні, так і проти цілих організацій, сприяючи витоку корпоративної інформації.

– *Прітекстінг* (Pretexting) є методом соціальної інженерії, що базується на попередньому створенні правдоподібної легенди, яка дозволяє зловмиснику здобути довіру жертви та отримати необхідну інформацію. Головною особливістю цієї техніки є не просто викрадення даних, а формування сценарію, який переконує людину добровільно надати інформацію, вважаючи, що вона взаємодіє з легітимною особою. Наприклад, шахрай може представитися співробітником банку та попросити підтвердити особисті дані або ж видати себе за технічного спеціаліста, який нібито допомагає вирішити проблему із безпекою акаунта. Оскільки прітекстінг часто використовує психологічні механізми впливу, його особливо складно виявити, а наслідки таких атак можуть бути критичними для інформаційної безпеки компаній.

– *Смішінг* (Smishing) є варіацією фішингових атак, які здійснюються через SMS-повідомлення. Головна ідея такого методу полягає у надсиланні жертві текстових повідомлень, що містять посилання на шкідливі ресурси або заклики до негайних дій, спрямованих на отримання конфіденційних даних. Часто зловмисники маскують свої повідомлення під офіційні звернення від банків, поштових служб, державних установ або компаній, що надають різноманітні послуги. Особливо ефективними такі атаки є у періоди підвищеного навантаження на комунікаційні канали, наприклад під час великих знижок або кризових ситуацій, коли люди менше звертають увагу на потенційні загрози.

– *Вішінг* (Vishing) є методом шахрайства, що використовує телефонні дзвінки для отримання конфіденційної інформації. Зловмисники можуть представлятися співробітниками банків, служб підтримки, державних органів або навіть правоохоронних структур, створюючи ситуацію, яка спонукає жертву передати особисті дані або вчинити певні дії, наприклад, здійснити фінансовий переказ. Часто вішінг-атаки супроводжуються психологічним тиском: зловмисники можуть лякати жертву блокуванням рахунків, фінансовими штрафами або іншими негативними наслідками. Використання технологій зміни голосу та підробки номерів телефонів робить такі атаки ще більш складними для виявлення.

– *Спір-фішінг* (Spear Phishing) відрізняється від звичайного фішингу тим, що атака спрямована не на масову аудиторію, а на конкретну особу чи організацію. Основна мета таких атак – отримати доступ до важливої інформації, зокрема даних про корпоративні мережі, фінансові операції або особистих даних високопосадовців. Такі атаки зазвичай ретельно сплановані, включають персоналізовані звернення та імітують реальні комунікації, які жертва може очікувати отримати. Для підвищення ефективності атак зловмисники можуть використовувати попередньо зібрану інформацію про людину, що робить фальшиві повідомлення ще більш правдоподібними.

– *Вейлінг* (Whaling) це різновидом спір-фішингу, проте його основною мішенню стають керівники вищої ланки, топ-менеджери та інші особи, що мають доступ до критично важливих даних компанії. Такі атаки часто маскуються під офіційні запити від урядових установ, фінансових служб або партнерів, і можуть включати вимоги надати конфіденційну інформацію, підписати документи або здійснити фінансові перекази.



Вейлінг є одним із найнебезпечніших видів соціальної інженерії, оскільки він може спричинити значні фінансові втрати або витоки критично важливих корпоративних даних.

Несанкціоноване заволодіння інформацією шляхом застосування методів фішингу, фармінгу, смішингу, вішингу, спір-фішингу та вейлінгу відбувається через використання маніпулятивних технік, заснованих на обманних діях та шахрайських схемах (див., наприклад, [2,14], табл. 1). Водночас створення штучних обставин у межах прітекстінгу ґрунтується на застосуванні таких прийомів, як афера, інтрига, містифікація та провокація, що дозволяють зловмисникам імітувати правдоподібні сценарії для введення жертви в оману.

Таблиця 1.

Відображення форм маніпулятивного впливу в різновидах соціоінженерних атак

№ п/п	Різнovid соціоінженерної атаки	Форми маніпулятивного впливу					
		Шахрайство	Обман	Афера	Інтрига	Містифікація	Провокація
1.	Фішинг	+	+	–	–	–	–
2.	Фармінг	+	+	–	–	–	–
3.	Прітекстінг	+	+	+	+	+	+
4.	Смішинг	+	+	–	–	–	–
5.	Вішинг	+	+	–	–	–	–
6.	Спір фішинг	+	+	–	–	–	–
7.	Вейлінг	+	+	–	–	–	–

При оцінюванні рівня захищеності інформації в комп'ютерних системах із застосуванням соціоінженерного підходу важливо враховувати форми маніпулятивного впливу. Методи соціальної інженерії спрямовані на моделювання дій порушників, які, наприклад, можуть бути націлені на персонал організації. Соціальний інженер, намагаючись отримати відомості про комп'ютерні системи, використовує контактну інформацію з відкритих джерел, зокрема дані про користувачів: їхні прізвища, імена, посади. На основі зібраної інформації визначаються потенційні вразливості, через які можуть реалізовуватися загрози соціальної інженерії.

Основу методів соціальної інженерії становлять такі аспекти:

- закономірності, що впливають на свідомість людини;
- особливості аудиторії або середовища діяльності;
- недостатня обізнаність у термінах і предметних сферах інформаційної безпеки;
- психологічна нестійкість особистості, що проявляється у поведінкових шаблонах і може бути використана для маніпуляції через базові потреби, слабкі сторони, цінності та прагнення.



Більшість соціальних інженерів діють за схожими схемами та сценаріями атак. Вивчення їхніх методів дозволяє виділити рівні взаємодії з об'єктом впливу, серед яких: домінування, маніпулювання, суперництво та партнерство.

Методи соціальної інженерії можна класифікувати на дві основні категорії [2]:

- технічні
- психологічні

Технічні методи включають використання фішингових електронних листів, підроблених вебсайтів, шкідливих програм або пристроїв для перехоплення даних. Психологічні методи базуються на маніпуляції емоціями людини – страхом, довірою, авторитетом чи бажанням допомогти – з метою змусити її виконати небезпечні або небажані дії."

Віддалена соціальна інженерія здійснюється за допомогою сучасних засобів зв'язку, зокрема телефонного спілкування, що дозволяє зловмиснику зберігати анонімність і водночас безпосередньо впливати на ціль. Така форма взаємодії є особливо ефективною, оскільки пряма розмова не залишає часу на ретельний аналіз ситуації чи обмірковування можливих дій. Людина змушена оперативно ухвалювати рішення під тиском, створеним тоном, інтонацією та манерою розмови. Ці вербальні характеристики свідомо підбираються відповідно до психологічного портрета жертви, що підвищує ймовірність успішного отримання конфіденційної інформації. Залежно від ролі, яку зловмисник імітує під час взаємодії з ціллю, можна виокремити кілька типових моделей поведінки, що використовуються у соціоінженерних атаках:

– *Начальник* – людина, яка звикла керувати, не терпить зволікань і орієнтована на результат. Спілкування жорстке, командне, із нотками зверхності. Проблему зображають як дріб'язкову, яку слід розв'язати негайно. Замість прохань – лише чіткі вказівки. Будь-які сумніви чи перевіірочні питання зустрічаються незадоволенням або залякуванням.

– *Секретар* – здебільшого молода жінка з приємним голосом, яка виконує доручення керівника, не ставлячи зайвих запитань. Володіє інформацією про начальника та його справи, використовує перевірені чи складні для перевірки факти. Розмова ведеться м'яко, іноді з елементами флірту (якщо співрозмовник – чоловік). У разі відмови – демонстративне розчарування та скарга, що начальство буде незадоволене.

– *Технічний співробітник* – представник організації, що демонструє компетентність через застосування професійних термінів. Налаштований доброзичливо, однак поблажливо щодо клієнта. Основний меседж – усунення проблеми у його ж інтересах. Сумніви співрозмовника викликають подив, адже співпраця вигідна передусім йому. Можливе залякування серйозними наслідками у разі бездіяльності.

– *Користувач* – співробітник, зайнятий виконанням своїх завдань, який опинився у стресовій ситуації через неочікувану проблему. Його основна мета – якнайшвидше вирішити ситуацію, не вдаючись у деталі. Зловмисник використовує цей стан, підкреслюючи невідкладність дій і демонструючи готовність «допомогти», що змушує жертву довіритися.

Таким чином, соціальні інженери майстерно використовують різні психологічні тактики, адаптуючи стиль комунікації залежно від ролі жертви, щоб підштовхнути її до прийняття потрібних рішень.

Соціальна інженерія активно застосовується у цифровому середовищі, використовуючи можливості глобальної мережі Інтернет. Одним із найпоширеніших методів є маніпуляція через електронну пошту, де зловмисники використовують переконливі повідомлення, що змушують жертву розголошувати конфіденційну інформацію або виконувати певні дії. Також значну роль відіграють месенджери, які



дозволяють соціальному інженеру вести діалог у реальному часі, створюючи ефект терміновості або довіри. Крім того, форуми, блоги та чати слугують ефективним майданчиком для довготривалої взаємодії, де можна поступово впливати на свідомість жертви, формуючи потрібну поведінкову реакцію.

Окремою категорією соціальної інженерії є маніпуляція під час особистого контакту, що вважається складним і ризикованим методом. У такому випадку важливу роль відіграє не лише зміст розмови, а й невербальні фактори. Для досягнення потрібного ефекту соціальний інженер ретельно підбирає зовнішній вигляд, обираючи відповідні кольори одягу та стиль поведінки, які мають викликати довіру або створювати відчуття авторитетності. Манери, жести та міміка також адаптуються до ситуації, щоб забезпечити максимальний психологічний вплив. Важливим аспектом є просторове розташування під час спілкування, яке може використовуватися для встановлення домінування або, навпаки, для створення дружньої атмосфери.

Ефективність соціальної інженерії значною мірою залежить від здатності маніпулятора розпізнати вразливості своєї жертви. За голосом, інтонацією чи поведінкою можна визначити, яка риса характеру може бути використана для досягнення поставленої мети. Довірливі люди схильні без зайвих сумнівів приймати інформацію за правду, тоді як страх може спровокувати імпульсивні рішення. Жадібність спонукає людину діяти у власних інтересах, не замислюючись про можливі ризики, тоді як відкритість до комунікації робить її легкою мішенню для маніпуляцій. Зверхність може підштовхнути до прийняття рішень на емоціях, а милосердя – до необачної допомоги, що може бути використано у шахрайських схемах.

Ми пропонуємо наступну класифікацію ознак реалізації атак соціальної інженерії (рис. 2).

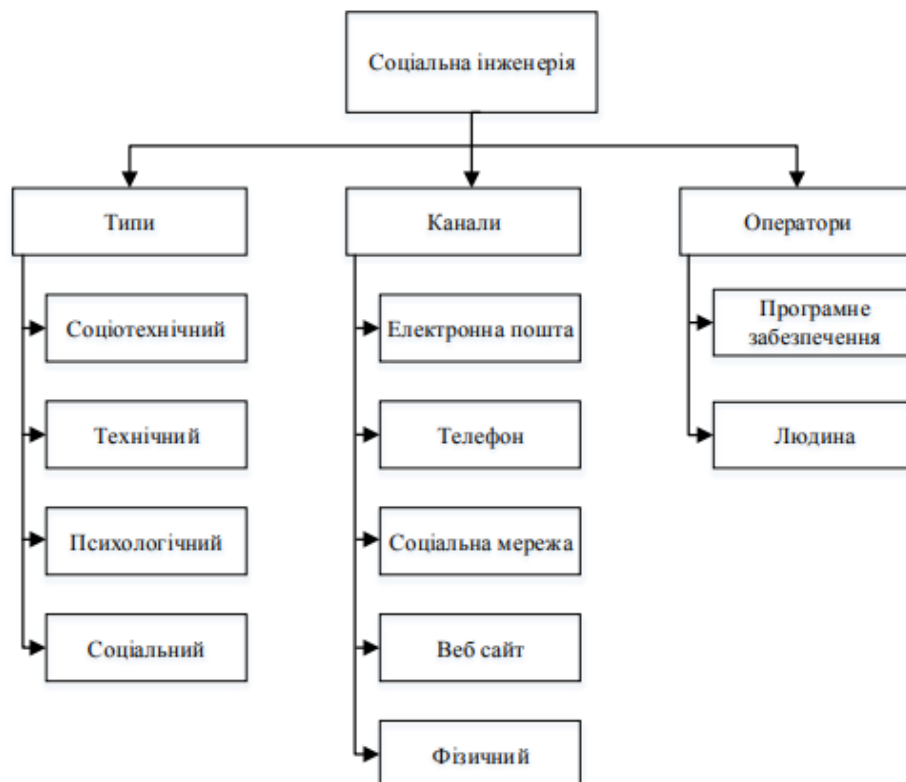


Рис. 2. Класифікація ознак реалізації атак соціальної інженерії



Вплив на жертву часто базується на її внутрішніх мотивах, серед яких – прагнення до самоствердження, бажання досягти успіху або матеріальна вигода. Соціальні інженери майстерно комбінують ці фактори, створюючи такі ситуації, у яких людина несвідомо погоджується виконати необхідні дії. Зважаючи на це, усвідомлення загроз соціальної інженерії є критично важливим для захисту від маніпулятивного впливу та збереження власної інформаційної безпеки.

Захист від атак соціальної інженерії є складним і багатограним процесом, що вимагає комплексного підходу для мінімізації ризиків і забезпечення інформаційної безпеки. Оскільки ці загрози спрямовані переважно на людський фактор, ефективний захист повинен включати як технічні, так і освітні заходи, що дозволяють знизити вразливість працівників до маніпуляцій та обману [15].

Одним із ключових методів захисту є навчання співробітників і підвищення їхньої обізнаності щодо різних типів атак соціальної інженерії. Розуміння потенційних загроз і навички їх розпізнавання є критично важливими для протидії шахрайським схемам. Регулярні тренінги з інформаційної безпеки сприяють розвитку стійкості до маніпуляцій та допомагають уникати типових помилок, таких як передавання конфіденційної інформації стороннім особам або взаємодія з підозрілими електронними листами та посиланнями. Додатково ефективним інструментом є моделювання реальних атак, наприклад, тестові фішингові кампанії, які дозволяють оцінити рівень готовності персоналу і виявити слабкі місця у знаннях та навичках співробітників [15].

Крім освітніх заходів, важливу роль відіграє впровадження чітких правил поведінки з конфіденційною інформацією в рамках політики безпеки підприємства. Ці правила повинні визначати порядок доступу до чутливих даних, використання надійних паролів і методи реагування на підозрілі контакти. Наприклад, працівники мають перевіряти особу, яка запитує важливу інформацію, та уникати її передавання через ненадійні канали зв'язку [16].

Технічні заходи також є важливим компонентом захисту від соціальної інженерії. Використання багатофакторної автентифікації значно ускладнює отримання несанкціонованого доступу до інформаційних систем навіть у випадку компрометації пароля. Ця технологія передбачає використання декількох рівнів перевірки особи, таких як введення пароля, підтвердження через відбиток пальця або введення одноразового коду, надісланого на мобільний пристрій. Такий підхід значно знижує ризик успішної атаки, навіть якщо зловмиснику вдалося заволодіти частиною облікових даних [16].

Інформаційні системи повинні бути оснащені сучасними засобами фільтрації електронної пошти, що дозволяють виявляти та блокувати фішингові повідомлення. Такі фільтри аналізують зміст листів і розпізнають характерні ознаки шахрайства, наприклад, спроби видавати себе за відомі компанії чи установи. Додатково, впровадження систем моніторингу поведінки користувачів дає змогу оперативно ідентифікувати підозрілу активність та запобігати потенційним загрозам [14].

Штучний інтелект, зокрема великі мовні моделі, може суттєво посилити захист від соціальної інженерії. Ілюстрацією ефективності використання штучного інтелекту для протидії атакам соціальної інженерії є інцидент, що стався у 2024 році в міжнародній компанії Arup. Співробітник фінансового відділу отримав запрошення на відеоконференцію, яка виглядала цілком легітимною: серед учасників були колеги, включаючи фінансового директора. Під час зустрічі йому надійшов запит на терміновий переказ значної суми – \$25 мільйонів на зазначений рахунок. Усі учасники, окрім самого працівника, виявилися цифровими двійниками, створеними за допомогою технологій deepfake – їхній зовнішній вигляд і голоси були згенеровані штучним інтелектом.



Після виявлення шахрайства компанія інтегрувала систему штучного інтелекту, що аналізує відео- та аудіозаписи на наявність ознак фальсифікації. Ця система виявляє аномалії у мові, інтонаціях, міміці та жестах, які можуть свідчити про використання deepfake-технологій. Подібні алгоритми, що працюють у режимі реального часу, дозволяють оперативно виявляти загрози та попереджати користувачів або служби безпеки про потенційні атаки.

Цей приклад яскраво демонструє практичну цінність інтеграції ШІ у систему кіберзахисту та підтверджує доцільність впровадження багаторівневого підходу до протидії сучасним соціоінженерним загрозам [17].

Вони здатні аналізувати текстову інформацію, виявляти маніпулятивні техніки та автоматично попереджати користувачів про можливі загрози. Завдяки алгоритмам машинного навчання такі системи постійно вдосконалюються та адаптуються до нових методів атак. Великі мовні моделі, такі як GPT, відіграють ключову роль у виявленні та протидії атакам соціальної інженерії, забезпечуючи новий рівень захисту завдяки своїм можливостям аналізу текстової інформації та розпізнавання маніпулятивних технік. Використовуючи передові алгоритми обробки природної мови (NLP), ці моделі здатні ідентифікувати характерні патерни шахрайських повідомлень, спроб залякування або переконання, що часто застосовуються зловмисниками для введення користувачів в оману [18].

Крім інформаційної безпеки, важливим є контроль доступу до фізичних ресурсів, таких як робочі місця, серверні приміщення або архіви. Використання електронних перепусток і систем відеоспостереження допомагає забезпечити контроль за переміщенням осіб у межах організації, знижуючи ризик інсайдерських загроз.

Формування культури інформаційної безпеки також відіграє значну роль у мінімізації ризиків. Це передбачає підтримку постійного діалогу щодо важливості захисту даних та впровадження ініціатив, які стимулюють працівників дотримуватися правил безпеки. Наприклад, можна запровадити систему заохочень за вчасне виявлення загроз або пропозиції щодо покращення заходів безпеки [19].

Сучасні технології стрімко змінюють ландшафт інформаційної безпеки, і разом із цим удосконалюються методи соціальної інженерії. Зловмисники активно використовують автоматизацію та штучний інтелект, що ускладнює виявлення та нейтралізацію атак. У відповідь на це організації повинні бути гнучкими та готовими швидко адаптувати свої системи захисту [19].

Перспективним напрямом у протидії соціальній інженерії є використання аналізу великих обсягів даних для виявлення аномалій у поведінці користувачів. Технології Big Data дозволяють відстежувати поведінкові шаблони та ідентифікувати потенційні загрози. Наприклад, якщо співробітник несподівано намагається отримати доступ до інформації, яка не належить до його компетенції, це може бути ознакою атаки [20].

Віртуальні помічники та чат-боти, що працюють на основі штучного інтелекту, можуть стати додатковим засобом захисту. Вони допомагають користувачам оперативно оцінювати підозрілі запити та пропонують оптимальні дії для запобігання загрозам. Це особливо актуально для великих компаній, де обсяг взаємодій є значним, а традиційні методи перевірки можуть бути недостатньо ефективними [20].

Ще одним важливим аспектом захисту від соціальної інженерії є використання технологій шифрування для захисту конфіденційних даних. Сучасні криптографічні протоколи забезпечують надійне кодування інформації, ускладнюючи її використання у разі витоку. Водночас важливо, щоб працівники розуміли принципи безпечного



зберігання та передавання зашифрованих даних, щоб уникнути помилок, які можуть призвести до компрометації інформації.

Важливим аспектом забезпечення інформаційної безпеки є розробка ефективних систем реагування на інциденти. План дій у разі виявлення соціальної інженерної атаки повинен бути чітко визначеним і передбачати швидке інформування всіх зацікавлених сторін, ізоляцію можливих вразливих систем, а також аналіз і усунення слабких місць у безпеці. Крім того, після кожного інциденту слід проводити детальний аналіз, щоб зрозуміти природу атаки та розробити заходи для запобігання подібним випадкам у майбутньому [21].

Регулярний аудит інформаційної безпеки відіграє ключову роль у виявленні потенційних вразливостей. Незалежні перевірки та тестування на проникнення дозволяють оцінити ефективність існуючих заходів захисту, виявити слабкі місця в політиках безпеки, контролі доступу та навчальних програмах. Оцінка рівня підготовки персоналу також є важливою складовою аудиту, оскільки дозволяє визначити, наскільки добре працівники знають методи соціальної інженерії та вміють їх розпізнавати [22].

Оскільки атаки соціальної інженерії часто базуються на емоційному впливі, наприклад, створенні почуття терміновості або страху, важливо навчати співробітників зберігати холонокровність та не піддаватися маніпуляціям. Зловмисники можуть використовувати складні психологічні техніки, щоб викликати почуття провини або співчуття, тому працівники повинні бути обізнані про такі загрози та знати, що у випадках сумнівних запитів слід негайно звертатися за консультацією до фахівців з інформаційної безпеки [22].

Таким чином, враховуючи технічні, організаційні та психологічні заходи протидії, пропонуємо узагальнений алгоритм захисту від атак соціальної інженерії, блок-схема якого зображена на рис. 3.

Дана блок-схема відображає універсальний підхід до реагування на потенційні загрози соціальної інженерії, орієнтований на працівників організацій та користувачів цифрових сервісів.

Основні етапи запропонованого алгоритму можна описати таким чином:

1. *Підозріле повідомлення або запит.* Вхідна точка алгоритму – поява повідомлення чи звернення, яке викликає сумніви або виглядає незвично.

2. *Оцінка емоційного тиску.* Перевіряється наявність психологічного впливу: терміновість, залякування, тиск авторитету чи довіри, які характерні для атак соціальної інженерії.

3. *Порівняння з внутрішніми політиками.* Зміст повідомлення повинен бути співставленим з правилами безпеки, прийнятими в організації. Наприклад, заборона передавати дані поштою або виконувати команди без підтвердження.

4. *Консультація з фахівцем з безпеки.* У разі сумнівів – обов'язкове звернення до уповноваженого спеціаліста (CISO, служба ІБ, менеджер з ризиків тощо).

5. *Прийняття рішення.* На основі аналізу можуть бути прийняті наступні рішення:

- ігнорувати запит;
- виконати дію (за погодженням);
- передати інформацію до служби безпеки.

6. *Документування інциденту та навчання.* Якщо загроза підтверджена – фіксація інциденту, внесення змін у політики/навчання, проведення тренінгів або розсилок з метою профілактики. Такий підхід дозволяє стандартизувати дії працівників, зменшує ризик людської помилки і забезпечує швидку реакцію на сучасні загрози.



Рис. 3. Загальний алгоритм захисту від атак соціальної інженерії

Одна з головних переваг великих мовних моделей полягає у здатності аналізувати великі обсяги тексту в режимі реального часу. Вони можуть сканувати електронну пошту, повідомлення в месенджерах та інші канали комунікації, виявляючи потенційні загрози на основі стилю викладу, мовних зворотів і контексту повідомлення. Наприклад, якщо текст містить ознаки терміновості чи неправдоподібні запити, модель може виявити ці ризики та попередити користувача або службу безпеки про можливу загрозу.

Завдяки здатності аналізувати контекст і взаємодію між користувачами, великі мовні моделі ефективно виявляють складні види атак, зокрема spear-phishing, коли зловмисники використовують персоналізовану інформацію для створення переконливих шахрайських повідомлень. Такі моделі здатні виявляти подібні загрози шляхом порівняння вхідного тексту з раніше зафіксованими прикладами атак і фіксації відхилень, що свідчать про потенційну маніпуляцію.

Окрім виявлення вже активних загроз, мовні моделі можуть виконувати превентивну функцію, сприяючи формуванню стратегій захисту ще до виникнення атаки. Аналізуючи історичні дані, шаблони поведінки користувачів та характерні ознаки ризику, ці системи здатні передбачати потенційні уразливості, знижуючи ймовірність успішного злому та зміцнюючи загальну кібербезпеку організацій.

Ще однією важливою функцією є здатність мовних моделей до самонавчання. Вони постійно вдосконалюють свої алгоритми, адаптуючись до нових тактик зловмисників і



підвищуючи ефективність виявлення загроз. Цей процес дозволяє їм залишатися актуальними та бути готовими до виявлення нових методів атак, які ще не були широко застосовані.

Крім того, великі мовні моделі сприяють автоматизації аналізу загроз, створюючи детальні звіти про потенційно небезпечні повідомлення та пропонуючи рекомендації щодо подальших дій. Це значно прискорює реагування на інциденти, що особливо корисно для служб підтримки та команд кібербезпеки, які обробляють великий потік запитів [23].

Ще один важливий аспект – використання мовних моделей для навчання співробітників. Вони можуть створювати симуляції соціоінженерних атак, що дозволяє працівникам тренувати навички розпізнавання загроз у реальних ситуаціях. Такі навчальні програми підвищують рівень обізнаності співробітників та зміцнюють загальну кібербезпеку організації.

Водночас, застосування великих мовних моделей у сфері інформаційної безпеки має певні виклики. Наприклад, вони можуть генерувати хибнопозитивні сповіщення, що призводить до зайвого навантаження на служби безпеки. Тому їх варто використовувати у комплексі з іншими інструментами та методами для забезпечення найбільш ефективного захисту.

Загалом, інтеграція великих мовних моделей у системи безпеки дозволяє організаціям застосовувати проактивний підхід до виявлення та запобігання атакам. Їх здатність до аналізу великих обсягів інформації, розпізнавання контекстуальних загроз і самонавчання значно підвищує ефективність кіберзахисту, допомагаючи випереджати дії зловмисників [24].

Ефективне впровадження великих мовних моделей у сферу кібербезпеки передбачає їхню гармонійну інтеграцію з уже наявними захисними механізмами та ретельне налаштування для досягнення найкращих результатів. Важливо, щоб ці моделі взаємодіяли з іншими засобами кіберзахисту, зокрема з системами виявлення загроз, антивірусними програмами та інструментами шифрування. Такий підхід дозволяє створити багаторівневу систему безпеки, що підвищує її ефективність у протидії сучасним загрозам.

Окрім цього, критично важливим є дотримання принципів конфіденційності та етичних норм при використанні великих мовних моделей. Впровадження штучного інтелекту для моніторингу інформаційних потоків і аналізу поведінки користувачів має відповідати чинному законодавству та внутрішнім регламентам компаній. Це включає забезпечення прозорості у використанні технологій та суворе дотримання вимог щодо захисту персональних даних. Для уникнення порушень прав користувачів і зміцнення довіри необхідно розробляти чіткі політики, що визначають мету та способи застосування мовних моделей [25].

Ще одним важливим аспектом є підготовка спеціалістів у галузі інформаційної безпеки до роботи з великими мовними моделями. Вони мають глибше розуміти принципи їхньої роботи, обмеження та ефективні методи застосування для виявлення загроз і нейтралізації потенційних атак. Належна підготовка фахівців сприяє ефективнішому використанню цих технологій і забезпечує швидку адаптацію до динамічного середовища кібербезпеки.

Постійний розвиток великих мовних моделей відкриває нові можливості у сфері кіберзахисту, демонструючи високий потенціал у запобіганні кібератакам. Водночас дослідження у цій галузі не можуть припинятися, оскільки зловмисники безперервно вдосконалюють методи обходу захисних систем. Очікується, що у майбутньому штучний



інтелект відіграватиме ще вагомішу роль у формуванні інтелектуальних систем безпеки, здатних оперативно реагувати на нові виклики у режимі реального часу [26].

Загалом, великі мовні моделі відкривають нові перспективи у виявленні та нейтралізації атак соціальної інженерії. Завдяки здатності до швидкої обробки значних обсягів інформації та аналізу контексту вони дозволяють оперативно ідентифікувати загрози, підвищуючи рівень захисту. Попри певні труднощі, пов'язані з їхнім впровадженням, ці моделі мають значний потенціал для вдосконалення кібербезпеки. Організації, які інвестують у їхню інтеграцію та розвиток, отримують перевагу у протистоянні сучасним загрозам і здатні ефективніше реагувати на нові виклики цифрової безпеки [27].

Завдяки можливості аналізувати текст у реальному часі, виявляти підозрілі патерни та адаптуватися до нових загроз, ці моделі забезпечують проактивний захист і підвищують ефективність реагування на інциденти. Автоматизоване створення звітів та навчання користувачів сприяють підвищенню обізнаності персоналу, що є важливим аспектом у запобіганні успішним атакам. Таким чином, використання великих мовних моделей значно зміцнює систему кібербезпеки і допомагає організаціям залишатися на крок попереду зловмисників.

Ми визначили ключові функції великих мовних моделей у боротьбі з атаками соціальної інженерії та їхні значні переваги (див. табл. 2)

Таблиця 2.

Роль великих мовних моделей у виявленні та нейтралізації атак соціальної інженерії

Функція великих мовних моделей	Опис можливостей	Результати застосування
Аналіз тексту в реальному часі	Перевірка електронних листів, повідомлень та комунікацій	Швидке виявлення потенційних загроз
Виявлення підозрілих патернів	Визначення характерних ознак обману або маніпуляцій	Зниження ризику успішних атак соціальної інженерії
Самонавчання та адаптація	Постійне вдосконалення на основі нових даних	Актуальність і здатність адаптуватися до нових методів
Автоматизоване створення звітів	Генерація детальних описів загроз та рекомендацій	Прискорення процесу реагування та інциденти
Навчання користувачів	Створення симуляцій атак для тренування співробітників	Підвищення обізнаності та загального рівня безпеки

Таким чином атаки соціальної інженерії залишаються однією з найсерйозніших загроз для інформаційної безпеки, оскільки вони орієнтовані на людський фактор і базуються на широкому спектрі психологічних маніпуляцій. Проведений аналіз свідчить, що ефективна протидія таким загрозам потребує комплексного підходу, який поєднує технологічні рішення, зокрема використання великих мовних моделей, із системним підвищенням обізнаності та навчанням користувачів. Водночас застосування мовних моделей супроводжується певними викликами – можливістю хибнопозитивних сповіщень, необхідністю контролю якості результатів і дотриманням норм



конфіденційності. Найбільшу ефективність досягається за умови інтеграції таких рішень з іншими засобами кіберзахисту та належної підготовки фахівців у сфері інформаційної безпеки.

ВИСНОВКИ

Соціальна інженерія залишається однією з найнебезпечніших загроз у сфері кібербезпеки, оскільки вона використовує людський фактор як основний вектор атаки. Незважаючи на постійне вдосконалення технологічних засобів захисту, саме психологічні маніпуляції дають змогу зловмисникам отримувати конфіденційну інформацію, змушуючи жертв самостійно передавати доступ до систем або виконувати певні дії. Аналіз сучасних методів соціальної інженерії, таких як фішинг, вішинг, смішинг, прітекстинг, спір-фішинг і вейлінг, показує, що успіх цих атак значною мірою залежить від експлуатації довіри, страху, терміновості, соціального впливу та когнітивних упереджень.

Ефективна протидія атакам соціальної інженерії потребує комплексного підходу, що включає як технічні, так і освітні заходи. Важливим аспектом є підвищення рівня цифрової грамотності користувачів, навчання персоналу розпізнаванню шахрайських схем і розробка чітких політик інформаційної безпеки. Використання багатофакторної автентифікації, систем виявлення загроз, аналізу поведінкових аномалій і технологій штучного інтелекту дозволяє значно знизити ризики компрометації даних. Зокрема, великі мовні моделі демонструють високу ефективність у розпізнаванні маніпулятивних технік та автоматизованому аналізі загроз, що сприяє оперативному реагуванню на потенційні атаки.

Незважаючи на розвиток захисних технологій, соціальні інженери постійно вдосконалюють свої методи, адаптуючи їх до нових реалій цифрового середовища. Це вимагає постійного оновлення стратегій кіберзахисту та гнучкого підходу до забезпечення інформаційної безпеки. Майбутні дослідження мають бути зосереджені на розробці інтегрованих систем захисту, що поєднують автоматизовані рішення на основі штучного інтелекту з людськими факторами, такими як критичне мислення та навички протидії маніпуляціям.

Загалом, результати дослідження підтверджують, що лише комбінація технологічних рішень та свідомої поведінки користувачів може створити ефективний захист від соціоінженерних атак. Інвестування в навчання, автоматизацію кібербезпеки та адаптивні методи аналізу загроз сприятиме значному підвищенню рівня захищеності як окремих користувачів, так і організацій загалом.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Zaoui, M., Yousra, B., Yassine, S., Yassine, M., & Karim, O. (2024). A comprehensive taxonomy of social engineering attacks and defense mechanisms: Toward effective mitigation strategies. *IEEE Access*, 12, 72224–72241. <https://doi.org/10.1109/ACCESS.2024.3403197>
2. Mokhor, V. V., Tsurkan, O. V., Herasymov, R. P., & Tsurkan, V. V. (2017). Information security assessment of computer systems by socio-engineering approach. In *Selected Papers of the XVII International Scientific and Practical Conference "Information Technologies and Security"* (pp. 92–98). Kyiv.
3. *Analysis of the cyber attack on the Ukrainian power grid*. (2019).
4. Edwards, M., Larson, R., Green, B., Rashid, A., & Baron, A. (2017). Panning for gold: Automatically analyzing online social engineering attack surfaces. *Computers & Security*, 69, 18–34. <https://doi.org/10.1016/j.cose.2017.05.003>



5. Fathollahi-Fard, M. A., Hajiaghaei-Keshteli, M., & Tavakkoli-Moghaddam, R. (2018). The social engineering optimizer (SEO). *Engineering Applications of Artificial Intelligence*, 72, 267–293. <https://doi.org/10.1016/j.engappai.2018.04.006>
6. Mouton, F., Leenen, L., & Vente, H. (2016). Social engineering attack examples, templates and scenarios. *Computers & Security*, 59, 186–209. <https://doi.org/10.1016/j.cose.2016.02.008>
7. Engbretson, P. (2013). *The basics of hacking and penetration testing*. Elsevier.
8. Heartfield, R., & Loukas, G. (2018). Detecting semantic social engineering attacks with the weakest link: Implementation and empirical evaluation of a human-as-a-security-sensor framework. *Computers & Security*, 76, 101–127. <https://doi.org/10.1016/j.cose.2018.02.010>
9. Thomas, V. (2014). *Building an information security awareness program*. Elsevier.
10. Ghafir, I., Prenosil, V., Alhejailan, A., & Hammoudeh, M. (2016). Social engineering attack strategies and defence approaches. In *Proceedings of the IEEE 4th International Conference on Future Internet of Things and Cloud* (pp. 145–149). Vienna, Austria. <https://doi.org/10.1109/FiCloud.2016.27>
11. Mitnick Security. (2022). *The top 5 most famous social engineering attacks of the last decade*. <https://www.mitnicksecurity.com/blog/the-top-5-most-famous-social-engineering-attacks-of-the-last-decade>
12. Infosec Institute. (n.d.). *The top ten most famous social engineering attacks*. <https://www.infosecinstitute.com/resources/security-awareness/the-top-ten-most-famous-social-engineering-attacks>
13. PhoenixNAP. (n.d.). *Examples of social engineering attacks*. <https://phoenixnap.com/blog/social-engineering-examples>
14. Krombholz, K., Nobel, H., Huber, M., & Weippl, E. (2014). Advanced social engineering attacks. *Journal of Information Security and Applications*, 19(3), 183–194. <https://doi.org/10.1016/j.jisa.2014.09.005>
15. Марченко, О. І. (2021). *Алгоритми обробки даних для кіберзахисту*. Вінниця: ВНТУ.
16. Литвиненко, В. А. (2023). *Використання штучного інтелекту в інформаційній безпеці*. Київ: Видавництво НАНУ.
17. World Economic Forum. (2024, February). *Deepfake scam: Employee tricked into transferring \$25 million during video call*. <https://www.weforum.org/stories/2025/02/deepfake-ai-cybercrime-arup>
18. Сидоренко, І. Г. (2021). *Психологія соціальної інженерії: механізми та захист*. Харків: Харківський університет.
19. Пашко, В. К. (2022). *Інформаційні технології: основи та застосування*. Тернопіль: ТНТУ.
20. Білий, А. С. (2020). *Захист інформації в комп'ютерних мережах*. Запоріжжя: ЗНУ.
21. Колесник, Д. О. (2021). *Кібербезпека в умовах цифрової трансформації*. Луцьк: ЛНТУ.
22. Шевченко, Р. П. (2022). *Моделювання загроз інформаційної безпеки*. Суми: СумДУ.
23. Стеценко, Н. Г. (2020). *Інформаційна культура та кібербезпека*. Херсон: ХДУ.
24. Голуб, О. В. (2023). *Інтеграція великих мовних моделей у системи захисту*. Миколаїв: МНУ.
25. Trellix. (2023). *Trellix 2024 threat predictions*. <https://www.trellix.com/about/newsroom/stories/research/trellix-2024-threat-predictions>
26. Tripathi, S. (2023). *Underground development of malicious LLMs*. <https://www.trellix.com/about/newsroom/stories/research/trellix-2024-threat-predictions>
27. Ajeeth, S. (2023). *The resurrection of script kiddies*. <https://www.trellix.com/about/newsroom/stories/research/trellix-2024-threat-predictions>

**Harasymchuk Oleh**

PhD, Associate Professor, Department of Information Security
Lviv Polytechnic National University, Lviv, Ukraine
ORCID: 0000-0002-8742-8872
oleh.i.harasymchuk@lpnu.ua

Oliarnyk Yulia

Student, Department of Information Security
Lviv Polytechnic National University, Lviv, Ukraine
ORCID: 0009-0005-1018-651X
yuliia.oliarnyk.kb.2023@lpnu.ua

Nestor Andrii

Student, Department of Information Security
Lviv Polytechnic National University, Lviv, Ukraine
ORCID: 0009-0004-1601-9598
andrii.nestor.kb.2023@lpnu.ua

Nakonechyy Tatas

Postgraduate Student, Department of Information Security
Lviv Polytechnic National University, Lviv, Ukraine
ORCID: 0009-0003-4487-9424
taras.i.nakonechnyi@lpnu.ua

PSYCHOLOGICAL METHODS OF FRAUD IN CYBERSPACE AND WAYS TO COUNTER THEM

Abstract. The article examines the methods of social engineering used by attackers to gain unauthorized access to confidential information and manipulate the behavior of victims. The main types of attacks, such as phishing, vishing, smishing, pretexting, spear-phishing and whaling, as well as their features, implementation mechanisms and methods of deceiving users, are considered. Particular attention is paid to the psychological aspects of social engineering, including the influence of fear, trust, urgency, social proof and cognitive biases on the decision-making process. Modern approaches to protection against social engineering attacks are outlined, which include a combination of technological and educational methods. Measures are proposed to increase the digital literacy of users, develop information security policies, use multi-factor authentication, user behavior analysis systems and artificial intelligence to detect threats. Particular attention is paid to the use of large language models to identify fraudulent schemes and automate cybersecurity. The results of the study indicate the need for a comprehensive approach to protection against social engineering attacks, which involves synergy between technological tools and the human factor. The proposed recommendations are aimed at minimizing risks and increasing the overall level of security in the digital environment.

Keywords: social engineering, psychological manipulation, cybersecurity, phishing, vishing, information security, artificial intelligence, protection against fraud.

REFERENCES

1. Zaoui, M., Yousra, B., Yassine, S., Yassine, M., & Karim, O. (2024). A comprehensive taxonomy of social engineering attacks and defense mechanisms: Toward effective mitigation strategies. *IEEE Access*, 12, 72224–72241. <https://doi.org/10.1109/ACCESS.2024.3403197>
2. Mokhor, V. V., Tsurkan, O. V., Herasymov, R. P., & Tsurkan, V. V. (2017). Information security assessment of computer systems by socio-engineering approach. In *Selected Papers of the XVII International Scientific and Practical Conference "Information Technologies and Security"* (pp. 92–98). Kyiv.
3. *Analysis of the cyber attack on the Ukrainian power grid.* (2019).



4. Edwards, M., Larson, R., Green, B., Rashid, A., & Baron, A. (2017). Panning for gold: Automatically analyzing online social engineering attack surfaces. *Computers & Security*, 69, 18–34. <https://doi.org/10.1016/j.cose.2017.05.003>
5. Fathollahi-Fard, M. A., Hajiaghahi-Keshteli, M., & Tavakkoli-Moghaddam, R. (2018). The social engineering optimizer (SEO). *Engineering Applications of Artificial Intelligence*, 72, 267–293. <https://doi.org/10.1016/j.engappai.2018.04.006>
6. Mouton, F., Leenen, L., & Vente, H. (2016). Social engineering attack examples, templates and scenarios. *Computers & Security*, 59, 186–209. <https://doi.org/10.1016/j.cose.2016.02.008>
7. Engebretson, P. (2013). *The basics of hacking and penetration testing*. Elsevier.
8. Heartfield, R., & Loukas, G. (2018). Detecting semantic social engineering attacks with the weakest link: Implementation and empirical evaluation of a human-as-a-security-sensor framework. *Computers & Security*, 76, 101–127. <https://doi.org/10.1016/j.cose.2018.02.010>
9. Thomas, V. (2014). *Building an information security awareness program*. Elsevier.
10. Ghafir, I., Prenosil, V., Alhejailan, A., & Hammoudeh, M. (2016). Social engineering attack strategies and defense approaches. In *Proceedings of the IEEE 4th International Conference on Future Internet of Things and Cloud* (pp. 145–149). Vienna, Austria. <https://doi.org/10.1109/FiCloud.2016.27>
11. Mitnick Security. (2022). *The top 5 most famous social engineering attacks of the last decade*. <https://www.mitnicksecurity.com/blog/the-top-5-most-famous-social-engineering-attacks-of-the-last-decade>
12. Infosec Institute. (n.d.). *The top ten most famous social engineering attacks*. <https://www.infosecinstitute.com/resources/security-awareness/the-top-ten-most-famous-social-engineering-attacks>
13. PhoenixNAP. (n.d.). *Examples of social engineering attacks*. <https://phoenixnap.com/blog/social-engineering-examples>
14. Krombholz, K., Hobel, H., Huber, M., & Weippl, E. (2014). Advanced social engineering attacks. *Journal of Information Security and Applications*, 19(3), 183–194. <https://doi.org/10.1016/j.jisa.2014.09.005>
15. Marchenko, O. I. (2021). *Data processing algorithms for cybersecurity*. Vinnytsia: VNTU.
16. Lytvynenko, V. A. (2023). *Use of artificial intelligence in information security*. Kyiv: NASU Publishing House.
17. World Economic Forum. (2024, February). *Deepfake scam: Employee tricked into transferring \$25 million during video call*. <https://www.weforum.org/stories/2025/02/deepfake-ai-cybercrime-arup>
18. Sydorenko, I. G. (2021). *Psychology of social engineering: Mechanisms and protection*. Kharkiv: Kharkiv University.
19. Pashko, V. K. (2022). *Information technologies: Fundamentals and applications*. Ternopil: TNTU.
20. Bily, A. S. (2020). *Information protection in computer networks*. Zaporizhzhia: ZNU.
21. Kolesnyk, D. O. (2021). *Cybersecurity in the context of digital transformation*. Lutsk: LNTU.
22. Shevchenko, R. P. (2022). *Modeling information security threats*. Sumy: SumDU.
23. Stetsenko, N. G. (2020). *Information culture and cybersecurity*. Kherson: KhDU.
24. Golub, O. V. (2023). *Integration of large language models into security systems*. Mykolaiv: MNU.
25. Trellix. (2023). *Trellix 2024 threat predictions*. <https://www.trellix.com/about/newsroom/stories/research/trellix-2024-threat-predictions>
26. Tripathi, S. (2023). *Underground development of malicious LLMs*. <https://www.trellix.com/about/newsroom/stories/research/trellix-2024-threat-predictions>
27. Ajeeth, S. (2023). *The resurrection of script kiddies*. <https://www.trellix.com/about/newsroom/stories/research/trellix-2024-threat-predictions>

