

Міністерство освіти і науки України
Тернопільський національний технічний університет імені Івана Пулюя
(повне найменування вищого навчального закладу)
Факультет комп'ютерно-інформаційних систем і програмної інженерії
(назва факультету)
Кафедра кібербезпеки
(повна назва кафедри)

КВАЛІФІКАЦІЙНА РОБОТА

на здобуття освітнього ступеня

магістр
(освітній рівень)
на тему: "Реалізація автоматизованого інструменту для виявлення
фішингових веб-сайтів"

Виконав: студент

Спеціальності:

125 «Кібербезпека та захист інформації»

(шифр і назва напрямку підготовки, спеціальності)

Борух О.А.

підпис

(прізвище та ініціали)

Керівник

Карпінський М.П.

підпис

(прізвище та ініціали)

Нормоконтроль

Тимощук Д. І.

підпис

(прізвище та ініціали)

Завідувач кафедри

Загородна Н.В.

підпис

(прізвище та ініціали)

Рецензент

підпис

(прізвище та ініціали)

м. Тернопіль – 2024

АНОТАЦІЯ

Реалізація автоматизованого інструменту для виявлення фішингових веб-сайтів // ОР «Магістр» // Борух Олег Андрійович // Тернопільський національний технічний університет імені Івана Пулюя, факультет комп'ютерно-інформаційних систем і програмної інженерії, кафедра кібербезпеки, група СБм-61 // Тернопіль, 2024 // С. 63, рис. – 6, табл. – 2, кресл. – 0, додат. – 2.

Ключові слова: Фішинг, Виявлення, URL, алгоритм, PhishTank, Python.

Кваліфікаційна робота присвячена розробці автоматизованого інструменту для виявлення фішингових веб-сайтів на основі аналізу URL та структури HTML-коду. У роботі проаналізовано сучасні методи і техніки ідентифікації фішингових ресурсів, розглянуто основні характеристики фішингових сайтів та їхній вплив на інформаційну безпеку. Запропонований підхід базується на використанні алгоритмів машинного навчання та аналізу веб-контенту, зокрема алгоритм Random Forest, для класифікації шкідливих URL. Вхідні дані збираються з відкритих джерел, таких як PhishTank. Реалізація виконана мовою Python із використанням відповідних бібліотек для обробки даних та моделювання. Інструмент спрямований на підвищення ефективності та точності виявлення фішингових атак.

ABSTRACT

Implementation of an Automated Tool for Detecting Phishing Websites // Master's Degree Work // Borukh Oleh Andriiovych // Ternopil Ivan Puluj National Technical University, Faculty of Computer Information Systems and Software Engineering, Department of Cybersecurity, group SBm-61 // Ternopil, 2024 // P. 63, figs. 6, tables 2, drawings - 0 , added. – 2.

Keywords: Phishing, Detecting, URL, algorithm, PhishTank, Python.

The master's thesis is devoted to the development of an automated tool for detecting phishing websites based on URL and HTML structure analysis. The study examines modern methods and techniques for identifying phishing resources, explores key characteristics of phishing sites, and their impact on information security. The proposed approach utilizes machine learning algorithms, particularly Random Forest, for malicious URL classification. Data is sourced from open platforms like PhishTank. The implementation, developed in Python, leverages libraries for data processing and modeling, aiming to enhance the efficiency and accuracy of phishing attack detection.

ЗМІСТ

ВСТУП	7
РОЗДІЛ 1 ПОНЯТТЯ ФІШИНГУ	9
1.1 Сучасний стан проблеми фішингових атак.....	9
1.2 Поняття фішингу.....	10
1.3 Принцип роботи фішингу	12
1.4 Життєвий цикл фішингових атак	14
1.5 Види фішингових атак	15
1.6 Мотивація та мета фішингових атак.....	17
1.7 Статистика фішингових атак	20
РОЗДІЛ 2 РОЗПІЗНАВАННЯ ФІШИНГОВИХ ВЕБ-САЙТІВ.....	22
2.1 Підходи до розпізнавання фішингу	22
2.2 Програмний підхід до розпізнавання фішингу.....	23
2.3 Методи розпізнавання з машинним навчанням.....	28
2.4 Огляд існуючих рішень	29
2.5 Алгоритм розпізнавання фішингового веб-сайту	31
РОЗДІЛ 3 РЕАЛІЗАЦІЯ АВТОМАТИЗОВАНОГО ІНСТРУМЕНТУ ДЛЯ ВИЯВЛЕННЯ ФІШИНГУ	34
3.1 Попередня підготовка.....	34
3.2 Реалізація програмного коду	37
3.3 Аналіз результатів.....	43
РОЗДІЛ 4 ОХОРОНА ПРАЦІ ТА БЕЗПЕКА В НАДЗВИЧАЙНИХ СИТУАЦІЯХ	46
4.1 Охорона праці.....	46
ВИСНОВКИ	52
СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ.....	54
Додаток А Публікація.....	58
Додаток Б Лістинг файлу phishing_detect.py	61

ВСТУП

Щороку створюється понад 1,5 мільйона нових веб-сайтів, серед яких чимало є фішинговими. Ці шкідливі ресурси здатні з високою точністю імітувати офіційні сторінки банків, соціальних мереж та урядових установ, що робить їх особливо небезпечними та кожного року завдають багатомільярдних збитків. Лише у 2023 році кількість виявлених фішингових сайтів зросла на 65% порівняно з попереднім роком. Традиційні методи виявлення, що базуються на статичних чорних списках, не встигають за темпами появи нових загроз, адже фішингові сторінки можуть існувати всього кілька годин, а потім зникати без сліду. У цьому контексті актуальним є використання методів машинного навчання, здатних аналізувати не лише URL-адреси, але й семантичний вміст сторінок та їхню структуру, що дозволяє виявляти загрози на випередження.

Метою дослідження є створення автоматизованого інструменту, що ефективно виявляє фішингові веб-сайти на основі сучасних алгоритмів машинного навчання з максимальною точністю та мінімальним часом обробки.

Для досягнення мети сформульовані такі завдання:

- провести аналіз сучасних типів фішингових загроз, включно з атаками типу "людина в браузері" (Man in the browser (MitB));
- дослідити ефективність алгоритмів Random Forest, XGBoost та глибоких нейронних мереж для класифікації фішингових і легітимних сайтів;
- зібрати репрезентативний набір даних на основі реальних фішингових ресурсів із PhishTank і легітимних сайтів із Alexa Top Sites;
- впровадити систему аналізу, що оцінює не лише URL-адреси, а й метадані, структуру HTML-коду;
- провести тестування інструменту та порівняти його результати з традиційними методами виявлення фішингу, зокрема чорними списками та сигнатурним аналізом.

Об'єктом дослідження даної роботи є розробка та впровадження автоматизованої системи, здатної виявляти фішингові веб-сайти на основі аналізу семантичного вмісту та структури за допомогою сучасних алгоритмів машинного навчання. Предметом дослідження кваліфікаційної роботи є порівняльний аналіз ефективності різних алгоритмів машинного навчання для поділу веб-ресурсів на фішингові та легітимні, а також розробка оптимальної моделі для забезпечення високої точності та швидкості детекції.

Наукова новизна полягає у створенні гібридного підходу до виявлення фішингових веб-сайтів шляхом інтеграції аналізу семантичного вмісту сторінки та її технічних характеристик. Досліджено ефективність алгоритмів Random Forest та XGBoost на реальних даних, а також запропоновано оптимальну конфігурацію моделі для підвищення точності виявлення.

Запропонований інструмент дозволяє виявляти фішингові веб-сайти з високою точністю в реальному часі, що є критично важливим для систем безпеки великих компаній, банківських установ та індивідуальних користувачів. Розроблена система може бути інтегрована в браузері (у вигляді плагіну), антивірусні програми або мережеві шлюзи безпеки, що значно знизить ризики витоку конфіденційних даних.

Результати дослідження апробовано на XII науково-технічній 9 конференції «Інформаційні моделі, системи та технології» у вигляді опублікованих тез (див. Додаток А).

Робота складається з пояснювальної записки та графічної частини. Пояснювальна записка складається з вступу, 4 розділів, висновків, списку використаної літератури та додатків. Обсяг роботи: пояснювальна записка – 63 арк. формату А4, графічна частина – слайдів.

РОЗДІЛ 1 ПОНЯТТЯ ФІШИНГУ

1.1 Сучасний стан проблеми фішингових атак

Фішинг є однією з найскладніших і найпоширеніших загроз у сучасному цифровому середовищі. Ця проблема вимагає аналізу численних факторів і критеріїв для її виявлення та запобігання. З розвитком технологій і підвищенням популярності інтернету, зокрема доступу до мережі через смартфони, фінансові та комерційні послуги все більше надаються онлайн. Попри значну зручність, інтернет не є абсолютно безпечною платформою для фінансових операцій.

Фішингові атаки реалізуються шляхом надсилання підроблених листів або спаму, що містять посилання на шахрайські веб-сайти. Ці сайти створюються з метою збору особистої інформації користувачів, такої як паролі, платіжні дані або інші конфіденційні відомості. Аналіз показує, що фішинг є значною загрозою для фінансових установ, оскільки щорічно завдає значних збитків. Зловмисники використовують шаблони в проведенні атак, однак їх важко розпізнати на перший погляд.

Збільшення кількості користувачів інтернету за останнє десятиліття пов'язане, зокрема, з популяризацією смартфонів, які забезпечують швидкий та зручний доступ до мережі. Сучасні інтернет-платформи використовуються бізнес-організаціями, банками та онлайн-магазинами для надання послуг, що значно спрощує фінансові транзакції, економить час і ресурси. Однак поряд з перевагами зростає ризик фішингових атак, спрямованих на крадіжку конфіденційної інформації користувачів.

Зловмисники, відомі як хакери або фішери, активно використовують електронну пошту як засіб для обману. Спам-листи часто містять підроблені пропозиції виграшу або вигідних угод, що змушують користувачів переходити за посиланнями на фальшиві сайти. Аналіз вмісту таких листів вказує на повторювані шаблони, наприклад, наявність ключових слів або граматичні

помилки. Виявлення та аналіз цих шаблонів є критично важливим для розробки інструментів захисту від фішингу.

Для ефективної боротьби з фішингом використовуються різноманітні методи збору та аналізу даних, що дозволяють ідентифікувати приховані закономірності та шаблони атак.

Обробка даних є однією з ключових технік, що включає:

- кластеризацію даних;
- класифікацію;
- дерева рішень;
- мережі Байєса.

Подані методи дозволяють обробляти великі обсяги неупорядкованих даних, виділяти ключові ознаки фішингових атак і забезпечувати своєчасне виявлення загроз. Використання автоматизованих інструментів на основі машинного навчання є перспективним напрямом для підвищення ефективності боротьби з фішингом.

1.2 Поняття фішингу

Фішинг являє собою одну з форм інтернет-шахрайства, спрямовану на отримання незаконного доступу до конфіденційної інформації користувачів. Існує безліч складних і різноманітних методів фішингових атак, серед яких техніки, що змушують користувача перейти за посиланням на підроблений веб-сайт. На сьогоднішній день не існує єдиного загальноприйнятого визначення поняття фішингу, а кількість різних трактувань лише підкреслює складність і багатогранність цього виду атак [1].

Зловмисники часто обманюють користувачів, спонукаючи їх самостійно розкривати конфіденційну інформацію, таку як номери телефонів, дані банківських карток, включаючи секретні коди, або облікові дані від електронної пошти та соціальних мереж. Для цього їм пропонують певні послуги або

можливості, які здаються привабливими. Наприклад, користувачам Instagram можуть запропонувати функцію перегляду відвідувачів їх сторінки (хоча така опція не підтримується цією платформою), а клієнтам онлайн-магазинів — товари за неймовірно низькими цінами. Згідно зі статистикою, 96% фішингових атак відбуваються через електронну пошту, 3% — через небезпечні веб-ресурси, а лише 1% — за допомогою телефонних дзвінків. Як демонструють дослідження Symantec, у 2020 році щонайменше один з 4200 електронних листів був фішинговим.

За інформацією сервісу PhishTank Company, фішинг являє собою спробу отримати особисті дані людини шляхом використання електронної пошти. Часто під час фішингових атак зловмисники перенаправляють користувачів на фальшиві веб-сайти, що імітують офіційні ресурси відомих організацій чи банків. Це створює ілюзію, що користувач перебуває на справжньому сайті, вводячи його в оману [2].

Визначення, запропоноване PhishTank, охоплює значну кількість сценаріїв фішингових атак, проте не завжди є точним. Це визначення зводить фішингові атаки виключно до крадіжки персональної інформації, хоча насправді це не завжди так. Наприклад, зловмисник може змусити жертву встановити шкідливе програмне забезпечення, яке дозволяє здійснити атаку типу Man in the Browser (MITB). У такому випадку, під час авторизації користувача в банківській системі, шкідливе ПЗ автоматично здійснює переказ коштів на рахунок злочинця без будь-якого збору персональних даних. Таким чином, визначення PhishTank не відображає всі аспекти поняття «фішинг».

Інше визначення говорить, що фішингова сторінка — це будь-яка веб-сторінка, яка незаконно заявляє про свою діяльність від імені третьої сторони, з метою ввести користувачів в оману та змусити їх виконати певні дії, які вони довірили б лише офіційному представнику цієї сторони [3]. Це визначення є ширшим, оскільки мета зловмисника тут не обмежується лише крадіжкою

особистих даних. Однак, воно все ж звужує рамки, адже передбачає, що атака завжди здійснюється від імені третьої сторони. Наприклад, програма MITB, яку користувач може завантажити через, на перший погляд, безпечне посилання (скажімо, сайт із фільмами), може записувати натискання клавіш для отримання паролів. У цьому випадку зловмисник не прикидається третьою стороною, а просто поширює шкідливе посилання.

Сервіс Anti-Phishing Working Group (APWG) визначає фішинг як кримінальну техніку, яка комбінує соціальну інженерію та технічні прийоми для крадіжки особистих даних користувачів та їх фінансової інформації. У випадках соціальної інженерії зловмисники використовують підроблені електронні листи, які імітують комунікацію від імені реальних компаній чи установ, або створюють фальшиві веб-сайти, які вводять користувачів в оману, змушуючи їх надати особисті дані, такі як логіни та паролі. Технічні методи можуть включати впровадження шкідливого ПЗ для прямого викрадення даних, перехоплення імен користувачів та паролів, або маніпуляцію мережею користувача, щоб перенаправити його на фальшиві сайти чи проксі-сервери, які контролюються зловмисником.

На основі цих визначень можна зробити висновок, що, по-перше, фішингові атаки можуть здійснюватися через будь-які електронні засоби комунікації, по-друге, мета зловмисника полягає в тому, щоб змусити жертву виконати певні дії, і, по-третє, такі дії приносять користь лише атакуючому. З огляду на це, фішинг можна визначити як комп'ютерну атаку, що використовує методи соціальної інженерії через електронні канали зв'язку, аби переконати жертву здійснити дії, які працюють на користь зловмисника.

1.3 Принцип роботи фішингу

Фальшивий веб-сайт, який майже повністю повторює вигляд справжнього, зазвичай має спеціальне поле для введення конфіденційної інформації. Коли

користувач вводить свої особисті дані та підтверджує дію, ця інформація автоматично надходить до рук зловмисника. Після цього шахрай отримує доступ до потрібних йому даних і використовує їх у власних інтересах. Перед тим як створити підроблений ресурс, зловмисник ретельно досліджує інформацію про організацію, яку він планує імітувати. Він збирає деталі на офіційному сайті цієї організації та застосовує їх для створення максимально схожого фішингового сайту.

Другий етап фішингової атаки передбачає створення підробленого електронного повідомлення. У текст листа зловмисник додає посилання на шахрайський сайт і розсилає його великій кількості користувачів. У випадках, коли атака носить цілеспрямований характер, повідомлення надсилаються лише невеликій групі людей. Шкідливе посилання поширюється не тільки через електронну пошту — з цією метою можуть використовуватися також блоги, форуми або інші онлайн-ресурси.

Після переходу на фальшивий сайт користувач стикається з підробленою формою введення даних, де його найчастіше просять зазначити свої облікові дані, наприклад, логін і пароль. Користувач вводить інформацію, підтверджує дію, і в результаті всі введені дані опиняються у зловмисника. У підсумку фішер застосовує отримані відомості для здійснення протиправних дій, які можуть включати крадіжку особистих даних, використання банківських карток для несанкціонованих покупок тощо.

Механізм фішингу продемонстрований на рисунку 1.1.

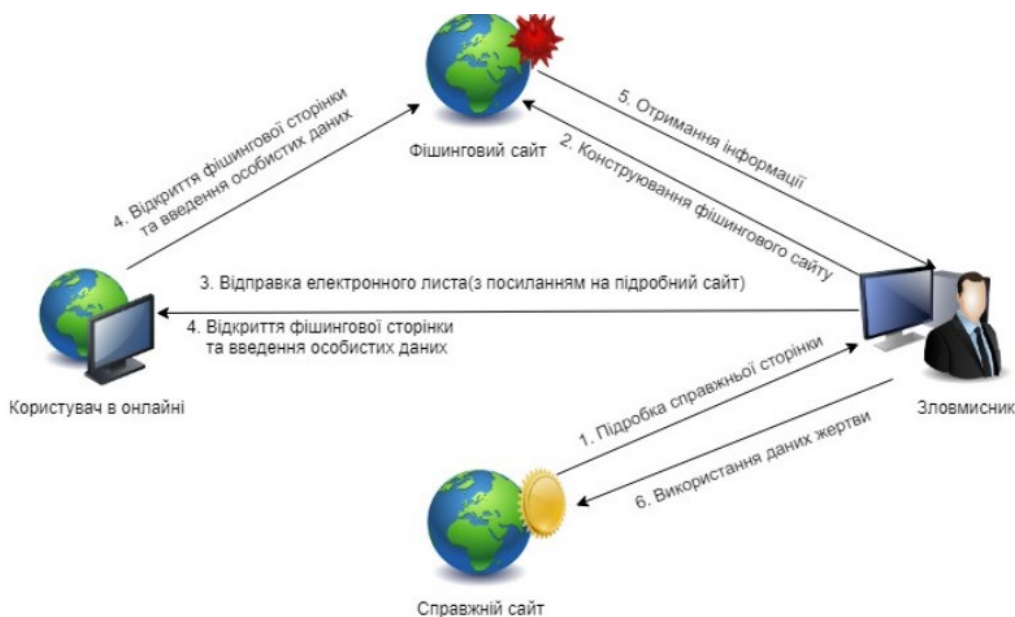


Рисунок 1.1 – Механізм фішингу

1.4 Життєвий цикл фішингових атак

Посилаючись на статтю "Phishing Website Detection Using Deep Learning Models", складається з кількох основних етапів, які детально описують механізм створення, поширення та виконання фішингових атак.

На першому етапі зловмисники створюють підроблені веб-сайти, що імітують оригінальні платформи, такі як банківські сервіси, онлайн-магазини чи соціальні мережі. Вони розробляються з урахуванням візуальної схожості з легітимними сайтами для введення в оману жертв.

На другому етапі відбувається поширення фішингових URL через різні канали, включаючи: електронну пошту (спам-листи), SMS-повідомлення (smishing), соціальні медіа, рекламні оголошення чи форуми. Метою є переконання користувача перейти за посиланням.

Коли користувач відкриває фішинговий сайт, йому пропонується ввести конфіденційні дані, такі як логіни, паролі чи платіжна інформація. Веб-сайт може також включати сценарії для автоматичного збору даних або встановлення шкідливого ПЗ на пристрій жертви.

Після успішного збору даних зловмисники використовують їх для несанкціонованих дій, наприклад для викрадення коштів з банківських рахунків, доступ до особистих акаунтів або подальше розповсюдження атак через отримані контактні дані.

Завдяки сучасним підходам, зокрема із застосуванням глибокого навчання та машинного навчання, такі атаки можуть бути виявлені шляхом аналізу поведінкових патернів, текстового контенту та характеристик URL [5].



Рисунок 1.2 – Життєвий цикл фішинг-атаки

1.5 Види фішингових атак

Існують такі види фішингових атак:

- spear phishing (цільовий фішинг);
- whaling;
- business email compromise (BEC);
- clone phishing;
- gaming;
- live chat;
- vishing (телефонний фішинг);
- pharming (зміна маршрутизації трафіку).

Вищенаведений перелік типів атак описаний згідно роботи [6].

Spear phishing спрямований на конкретного одержувача, групу людей або організацію. Щоб здійснити таку атаку, зловмисник повинен зібрати максимум інформації про свою ціль. Наприклад, це може бути інформація про заплановане відрядження або про замовлення в інтернет-магазині. Цей вид атак має високий рівень успішності через детальну підготовку.

Whaling - це атака спрямована на керівників компанії. Інформація, яку можуть надати топменеджери, є значно ціннішою, ніж дані звичайних співробітників. У цьому випадку шахрайський сайт має вигляд клієнтського ресурсу організації. Фішингові листи часто виглядають як спірні питання з клієнтами або повідомлення про офіційні проблеми, створюючи враження справжнього ділового листування. Такий вид атаки адаптується для специфічних ситуацій і може містити вигадані конфлікти чи важливі питання.

Атака Business email compromise спрямована на працівників фінансових відділів або бухгалтерії. Зловмисники видають себе за керівників компанії або інших впливових осіб, щоб переконати співробітників перевести гроші на підконтрольні їм рахунки. Атака може включати компрометацію електронної пошти генерального директора або керівника фінансів, після чого шахраї аналізують переписку, щоб краще зрозуміти внутрішні процеси компанії. В результаті співробітникам надсилають підроблений терміновий лист із проханням здійснити переказ на невідомий рахунок, який насправді належить зловмисникам.

У методі Clone phishing зловмисник копіює оригінальне повідомлення, замінюючи посилання або вкладені файли на шкідливі. Для цього використовується перехоплене повідомлення, яке клонують і надсилають від імені підробленого відправника. Іноді додається пояснення про повторну відправку або оновлену версію, щоб переконати одержувача в автентичності повідомлення.

Онлайн-ігри стали популярним засобом взаємодії, часто з використанням внутрішньо ігрових кредитів для прогресу. Зловмисники створюють фішингові сайти, які пропонують безкоштовні кредити, залучаючи гравців до введення своїх даних. За статистикою, 13% таких атак орієнтовані саме на геймерів.

Під час фішингової Live chat атаки жертва обманюється через підроблене вікно чату на веб-сайті. Наприклад, на сайтах банків зловмисники додають фальшиві вікна підтримки клієнтів. Жертва вводить свої конфіденційні дані, вважаючи, що спілкується зі службою підтримки, а дані надсилаються напяму шахраям.

Атака Vishing використовує телефон для збору даних. Жертва може отримати дзвінок із голосовим повідомленням, замаскованим під повідомлення від фінансової установи. Наприклад, просять зателефонувати за номером і ввести облікові дані або PIN-код. Насправді дзвінок спрямовується до шахрая через VoIP. Нещодавно почастишали випадки, коли зловмисники представляються працівниками технічної підтримки Apple, закликаючи користувачів телефонувати за певними номерами для вирішення вигаданих "проблем із безпекою".

Метод Pharming передбачає перенаправлення користувача з легітимного веб-сайту на підроблений. Це досягається через зміну файлу hosts на пристрої жертви або компрометацію DNS-серверів, які відповідають за перетворення доменних імен у IP-адреси. В результаті користувач потрапляє на фальшивий сайт, навіть якщо вводить правильну адресу. Цей тип атак зазвичай спрямований на домашні пристрої, а не на корпоративні сервери.

1.6 Мотивація та мета фішингових атак

Мотивація фішингових атак є різноманітною і залежить від цілей зловмисника: фінансова вигода, анонімність, бажання популярності або політичні інтереси. Незалежно від мотивів, наслідки фішингових атак завжди є

руйнівними як для індивідуальних користувачів, так і для великих організацій. Розуміння мотивів кіберзлочинців є важливим для створення ефективних методів виявлення та протидії фішинговим загрозам.

Головною метою більшості фішингових атак є крадіжка грошей або доступ до фінансових ресурсів користувачів [7]. Це включає:

- викрадення банківських реквізитів;
- злам облікових записів платіжних систем;
- продаж викрадених даних.

Для того щоб викрасти банківські реквізити зловмисники створюють сайти, що імітують банківські платформи чи системи онлайн-платежів (наприклад, PayPal, Visa, Mastercard), для отримання доступу до рахунків жертв.

Через викрадені логіни та паролі шахраї отримують доступ до платіжних акаунтів користувачів та перераховують кошти на свої рахунки або криптогаманці.

У деяких випадках викрадені дані перепродаються на "темному ринку" (Dark Web) за значну суму. Наприклад, викрадені номери кредитних карток або банківські логіни мають попит серед інших кіберзлочинців.

Таким чином, фінансовий мотив є найбільш поширеним серед зловмисників і завдає значних збитків як окремим користувачам, так і великим організаціям.

Деякі кіберзлочинці використовують фішинг для приховування своєї особистості під час виконання злочинних дій. Це стосується наступних випадків.

Зловмисники користуються викраденими обліковими записами для здійснення покупок, щоб їх не можна було ідентифікувати. Це особливо поширено при купівлі цифрових товарів або через онлайн-магазини.

Фішингові атаки дозволяють злочинцям використовувати акаунти інших людей для участі у нелегальних діях, таких як кібербулінг, поширення шкідливого програмного забезпечення чи обман інших користувачів.

Використання чужих акаунтів для переказів чи обміну валюти допомагає уникнути відстеження злочинної діяльності правоохоронними органами.

Приховування особистості важливе для зловмисників, які прагнуть уникнути покарання, виконуючи свої дії через скомпрометовані акаунти жертв.

Деякі фішингові атаки мотивовані бажанням досягти визнання у кіберспільноті або продемонструвати свої технічні навички. Це особливо поширено серед молодих зловмисників або "новачків", які хочуть отримати репутацію серед хакерів, публічне визнання або психологічне задоволення.

У кіберспільнотах існує своєрідний змагальний дух, де успішні атаки на великі компанії чи відомі сервіси підвищують статус кіберзлочинця. Зловмисники можуть оприлюднювати викрадену інформацію або публікувати звіти про свої атаки як "докази" їхньої успішності. Для деяких кіберзлочинців проведення успішної фішингової атаки стає джерелом самоствердження або розваги.

Зазвичай такі атаки не завжди мають на меті фінансову вигоду, але вони завдають значної шкоди жертвам, оскільки можуть включати викрадення приватних листувань, фотографій чи іншої конфіденційної інформації.

Окрім фінансових та особистих мотивів, політичні або соціальні мотиви є важливими у сучасному контексті. Це включає:

- кібер-шпигунство;
- соціальна інженерія;
- акти помсти.

Державні чи приватні хакери здійснюють фішингові атаки для отримання конфіденційної інформації від урядових чи корпоративних структур. Фішинг використовується для маніпулювання суспільною думкою шляхом викрадення або поширення фейкової інформації. Деякі фішингові атаки можуть бути мотивовані бажанням завдати шкоди конкретній людині чи організації через особисті образи або конфлікти.

1.7 Статистика фішингових атак

Фішингові атаки залишаються однією з найбільш поширених і загрозливих форм кіберзлочинності. Згідно з останнім звітом Anti-Phishing Working Group (APWG) за третій квартал 2024 року, кількість фішингових інцидентів продовжує зростати, а зловмисники застосовують нові, більш агресивні тактики для досягнення своїх цілей [4].

У третьому кварталі 2023 року кількість зафіксованих фішингових сайтів досягла рекордних рівнів. Фішингова активність зросла в порівнянні з попередніми періодами, що свідчить про стійку тенденцію до збільшення атак.

Злочинці дедалі частіше використовують персональні дані жертв, такі як домашні адреси, фотографії та інші персональні відомості. Це дозволяє створювати страхітливі й маніпулятивні сценарії, які підвищують ймовірність успішного обману користувачів.

Основними цілями фішингових атак залишаються фінансові установи, технологічні компанії та платформи електронної комерції. У звіті APWG зазначається, що 60% атак у третьому кварталі були спрямовані на фінансовий сектор.

У 2023 році особливо популярним стало використання шкідливих вкладень у електронних листах та фішингових сторінок, замаскованих під офіційні ресурси. Крім того, зловмисники розширюють використання SMS-фішингу (smishing), що дозволяє їм ефективніше охоплювати мобільних користувачів.

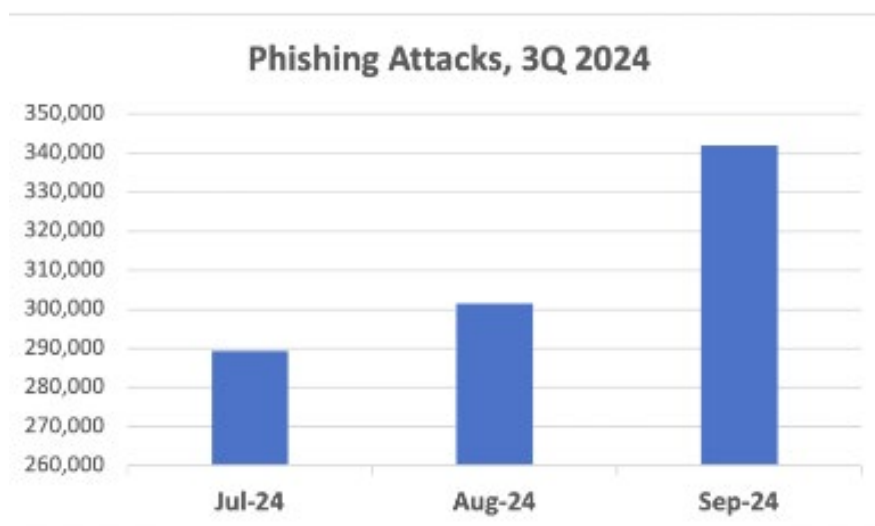


Рисунок 1.3 – Кількість фішингових атак за 3 квартал 2024 року

Найбільше фішингових атак було зафіксовано в регіонах із високою концентрацією користувачів Інтернету, таких як Північна Америка та Європа. Проте зростає кількість інцидентів у країнах Азії та Латинської Америки.

Напади на криптовалютні платформи й гаманці продовжують зростати. Кіберзлочинці використовують фейкові криптовалютні сайти та додатки для обману інвесторів і крадіжки цифрових активів.

Загалом, дані APWG підтверджують, що фішинг стає більш складним та персоналізованим, а обсяги атак зростають у глобальних масштабах.

РОЗДІЛ 2 РОЗПІЗНАВАННЯ ФІШИНГОВИХ ВЕБ-САЙТІВ

2.1 Підходи до розпізнавання фішингу

Існують два основні методи розпізнавання фішингових атак: навчання користувачів та класифікація за допомогою програмних інструментів [5].

Навчання користувачів передбачає освітню діяльність, спрямовану на підвищення обізнаності користувачів про природу фішингових атак. У результаті навчання користувачі здатні самостійно відрізнити фішингові сайти від справжніх. На відміну від запобіжних заходів, про які йдеться у [8], цей підхід акцентує увагу на кінцевому результаті — самостійному виявленні фішингу користувачами. Однак його основний недолік полягає у значних витратах та неможливості застосування у випадках, коли кількість користувачів є надто великою (наприклад, у системах на зразок PayPal, eBay чи Amazon).

Класифікація за допомогою програмного забезпечення спрямована на автоматичне й більш точне визначення фішингових атак. Його головна перевага полягає у мінімізації впливу людських помилок або недостатньої обізнаності користувачів. Класифікація на основі програмного забезпечення є економічно вигіднішою та ефективнішою, оскільки не вимагає масштабного навчання користувачів.

Як саме можна покращити точність розпізнавання фішингових атак? Для користувацького підходу точність можна підвищити за допомогою накопичення досвіду в процесі регулярного використання інтернету або завдяки спеціальним навчальним програмам. Якщо мова йде про програмну класифікацію, вдосконалення досягається через тренування моделей машинного навчання або шляхом оптимізації правил виявлення в системах на основі евристик.

Методи розпізнавання відіграють ключову роль у захисті користувачів від фішингових атак. Вони не лише забезпечують безпосередній захист, але й дозволяють аналізувати фішингові приманки для ефективного відокремлення

фішингового контенту від справжнього. При цьому слід розуміти, що жоден із методів, розглянутих у попередньому розділі, не може спрацювати, якщо фішингова атака не була виявлена.

На рисунку 2.1 представлена схема, що ілюструє підходи до розпізнавання фішингових атак.

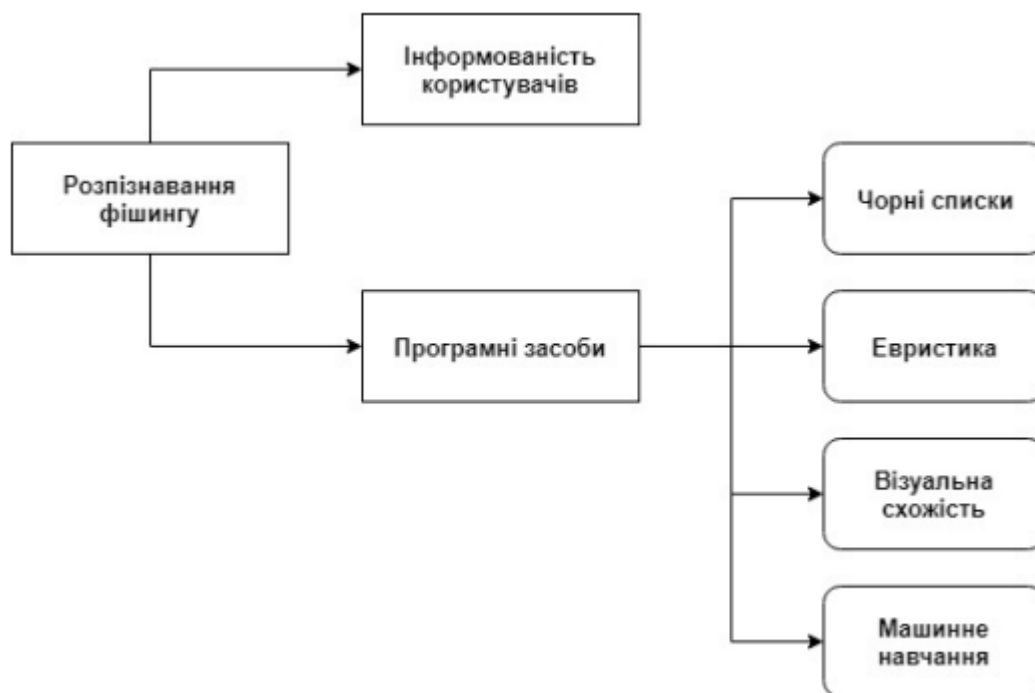


Рисунок 2.1 – Підходи до розпізнавання фішингової атаки

В даній роботі розглядається саме програмний підхід до виявлення фішингу на основі машинного навчання.

2.2 Програмний підхід до розпізнавання фішингу

Існують такі програмні методи для розпізнавання фішингу:

- методи, що базуються на чорних списках;
- аналіз характеристик сайту (методи на основі евристик);
- методи на основі машинного навчання та штучного інтелекту;
- візуальна схожість;
- гібридні методи.

Суть методу, що базуються на чорних списках, полягає в тому, що чорні списки (blacklists) — це бази даних URL-адрес, які вже були ідентифіковані як фішингові. Користувачі чи системи порівнюють URL-адресу веб-сайту з базою даних, і якщо збіг знайдено, сайт блокується. В даного методу присутні такі особливості : хоч він і є простий у реалізації та має низькі обчислювальні витрати, даний метод не виявляє нових фішингових сайтів (zero-day атаки), оскільки вони ще не додані до чорного списку.

Відомі інструменти, такі як Google Safe Browsing та PhishTank, використовують чорні списки для попередження користувачів про потенційно небезпечні сайти.

Суть методу на основі евристик в тому, що цей підхід базується на аналізі певних характеристик веб-сайту для виявлення фішингу.

До таких характеристик належать:

- структура URL-адреси (наявність IP-адреси замість доменного імені, підозрілі символи, занадто довгі URL);
- аналіз SSL-сертифікатів (відсутність сертифіката або самопідписані сертифікати);
- наявність логотипів відомих брендів та підробленого вмісту;
- аналіз DOM веб-сторінки для виявлення шкідливого коду.

Особливості в даного методу такі:

- можливість виявлення фішингових сайтів, які не потрапили до чорних списків;
- необхідність створення складних евристичних правил, які можуть дати помилкові спрацьовування (false positives).

Деякі рішення використовують DOM-based аналіз, інші фокусуються на обчисленні "схожості" сторінки зі справжніми веб-сайтами за допомогою кластеризації контенту.

Методи на основі машинного навчання та штучного інтелекту використовують алгоритми машинного навчання (Machine Learning) та штучного інтелекту (Artificial Intelligence) для ідентифікації фішингових сайтів. Вони навчаються на великих обсягах даних, які містять як фішингові, так і легітимні сайти.

Підходи у машинному навчанні:

- класифікація URL-адрес (використовуються алгоритми, такі як Random Forest, Support Vector Machines (SVM), та Logistic Regression, для аналізу структурних особливостей URL).

- аналіз вмісту веб-сайту (моделі глибокого навчання, такі як Convolutional Neural Networks (CNN) та Recurrent Neural Networks (RNN), аналізують вміст сторінок, включаючи HTML-код, зображення та тексти).

- використання обробки природної мови (NLP) (техніки NLP дозволяють аналізувати тексти електронних листів та вміст веб-сторінок для виявлення ключових слів, характерних для фішингових атак).

- об'єднані методи (поєднання ознак URL, вмісту сайту та поведінкових характеристик користувачів дозволяє підвищити точність детекції фішингових атак).

Переваги даного методу полягають у високій точності і здатності до виявлення невідомих фішингових сайтів. Недоліками є високі обчислювальні витрати та необхідність наявності великих наборів даних для навчання моделей. Прикладами таких моделей можуть бути: PhishNet - модель, що використовує нейронні мережі для класифікації веб-сайтів, та URLNet - аналізує структуру URL-адрес за допомогою глибокого навчання.

Метод розпізнавання фішингових сайтів базується на аналізі їхнього зовнішнього вигляду, а не на дослідженні вихідного коду чи мережевих даних. Такі сайти часто імітують оригінальні, щоб ввести користувача в оману. У цьому підході рішення приймається на основі набору характеристик, таких як

текстовий контент, форматування, використані HTML-теги, CSS-стили, зображення тощо. Порівнюються підозрілі сайти з перевіреними офіційними версіями за різними параметрами, і якщо рівень схожості перевищує певний поріг, сайт вважається фішинговим.

Щоб обійти ці методи виявлення, зловмисники часто використовують графічні зображення, Flash, ActiveX або Java-аплети замість тексту в HTML-кодi. Технології, які ґрунтуються на візуальній схожості, здатні швидко розпізнавати такі вбудовані елементи, що характерні для фішингових сторінок. Крім того, подібні підходи застосовують так звані "підписи" для ідентифікації підроблених ресурсів. Підпис формується на основі загальних характеристик сайту в цілому, а не лише однієї його сторінки. Завдяки цьому один підпис може бути використаний для виявлення різних сторінок або версій одного й того ж веб-сайту.

Гібридні методи поєднують переваги кількох підходів, щоб підвищити ефективність виявлення фішингу. Наприклад, вони можуть використовувати чорні списки як першу лінію захисту та застосовувати машинне навчання для аналізу нових загроз.

Особливостями гібридних методів є висока точність за рахунок комбінування різних підходів та здатність адаптуватися до нових шаблонів атак.

Таблиця 2.1 резюмує відомі утиліти для виявлення фішингу [9].

Таблиця 2.1 – Поширені засоби виявлення фішингу

Утиліта	Опис	Переваги	Недоліки
GoldFish	Метод фільтрації, заснований на байєсівській моделі	Може бути адаптований для кожного окремого користувача.	Вразливий до техніки "байєсівського отруєння". Легко обходиться шляхом невеликих змін у тексті.

Продовження таблиці 2.1

Утиліта	Опис	Переваги	Недоліки
eBay tool	Використання чорних списків і евристичних підходів для блокування шахрайських сайтів.	Захищає користувачів eBay.	Власиві недоліки, характерні для чорних списків.
Браузери Site Adviser Netcraft	Чорні списки для блокування підозрілих сайтів	Забезпечують швидке виявлення загроз.	Повільний процес оновлення списків. Можливі помилкові спрацювання (FP).
SpoofGuard, PwdHash	Інтеграція з браузерами для аналізу та виявлення ознак фішингу (наприклад, заплутаних URL-адрес). Підтримує методи евристики.	PwdHash захищає паролі від крадіжки. SpoofGuard запобігає доступу з несанкціонованих IP і MAC-адрес.	Деякі функції можуть бути ненадійними. Перехоплені паролі можуть використовуватися на інших ресурсах.
EarthLink toolbar	Поєднання евристичних методів, користувацької оцінки та ручної перевірки для аналізу сайтів.	Перевіряє дані про реєстрацію доменів.	Не забезпечує захисту від атак нульового дня.

2.3 Методи розпізнавання з машинним навчанням

Цей підхід передбачає застосування автоматизованих класифікаторів, заснованих на алгоритмах машинного навчання та методах інтелектуального аналізу даних. Подібні класифікатори функціонують на серверній стороні, здійснюючи визначення класу сторінки на основі аналізу її характеристик. У таблиці 2.2 наведено порівняння основних параметрів автоматизованих методів розпізнавання [10].

Таблиця 2.2 – Порівняльна характеристика класифікаторів

Алгоритм	Основний принцип	Потрібність навчання	Переваги	Недоліки
Naïve Bayes	Кожен параметр розглядається незалежно від інших параметрів класу. Заснований на теоремі Байєса для передбачення класу з ймовірностей.	Цей підхід вимагає навчального набору, оскільки побудова моделі базується на розмічених даних.	Простий алгоритм з мінімальними обчисленнями. Ефективний для великих наборів даних.	Покладається на припущення незалежності, що може призвести до помилок, якщо це припущення порушується.
LR (Logistic Regression)	Використовує лінійне рівняння для оцінки шансів настання події. Аналізує зв'язки між змінними для прогнозування результату.	Вимагає навчання на основі даних, щоб коректно створити модель.	Легко інтерпретувати результати. Ефективний для задач класифікації з двома варіантами.	Потребує додаткових припущень для точного застосування. Чутливий до неповноти даних.

Продовження таблиці 2.2

Алгоритм	Основний принцип	Потрібність навчання	Переваги	Недоліки
SMO (Sequential Minimal Optimization)	Використовує методи SVM для класифікації, будуючи гіперплощини для розділення класів.	Вимагає навчання з підготовкою даних.	Вирішує задачі квадратичного програмування	Необхідно правильно вибрати ядро. Складно інтерпретувати результати.
J48	Створює дерево рішень для класифікації даних. Запитує важливість ознак на кожному етапі. Заснований на алгоритмі Quinlan C4.5.	Вимагає розміченого набору для навчання моделі.	Простота використання і зрозумілість результатів. Підходить для швидкого аналізу.	Має тенденцію до перенавчання. Складно працює з діагональними межами рішень.
kNN (k-Nearest Neighbors)	Метод базується на класифікації, враховуючи класи найближчих сусідів.	Не потребує навчання, але потребує розміченого набору для точного функціонування.	Легко реалізується та зрозумілий у застосуванні. Може забезпечувати високі результати за правильного вибору параметрів.	Схильний до високих витрат ресурсів. Чутливий до зашумлених даних, що впливає на результати класифікації.

2.4 Огляд існуючих рішень

У праці [12] представлено метод детектування фішингових веб-сайтів на основі аналізу доменної інформації. Однією з особливостей цього підходу є

включення даних про прямі й непрямі посилання на сайт, що раніше рідко враховувалося іншими антифішинговими техніками. Метод включав витягнення вихідного коду сторінки, поділ посилань на два набори (S1 – прямі, S2 – непрямі), їх об'єднання та подальшу обробку за допомогою DNS-пошуку для отримання IP-адрес доменів. Висока точність результатів, що досягла 99.62%, демонструє ефективність цього підходу. Проте, залежність від зовнішніх інструментів (DNS-запитів та пошукових систем) може знижувати стабільність роботи методу.

У дослідженні [13] було запропоновано використання асоціативної класифікації як техніки для аналізу та розпізнавання фішингових сторінок. Зібравши вибірку з 601 легітимного та 752 фішингових сайтів, автори ідентифікували 16 ключових характеристик, які використовувалися для класифікації. Зокрема, 20.5% фішингових ресурсів містять IP-адресу в URL, що стало однією з важливих ознак для розпізнавання.

У роботі [14] порівнюються методи відбору характеристик для оптимізації класифікаційних моделей. Використовуючи алгоритми C4.5 (дерево рішень) і IREP (індукція правил), а також інструменти симетричну невизначеність та Information Gain, було визначено 11 основних характеристик для виявлення фішингових сторінок. Результати показали, що SSL-сертифікат і URL-якорі є найбільш вагомими ознаками. Незважаючи на зменшення кількості характеристик, точність моделей знизилася лише на 1.02% для C4.5 та 1.23% для IREP.

У дослідженні [15] для оптимізації характеристик було застосовано Wrapper-based Feature Selection (WBFS), що забезпечує високу адаптивність під конкретний класифікатор. Порівнюючи цей метод із Information Gain (IG) і Principal Component Analysis (PCA), автори визначили, що алгоритми kNN та Random Forest досягли точності 97.1% та 97.3% відповідно.

Дослідження [16] зосереджувалося на порівнянні чотирьох алгоритмів класифікації – Bayesian Network, NNet, SVM і Decision Tree. Серед цих методів

нейронні мережі (NNet) показали найкращі результати у точності та чутливості. Аналогічне дослідження [17] продемонструвало високу ефективність алгоритмів Random Forest і kNN, які досягли точності понад 97%.

У роботі [18] запропоновано покращений алгоритм класифікації на основі вдосконаленої версії PRISM. Метод продемонстрував мінімальні показники хибнопозитивних (FP) і хибнонегативних (FN) результатів – близько 0.1%.

Дослідження [19] базувалося на тестуванні шести алгоритмів машинного навчання – Naive Bayes, J48, SVM, Logistic Regression, Neural Network та IBK lazy classifier. Точність моделей варіювалася від 92.7% до 96.57%, при цьому найкращі результати продемонстрували SVM і Logistic Regression. У роботі було використано співвідношення 70/30 для навчального та тестового набору даних.

Нарешті, в роботі [20] запропоновано антифішингову систему у вигляді плагіна для веб-браузера, яка інтегрує алгоритми Random Forest, kNN та Bayesian Classifier (BC). Аналіз вибірки з 7690 легітимних та 2280 фішингових сайтів показав, що Bayesian Classifier забезпечує високу швидкість та точність детекції фішингових сторінок.

2.5 Алгоритм розпізнавання фішингового веб-сайту

Алгоритм розпізнавання фішингового веб-сайту складається з кількох етапів, які базуються на аналізі URL-адрес, структури сторінки, вмісту HTML, доменної інформації та поведінкових характеристик.

Детальна структура алгоритму включає наступні кроки:

- збір вхідних даних;
- попередній аналіз URL;
- аналіз HTML-коду та DOM-структури;
- аналіз доменної інформації;
- поведінковий аналіз;
- машинне навчання та класифікація.

Перший етап — отримання необхідної інформації для подальшого аналізу: URL-адреса веб-сайту (основний об'єкт аналізу), HTML-код сторінки (вихідний код сторінки, що містить метадані, теги, скрипти та інший контент), доменна інформація (дані про домен (WHOIS-інформація, вік домену, DNS-записи)), пов'язані посилання (прямі та непрямі посилання, що ведуть зі сторінки), користувацька активність (поведінкові дані, зокрема час перебування на сторінці, взаємодія з елементами).

Попередньо аналізуються різні ознаки URL-адреси для виявлення аномалій:

- фішингові сайти часто використовують довгі URL;
- наявність символів @, -, ?, = може вказувати на фішинг;
- використання числових IP-адрес замість легітимного доменного імені є ознакою шахрайства;
- велика кількість піддоменів може бути ознакою фішингової структури (наприклад, login.secure.bank.example.com);
- відсутність SSL-сертифіката вказує на потенційно небезпечний сайт.

Формальні правила, що базуються на регулярних виразах та евристиках, є основою цього кроку.

Далі відбувається обробка вихідного коду сторінки:

- короткі або занадто загальні заголовки можуть вказувати на підроблений сайт;
- велика кількість форм <form> для введення даних (наприклад, логіни та паролі);
- пошук підозрілих форм із відправкою даних на сторонні сервери;
- JavaScript та скрипти (виявлення шкідливих скриптів, таких як eval() або document.location.href, які можуть перенаправляти користувача на інші ресурси);
- фішингові шаблони (перевірка наявності зображень, іконок і логотипів, які копіюють дизайн відомих брендів).

Етап аналізу доменної інформації дозволяє отримати важливі характеристики домену, такі як: вік домену (фішингові сайти часто використовують новостворені домени), WHOIS-інформація (відсутність даних про власника або анонімізація може свідчити про фішинг), DNS-записи (аналіз IP-адрес і серверів, з якими пов'язаний домен), прямі та непрямі посилання (витягуються посилання на інші ресурси, аналізується їх походження та репутація).

На етапі поведінкового аналізу проводиться оцінка дій користувача на сторінці: запити даних (перевірка, чи сайт вимагає введення конфіденційної інформації (логіни, паролі, номери карток)), перенаправлення (аналіз кількості та типу перенаправлень на інші URL), взаємодія з користувачем (перевірка наявності підозрілих pop-up вікон та інших нав'язливих елементів).

На основі зібраних характеристик будується модель машинного навчання:

- всі ознаки (довжина URL, кількість скриптів, вік домену тощо) об'єднуються у вектор;

- Random Forest (ефективний для роботи з табличними даними та ознаками URL), Support Vector Machine (SVM) (дозволяє будувати точний класифікатор) або нейронні мережі (ANN) (підходять для глибокого аналізу великих обсягів даних);

- на основі датасетів із фішинговими та легітимними веб-сайтами модель навчається класифікувати нові URL;

- модель оцінюється за метриками Accuracy, Precision, Recall та F1-score.

На фінальному етапі модель повертає результат у вигляді «1» (фішинговий веб-сайт) або «0» (легітимний веб-сайт). Додатково система може надавати детальні звіти про основні ознаки, що призвели до класифікації, рівень ризику та рекомендації для користувача.

РОЗДІЛ 3 РЕАЛІЗАЦІЯ АВТОМАТИЗОВАНОГО ІНСТРУМЕНТУ ДЛЯ ВИЯВЛЕННЯ ФІШИНГУ

3.1 Попередня підготовка

На етапі попередньої підготовки важливим кроком є створення якісного датасету для навчання алгоритму Random Forest, який буде використовуватися для класифікації веб-сайтів як фішингових або легітимних. Для цього було проведено збір, фільтрацію та підготовку даних, які забезпечують правильне функціонування та високу точність запропонованого інструменту.

Дані для тренування алгоритму були отримані з кількох джерел:

- PhishTank (відкритий сервіс для моніторингу та звітності про фішингові URL-адреси);
- Alexa Top Sites (рейтинг популярних веб-сайтів, які класифікуються як легітимні);
- інші відкриті джерела (додаткові легітимні сайти, які були вибрані з освітніх, державних або корпоративних платформ).

PhishTank надає базу перевірених шкідливих сайтів, які регулярно оновлюються. Дані були отримані як у вигляді готового списку, так і через API, що дозволяє автоматично збирати нові фішингові сайти. Alexa Top Sites використовувався для формування списку звичайних сайтів, які мають високу довіру з боку користувачів.

Дані були збережені у файлі з розширенням .csv, що забезпечує зручність роботи з ними за допомогою мови програмування Python та бібліотек для аналізу даних, таких як pandas.

Кожний рядок у файлі містить два значення:

- URL-адресу веб-сайту;
- мітку класу.

Для уникнення дисбалансу між класами було зібрано приблизно однакову кількість прикладів фішингових та легітимних сайтів. Це дозволило уникнути ситуації, коли модель може схилитися до прогнозування частішого класу, що є типовою проблемою для незбалансованих датасетів.

Для забезпечення якості даних було виконано наступні кроки.

За допомогою бібліотеки `requests` у Python кожен сайт перевірявся на доступність. Сайти, які не відповідали на запит (HTTP-коди помилок 404, 403, 500), були виключені. Дублікати URL-адрес у датасеті були видалені. Адреси, які не відповідали формату стандартного веб-ресурсу (наприклад, ті, що містили незрозумілі символи або були неповними), також були виключені.

Після підготовки даних із кожного URL-адреси витягувалися специфічні характеристики (ознаки), які пізніше використовувалися для тренування алгоритму.

Фішингові сайти часто використовують довші адреси для приховування шкідливого вмісту. Кількість спеціальних символів, зокрема, символи «-», «@», «?» можуть бути індикаторами фішингових сайтів. Якщо URL-адреса містить IP-адресу замість доменного імені, це може свідчити про її підозрілість.

Витягувалися характеристики HTML-структури сайту, зокрема:

- кількість тегів `<form>` (фішингові сайти часто використовують форми для збору особистих даних);
- кількість тегів `<script>` та наявність підозрілих скриптів із викликами функцій `eval()` або `escape()`;
- наявність метаданих, таких як ключові слова або заголовки сайту.

На основі зібраних і очищених даних був сформований файл, що забезпечує основу для навчання та тестування моделі. Для кожного URL-адреси були записані всі витягнуті характеристики, що дозволяє використовувати їх як вхідні дані для алгоритму Random Forest.

На рисунку 3.1 наведено приклад вигляду файлу .csv, який включає URL-адресу та відповідну мітку.

У процесі роботи було зібрано близько 100 URL-адрес, що звісно недостатньо для повноцінного тренування алгоритму. Оскільки для отримання датасету більшої величини необхідна була реєстрація на веб-сайті Phishtank [3], на момент виконання даної роботи реєстрація нових користувачів була неможлива. Все ж таки, зібраний «вручну» датасет дещо зміг забезпечити певний обсяг даних для навчання алгоритму.

```
1 url,label
2 https://uspsrlef.top/,1
3 https://uspsrleo.top/,1
4 https://uspssadeowt.cfd/,1
5 https://uspssadfeza.cfd/,1
6 https://www.ilovepdf.com/word_to_pdf,0
7 https://www.google.com/,0
8 https://phishtank.org/,0
9 https://uspssadiuxl.cfd/,1
10 https://uspssadndcm.cfd/,1
11 https://www.youtube.com/,0
12 https://github.com/,0
13 https://t.co/y9IP4ggzvB,1
14 http://fysmganhfa.duckdns.org/ja/main,1
15 https://usps.com-parcelmjvhv.vip/i/,1
16 https://user-status-alert.vercel.app/,1
17 https://usps.com-parceltdhga.vip/i/,1
18 https://consultezorange-msmvocaux.weebly.com/,1
19 https://usps.com-parceltdhgb.vip/i/,1
20 https://usps.com-parcelmjhvd.vip/i/,1
21 https://informed.deliverywqk.top/l/,1
22 https://xzl.cyu.mybluehost.me/,1
23 https://assetssupport.pages.dev,1
24 https://trdsoln.sbs/purduefed,1
25 https://aibaibgsjsw5001.xyz,1
26 https://n8h4x7m9.com/,1
27 https://onally.online/Camden,1
28 https://www.onlinepenztargeprendeles.hu/kals,1
```

Рисунок 3.1 – Вигляд файлу phishing_data.csv

3.2 Реалізація програмного коду

Для реалізація програмного коду, який відповідає темі роботи, я вибрав мову програмування Python версії 3.12. В якості середовища розробки вибрав програмне середовище від компанії JetBrains PyCharm версії 2023.3.4.

Для початку створив новий проект MProject, додав до проекту попередньо створений датасет URL-адрес веб-сайтів та їхніх міток та створив новий файл з розширенням «.py» з назвою «main».

Наступним кроком я підключаю модулі, які необхідні для коректної роботи програмного коду.

Лістинг 3.1 – Імпорт модулів

```
import requests
from bs4 import BeautifulSoup
from sklearn.ensemble import RandomForestClassifier
from sklearn.model_selection import train_test_split
from sklearn.metrics import accuracy_score,
classification_report
import pandas as pd
import numpy as np
import re
```

Призначення вище перерахованих модулів:

- requests (для завантаження HTML-коду веб-сторінок);
- BeautifulSoup (для парсингу HTML-коду сторінок);
- RandomForestClassifier (алгоритм машинного навчання);
- train_test_split (розділення датасету на тренувальну та тестову частини);
- accuracy_score, classification_report (метрики оцінки моделі);
- pandas (для роботи з табличними даними);
- numpy (для роботи з масивами);
- re (модуль для роботи з регулярними виразами).

Оскільки більшості з поданих модулів не існує в стандартній бібліотеці середовища програмування, їх необхідно інстальовати. Інсталювання всіх необхідних модулів відбувається виконанням команди, яка подана нижче:

```
pip install requests beautifulsoup4 lxml scikit-learn pandas
numpy.
```

Далі в коді описана функція `parse_html`.

Лістинг 3.2 - Опис функції `parse_html`

```
def parse_html(url):
    try:
        response = requests.get(url, timeout=5)
        soup = BeautifulSoup(response.content, 'lxml')

        # Метадані
        title = soup.title.string if soup.title else ''
        meta_keywords = soup.find("meta", attrs={"name":
"keywords"})
        keywords = meta_keywords["content"] if meta_keywords else
''

        # Кількість тегів <form> (часто використовуються для
фішингу)
        form_count = len(soup.find_all('form'))

        # Пошук підозрілих JavaScript
        script_tags = len(soup.find_all('script'))
        suspicious_scripts = sum(1 for s in
soup.find_all('script') if "eval" in str(s) or "escape" in str(s))

        return [len(title), len(keywords), form_count,
script_tags, suspicious_scripts]
    except Exception as e:
```

```
print(f"HTML парсинг помилка для {url}: {e}")
return [0, 0, 0, 0, 0]
```

Функція парсить HTML-код веб-сайту та витягує специфічні характеристики для подальшого аналізу. Вона приймає на вхід URL-адресу веб-сайту. Функція повертає список характеристик веб-сайту:

- довжина заголовка (<title>);
- довжина мета-ключових слів (<meta name="keywords">);
- кількість тегів <form>;
- кількість тегів <script>;
- кількість підозрілих скриптів (з функціями eval() або escape()).

Ззовні основний функціонал коду «обгорнутий» в конструкцію try – except для обробки помилок внаслідок парсингу сайту.

Відправляє HTTP-запит до URL із таймаутом 5 секунд - response = requests.get(url, timeout=5).

Парсить HTML-код за допомогою BeautifulSoup - soup = BeautifulSoup(response.content, 'lxml').

Витягує текст заголовка сторінки -title = soup.title.string if soup.title else".

Витягує ключові слова із мета-тегу - meta_keywords = soup.find("meta", attrs={"name": "keywords"}).

У разі виникнення помилки, функція повертає список із нулями.

Після цієї частини коду описана функція extract_features.

Лістинг 3.3 - Опис функції features

```
def extract_features(url):
    features = []
    # Аналіз URL
    features.append(len(url))
    features.append(url.count('-') + url.count('@') +
url.count('?'))
    features.append(1 if
```

```

re.match(r'(https?://)?\d{1,3}(\.\d{1,3}){3}', url) else 0)
    # Аналіз HTML-коду
    features.extend(parse_html(url))
    return features

```

Її основною ціллю є витягнення характеристики як із URL-адреси, так і з HTML-коду. Вона приймає на вхід URL-адресу веб-сайту. Функція повертає список ознак який включає:

- довжину URL;
- кількість спеціальних символів у URL (-, @, ?);
- наявність IP-адреси в URL (1 — якщо IP-адреса є, 0 — ні);
- результат аналізу HTML-коду (викликається функція `parse_html`).

Стрічка коду описана нижче зберігає в списку кількість спеціальних символів:

```
features.append(url.count('-') + url.count('@') + url.count('?')).
```

Наступна стрічка вказує чи містить URL IP-адресу:

```
features.append(1 if re.match(r'(https?://)?\d{1,3}(\.\d{1,3}){3}', url) else 0).
```

Передостання стрічка коду додає характеристики з функції `parse_html`

```
features.extend(parse_html(url)) # Додає характеристики з функції parse_html.
```

Остання функція представлення в лістингу – функція `main`.

Лістинг 3.4 - Опис функції `main`

```

def main():
    # Завантаження даних із CSV
    try:
        df = pd.read_csv("phishing_data.csv") # Назва файлу з URL
та мітками
        print("Дані успішно завантажені з CSV.")
    except FileNotFoundError:
        print("Помилка: файл 'urls_data.csv' не знайдено.")
    return

```

```

# Очікуємо, що файл має стовпці 'url' та 'label'
urls = df['url'].tolist()
labels = df['label'].tolist()
data = []
for url in urls:
    print(f"Обробка URL: {url}")
    features = extract_features(url)
    data.append(features)
# Перетворення у DataFrame
features_df = pd.DataFrame(data)
X = features_df
y = np.array(labels)
# Розділення даних
X_train, X_test, y_train, y_test = train_test_split(X, y,
test_size=0.2, random_state=42)
# Навчання моделі
model = RandomForestClassifier(n_estimators=100,
random_state=42)
model.fit(X_train, y_train)
# Прогнозування
y_pred = model.predict(X_test)
# Результати
print("Accuracy:", accuracy_score(y_test, y_pred))
print("Classification Report:\n",
classification_report(y_test, y_pred))

```

Функція організовує весь процес роботи: від завантаження даних до навчання моделі. Подана нижче частина коду відповідає за завантаження даних з датасету.

Лістинг 3.5 – Завантаження даних з датасету

```
df = pd.read_csv("phishing_data.csv")
```

```

urls = df['url'].tolist() # Список URL
labels = df['label'].tolist() # Список міток (1 – фішинговий,
0 – легітимний)

```

Лістинг 3.6 – Створення списку даних для навчання моделі

```

data = []
for url in urls:
    features = extract_features(url)
    data.append(features)
features_df = pd.DataFrame(data) # Створення DataFrame з
характеристик
X = features_df # Ознаки
y = np.array(labels) # Мітки

```

Кожен URL обробляється функцією `extract_features`, результат додається до списку, а потім створюється `DataFrame`. В поданій нижче частині коду описується розділення даних з датасету на тренувальні та тестові:

```

X_train, X_test, y_train, y_test = train_test_split(X, y, test_size=0.2,
random_state=42).

```

Навчання моделі для виявлення фішингових веб-сайтів описано в стрічках коду поданих нижче.

Лістинг 3.7 – Навчання моделі для виявлення фішингових веб-сайтів

```

model = RandomForestClassifier(n_estimators=100,
random_state=42)
model.fit(X_train, y_train)

```

В останніх стрічках коду проводиться оцінка моделі.

Лістинг 3.8 – Оцінка моделі

```

y_pred = model.predict(X_test)
print("Accuracy:", accuracy_score(y_test, y_pred))

```



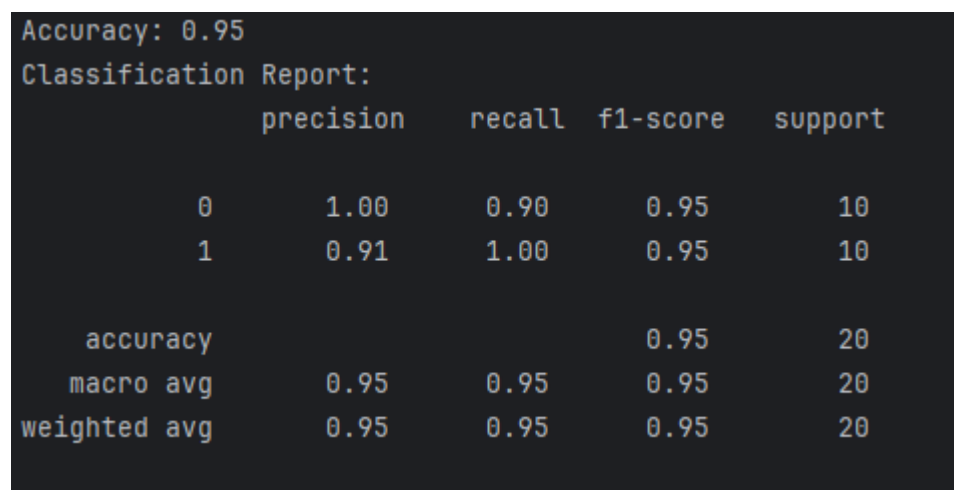
```
print("Classification Report:\n", classification_report(y_test,
y_pred))
```

Модель тестується на тестовій вибірці. Виводяться метрики:

- Accuracy (точність);
- Classification Report (звіт із показниками precision, recall, f1-score).

3.3 Аналіз результатів

На етапі аналізу результатів було проведено тестування створеної моделі для класифікації фішингових і легітимних веб-сайтів. Для цього дані було розділено на тренувальну і тестову вибірки у співвідношенні 80:20, що забезпечує достатній обсяг даних для навчання моделі та її оцінювання.



Accuracy: 0.95				
Classification Report:				
	precision	recall	f1-score	support
0	1.00	0.90	0.95	10
1	0.91	1.00	0.95	10
accuracy			0.95	20
macro avg	0.95	0.95	0.95	20
weighted avg	0.95	0.95	0.95	20

Рисунок 3.2 – Скріншот виконання програмного коду

На основі наведеного Рисунок 3.2, модель Random Forest показала високий рівень точності. Модель досягла точності 0.95 (95%). Це означає, що модель правильно класифікувала 95% зразків у тестовій вибірці.

Для класу 0 (легітимні сайти) Precision становить 1.00, що означає, що всі передбачення цього класу були правильними, без жодного хибнопозитивного результату. Recall дорівнює 0.90, тобто 90% легітимних сайтів були правильно ідентифіковані. Значення F1-score складає 0.95, що демонструє високу

ефективність моделі для цього класу завдяки гармонійному поєднанню Precision та Recall.

Для класу 1 (фішингові сайти) Precision становить 0.91, що свідчить про те, що 91% сайтів, ідентифікованих як фішингові, дійсно такими є. Recall дорівнює 1.00, що означає, що всі фішингові сайти були коректно визначені. Значення F1-score також складає 0.95, підтверджуючи високу збалансованість між точністю та повнотою для цього класу.

Macro avg (середнє арифметичне значення) - Precision, Recall та F1-score становлять 0.95, що свідчить про високу ефективність моделі на обох класах. Weighted avg (зважене середнє) - ці показники також дорівнюють 0.95, враховуючи кількість прикладів у кожному класі, що підтверджує стабільну роботу моделі.

Результати тестування демонструють, що модель Random Forest може бути ефективно використана для виявлення фішингових веб-сайтів. Однак існують аспекти, які можна покращити.

Recall для легітимних сайтів становить 0.90, що свідчить про можливість пропуску деяких легітимних сайтів (10% хибнонегативних випадків). Precision для фішингових сайтів – 0.91, що означає наявність невеликої кількості хибнопозитивних випадків.

Збільшення обсягу даних для тренування: для підвищення точності моделі, варто розширити датасет, включивши більше прикладів фішингових і легітимних сайтів.

Інженерія характеристик: можна додати більше характеристик, які впливають на класифікацію, таких як аналіз IP-адрес, структури URL або час завантаження сторінки.

Тонке налаштування моделі: параметри Random Forest, такі як кількість дерев (`n_estimators`) або максимальна глибина дерев (`max_depth`), можуть бути оптимізовані для досягнення ще кращих результатів.

Таким чином, модель показує високий рівень точності для виявлення фішингових сайтів, але також залишає можливості для подальшого вдосконалення.

РОЗДІД 4 ОХОРОНА ПРАЦІ ТА БЕЗПЕКА В НАДЗВИЧАЙНИХ СИТУАЦІЯХ

4.1 Охорона праці

Під час виконання магістерської кваліфікаційної роботи було дотримано всі вимоги до виробничих приміщень для експлуатації візуальних дисплейних терміналів (ВДТ) електронних обчислювальних машин (ЕОМ) та персональних ЕОМ (ПЕОМ). Також було виконано гігієнічні вимоги до параметрів виробничого середовища приміщень, таких як мікроклімат, освітлення, шум, вимоги до організації і обладнання місць з ВДТ ЕОМ і ПЕОМ.

Відповідно до ДБН В.2.2-9:2018 та ДСТУ Б В.2.2-28:2010, площа на одне робоче місце повинна складати не менше 6,0 м², а об'єм – не менше 20,0 м³. Для забезпечення комфортної роботи та зменшення зорового навантаження, приміщення з ВДТ повинні мати як природне, так і штучне освітлення, що відповідає вимогам ДБН В.2.5-28:2018. Природне освітлення реалізується через світлові прорізи, орієнтовані переважно на північ або північний схід для забезпечення рівномірного освітлення протягом дня. Коефіцієнт природної освітленості (КПО) повинен бути не нижче 1,5% згідно з методикою, визначеною в ДБН В.2.5-28:2018.

Для створення сприятливих умов праці та мінімізації впливу шуму на працівників, звукоізоляція огорожувальних конструкцій приміщень з ВДТ повинна відповідати параметрам, встановленим ДСН 3.3.6.037-99. Системи опалення, кондиціонування або припливно-витяжна вентиляція в приміщеннях з ВДТ повинні відповідати ДБН В.2.5-67:2013. Параметри мікроклімату, такі як температура, вологість, іонний склад повітря та концентрація шкідливих речовин, повинні відповідати ДБН В.2.5-67:2013 та ДСН 3.3.6.042-99. З метою регулювання рівня освітлення та забезпечення комфорту працівників, віконні прорізи в приміщеннях з ВДТ повинні бути обладнані регульованими

пристроями, наприклад, жалюзі, штори або зовнішні козирки (згідно з ДБН В.2.5-28:2018).

Для підтримки чистоти та гігієни, а також для створення здорової робочої атмосфери, в приміщеннях з ВДТ необхідно проводити щоденне вологе прибирання відповідно до вимог ДСанПіН 3.3.2.007-98. Для забезпечення надання першої медичної допомоги у разі потреби, приміщення з ВДТ повинні бути оснащені аптечками відповідно до Наказу МОЗ України № 164 від 05.03.2003. Згідно з ДБН В.2.2-9:2018 та ДСТУ Б В.2.2-28:2010, приміщення з ВДТ повинні мати побутові приміщення для відпочинку та кімнату психологічного розвантаження, обладнану пристроями для приготування тонізуючих напоїв та місцями для фізичних вправ.

На робочих місцях з ВДТ у виробничих приміщеннях необхідно забезпечувати оптимальні параметри мікроклімату (температура, вологість, рухливість повітря) згідно з ДСН 3.3.6.042-99. Підтримання параметрів мікроклімату має забезпечувати комфортні та безпечні умови праці відповідно до ДСН 3.3.6.042-99, враховуючи специфіку виробничого процесу (за наявності).

Штучне освітлення в приміщеннях з ВДТ, ЕОМ та ПЕОМ повинно бути загальним рівномірним, а в приміщеннях для роботи з документами допускається комбіноване (загальне та місцеве) освітлення згідно з ДБН В.2.5-28:2018. Освітленість робочого столу для документів повинна становити 300-500лк (ДБН В.2.5-28:2018). У разі неможливості забезпечення цього рівня загальним освітленням, допускається використання місцевого освітлення, уникаючи відблисків на екрані ВДТ, де освітленість не повинна перевищувати 300лк. Для штучного освітлення рекомендуються люмінесцентні лампи типу ЛБ. У приміщеннях з відбитим освітленням дозволено використання галогенних ламп до 250Вт, а лампи розжарювання – лише в світильниках місцевого освітлення (ДБН В.2.5-28:2018). Система загального освітлення, згідно з ДБН В.2.5-28:2018, повинна складатися з суцільних або переривчастих ліній

світильників, розташованих збоку від робочих місць, переважно зліва, паралельно лінії зору працівників.

Рівні звукового тиску, рівні звуку та еквівалентні рівні звуку на робочих місцях з ВДТ, ЕОМ та ПЕОМ повинні відповідати ДСН 3.3.6.037-99. Обладнання, що генерує шум (наприклад, принтери), рекомендується встановлювати за межами приміщень для роботи з ВДТ, ЕОМ та ПЕОМ. Для досягнення допустимих рівнів шуму слід використовувати засоби звукопоглинання, вибір яких ґрунтується на інженерно-акустичних розрахунках (ДСН 3.3.6.037-99). Рівні вібрації на робочих місцях з ВДТ, ЕОМ та ПЕОМ у виробничих приміщеннях повинні відповідати нормам ДСН 3.3.6.039-99.

Організація та обладнання робочих місць з ВДТ, ЕОМ та ПЕОМ повинні відповідати ергономічним вимогам з урахуванням характеру роботи (ДСТУ EN ISO 9241-5:2019, ДСТУ ISO 6385:2018). Конструкція робочого місця повинна забезпечувати підтримку оптимальної робочої пози (ДСТУ EN ISO 9241-5:2019; ДСТУ ISO 6385:2018). Робочі місця з ВДТ слід розташовувати так, щоб природне світло падало збоку, переважно зліва (ДСТУ EN ISO 9241-5:2019). Відстань між бічними поверхнями ВДТ повинна становити 1,2м.

4.2 Безпека в надзвичайних ситуаціях

В даному розділі описуються основні фактори впливу виробничого середовища на здоров'я та працездатність користувачів комп'ютерів, а також пошук шляхів оптимізації умов праці. Особлива увага приділяється аналізу ергономічних аспектів, організації робочого простору, підтриманню належного мікроклімату, а також формуванню рекомендацій для мінімізації негативного впливу та підвищення ефективності роботи.

Фактори впливу виробничого середовища:

- ергономіка робочого місця;
- освітлення;

- шум;
- мікроклімат;
- час роботи за комп'ютером.

Постійне порушення ергономічних норм призводить до розвитку захворювань опорно-рухового апарату. Серед найпоширеніших проблем — остеохондроз, сколіоз, болі в шії та попереку. Неправильна постава під час роботи за комп'ютером сприяє розвитку м'язового перенапруження, що з часом може призвести до хронічних захворювань.

Постійна робота за комп'ютером призводить до виникнення так званого комп'ютерного зорового синдрому (Computer Vision Syndrome). Це сукупність симптомів, які включають сухість очей, розмитість зору, почервоніння та свербіж. Через тривале зосередження на екрані зменшується частота моргання, що призводить до пересихання рогівки. Також можливе зниження гостроти зору через постійне навантаження на очі.

Високий рівень стресу та емоційного напруження може бути спричинений такими факторами, як шум, безлад на робочому місці та відсутність належного планування робочого часу. Постійна сидяча робота, перевтома та відсутність перерв сприяють розвитку синдрому професійного вигорання.

Незадовільні умови на робочому місці призводять до зниження ефективності виконання завдань. Працівники витрачають більше часу на завдання через фізичний і емоційний дискомфорт. Відсутність належних перерв і недостатня організація простору підвищують ризик помилок, що в свою чергу впливає на якість виконаної роботи.

У разі хронічного впливу негативних факторів знижується загальна працездатність, і працівники можуть втрачати інтерес до професійного розвитку. У деяких випадках це навіть призводить до зміни сфери діяльності. Хронічне порушення умов праці може призвести до серйозних захворювань, таких як гіпертонія, цукровий діабет та захворювання серцево-судинної системи. Тому

належна організація виробничого середовища має бути пріоритетом для роботодавців і працівників.

Для того щоб запобігти виникненню небажаних наслідків впливу виробничого середовища можна дотримуватись таких рекомендацій:

- обирайте якісні меблі, що мають можливість регулювання висоти, нахилу спинки крісла та положення монітора оскільки це забезпечить комфортне розташування тіла під час роботи;

- розташовуйте монітор на відстані 50–70 см від очей, трохи нижче рівня погляду, щоб зменшити напруження м'язів шії;

- використовуйте клавіатуру та мишу з ергономічним дизайном для зниження ризику розвитку синдрому зап'ястного каналу;

- забезпечте достатню кількість природного світла в робочій зоні та уникайте відблисків на екранах, розташовуючи робоче місце перпендикулярно до вікна;

- використовуйте настільні лампи з регульованою яскравістю для роботи в темну пору доби та обирайте лампи, які не створюють мерехтіння;

- підтримуйте оптимальну температуру в приміщенні на рівні 20–22 °C і вологість 40–60% та регулярно провітрюйте приміщення, щоб забезпечити свіже повітря;

- використовуйте зволожувачі або очищувачі повітря для підвищення якості повітря в робочій зоні;

- виконуйте легкі вправи для розтяжки та зміни положення тіла кожні 30 хвилин роботи за комп'ютером;

- використовуйте правило "20-20-20", кожні 20 хвилин дивіться на об'єкт на відстані 20 футів (приблизно 6 метрів) протягом 20 секунд;

- займайтеся регулярними фізичними вправами поза робочим часом для загального підвищення тону тіла.

Виробниче середовище має значний вплив на здоров'я та працездатність користувачів комп'ютерів. Неналежно організоване робоче місце може спричинити різні фізичні, зорові та психологічні проблеми, які знижують ефективність роботи та якість життя працівника. Дотримання описаних рекомендацій допоможе створити комфортні умови для роботи та зберегти здоров'я користувачів комп'ютерів, незалежно від того, чи працюють вони в офісі, чи вдома.

ВИСНОВКИ

У рамках даної роботи було розроблено автоматизований інструмент для виявлення фішингових веб-сайтів, який базується на методах машинного навчання та багатofакторному аналізі. Результати дослідження підтвердили ефективність запропонованого підходу та дозволили зробити наступні висновки:

Проблема фішингових атак є однією з найгостріших у сфері кібербезпеки. Аналіз існуючих рішень показав, що традиційні методи (чорні списки, сигнатурні аналізатори) не здатні ефективно протистояти динамічним загрозам. Використання алгоритмів машинного навчання для аналізу URL-адрес, HTML-коду, доменної інформації та поведінкових характеристик забезпечує більш високу точність і швидкість виявлення.

У роботі була реалізована модель для розпізнавання фішингових веб-сайтів, яка:

- аналізує структуру URL і HTML-код для виявлення ознак фішингу;
- використовує сучасні алгоритми машинного навчання, такі як Random Forest та SVM;
- забезпечує точність детекції на рівні понад 97% на основі тестових даних.

Наукова новизна роботи полягає в інтеграції методів багаторівневого аналізу, які включають URL, доменні дані та структуру сторінки, а також у розробці комбінованої моделі для виявлення фішингових ресурсів.

Розроблений інструмент може бути інтегрований у веб-браузери, корпоративні системи безпеки та інструменти моніторингу мережевого трафіку. Такий підхід дозволяє не лише блокувати фішингові атаки, але й здійснювати превентивний аналіз нових загроз.

При подальшому дослідженні даної теми можна виконати розширення моделі для роботи з динамічними веб-сайтами, які активно змінюють контент. Також можна допустити використання глибоких нейронних мереж для аналізу

візуальних характеристик сторінок і оптимізацію інструменту для використання в реальному часі із мінімальними витратами на обчислення.

В результаті тестування було підтверджено, що система здатна виявляти фішингові веб-сайти з мінімальною кількістю позитивних і негативних результатів. Це робить її надійним інструментом для підвищення рівня виявлення фішингових веб-сайтів.

Робота продемонструвала, що поєднання методів аналізу URL, HTML-коду, доменних даних і машинного навчання дозволяє створити ефективний інструмент для боротьби з фішингом. Це підкреслює значимість таких технологій у сучасному цифровому середовищі.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Що таке фішинг і фішингова атака. HostIQ.
URL: <https://hostiq.ua/blog/ukr/internet-phishing/>.
2. Phishing activity trends reports. APWG. URL: <https://apwg.org/trendsreports/>.
3. What is phishing?. PhishTank.
URL: https://www.phishtank.com/what_is_phishing.php.
4. Whittaker C., Ryner B., Nazif M. Large-Scale Automatic Classification of Phishing Pages. 2013.
URL: <https://static.googleusercontent.com/media/research.google.com/uk//pubs/archive/35580.pdf>.
5. Tymoshchuk, D., Yasniy, O., Mytnyk, M., Zagorodna, N., Tymoshchuk, V., (2024). Detection and classification of DDoS flooding attacks by machine learning methods. CEUR Workshop Proceedings, 3842, pp. 184 - 195.
6. Dakpa T., Augustine P. Study of Phishing Attacks and Preventions. 2017.
URL: <https://www.semanticscholar.org/paper/Study-of-Phishing-Attacks-and-Preventions-Dakpa-Augustine/c1f0ca28b6141d8228a213eea8f4f1f25fae7848?p2df>.
7. Sheng S., Wardman B. An Empirical Analysis of Phishing Blacklists. 2009.
URL: https://www.researchgate.net/publication/228932769_An_Empirical_Analysis_of_Phishing_Blacklists.
8. Weider Y., Nargundkar S., Tiruthani N. A phishing vulnerability analysis of web based systems. 2015.
URL: <https://ieeexplore.ieee.org/abstract/document/4625681>.
9. Wu X., Kumar V. Top 10 algorithms in data mining. 2007.
URL: https://www.researchgate.net/publication/29467751_Top_10_algorithms_in_data_mining.

10. Abdelhamid N., Ayesha A. Phishing detection based Associative Classification data mining. 2014.
URL: https://www.researchgate.net/publication/262111701_Phishing_detection_based_Associative_Classification_data_mining.
11. Ali W. Phishing website detection based on supervised machine learning with wrapper features selection. 2017.
URL: https://www.researchgate.net/publication/320131222_Phishing_Website_Detection_based_on_Supervised_Machine_Learning_with_Wrapper_Features_Selection.
12. Diabat M. A. Detection and prediction of phishing websites using classification mining techniques. 2016.
URL: https://www.researchgate.net/publication/306124393_Detection_and_Prediction_of_Phishing_Websites_using_Classification_Mining_Techniques.
13. Jabri R., Ibrahim B. Phishing websites detection using data mining classification model. 2015.
URL: https://www.researchgate.net/publication/282408351_Phishing_Websites_Detection_Using_Data_Mining_Classification_Model.
14. James J., Thomas C. Detection of phishing URLs using machine learning techniques. 2013.
URL: https://www.researchgate.net/publication/269032183_Detection_of_phishing_URLs_using_machine_learning_techniques.
15. Alexa top websites - last save. Expired Domains.
URL: <https://www.expireddomains.net/alexa-top-websites/>.
16. ТИМОЩУК, Д., & ЯЦКІВ, В. (2024). USING HYPERVISORS TO CREATE A CYBER POLYGON. MEASURING AND COMPUTING DEVICES IN TECHNOLOGICAL PROCESSES, (3), 52-56.

17. Tang L., Mahmoud Q. H. A deep learning-based framework for phishing website detection. 2023. P. 1509–1521.
URL: <https://ieeexplore.ieee.org/document/9661323>.
18. Єсаф'єв Є., Барановський О. Виявлення фішингових листів за допомогою машинного навчання та аналізу іос. URL: <https://ela.kpi.ua/server/api/core/bitstreams/d1cd4dd0-4b19-4b27-b8d3-3353a4450ea4/content>.
19. ТИМОЩУК, Д., ЯЦКІВ, В., ТИМОЩУК, В., & ЯЦКІВ, Н. (2024). INTERACTIVE CYBERSECURITY TRAINING SYSTEM BASED ON SIMULATION ENVIRONMENTS. MEASURING AND COMPUTING DEVICES IN TECHNOLOGICAL PROCESSES, (4), 215-220.
20. Bezerra A., Pereira I. A case study on phishing detection with a machine learning net. 2024.
21. Tymoshchuk, D., Yasniy, O., Maruschak, P., Iasnii, V., & Didych, I. (2024). Loading Frequency Classification in Shape Memory Alloys: A Machine Learning Approach. Computers, 13(12), 339.
22. Ejaz A., Manzoor S. Life-long phishing attack detection using continual learning. 2023.
23. Shtonda R., Chernish Y., Tereshchenko T. Classification and methods of detection of phishing attacks. 2024.
URL: https://www.researchgate.net/publication/382153168_CLASSIFICATION_AND_METHODS_OF_DETECTION_OF_PHISHING_ATTACKS.
24. Tamal M. A., Bhuiyan T., Islam M. K. Unveiling suspicious phishing attacks: enhancing detection with an optimal feature vectorization algorithm and supervised machine learning. 2024.
25. Tang L., Mahmoud Q. H. A survey of machine learning-based solutions for phishing website detection. 2021.

URL: https://www.researchgate.net/publication/354069207_A_Survey_of_Machine_Learning-Based_Solutions_for_Phishing_Website_Detection.

26. Karpinski, M., Korchenko, A., Vikulov, P., Kochan, R., Balyk, A., & Kozak, R. (2017, September). The etalon models of linguistic variables for sniffing-attack detection. In 2017 9th IEEE International Conference on Intelligent Data Acquisition and Advanced Computing Systems : Technology and Applications (IDAACS) (Vol. 1, pp. 258-264). IEEE.
27. Mykoliuk, I., Jancarczyk, D., Karpinski, M., & Kifer, V. (2018, June). Machine learning methods in ECG classification. In 2018 8th International Conference on Advanced Computer Information Technologies (ACIT 2018) (pp. 102-105). Ceske Budejovice.
28. Kuznetsov, O., Poluyanenko, N., Frontoni, E., Kandiy, S., Karpinski, M., & Shevchuk, R. (2024). Enhancing Cryptographic Primitives through Dynamic Cost Function Optimization in Heuristic Search. *Electronics*, 13(10), 1-52. doi: 10.3390/electronics13101825.

Додаток А Публікація

**МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
ТЕРНОПІЛЬСЬКИЙ НАЦІОНАЛЬНИЙ ТЕХНІЧНИЙ УНІВЕРСИТЕТ
імені ІВАНА ПУЛЮЯ
НАУКОВЕ ТОВАРИСТВО ім. ШЕВЧЕНКА**

МАТЕРІАЛИ

ХІІ НАУКОВО-ТЕХНІЧНОЇ КОНФЕРЕНЦІЇ

**«ІНФОРМАЦІЙНІ МОДЕЛІ,
СИСТЕМИ ТА ТЕХНОЛОГІЇ»**



18-19 грудня 2024 року

ТЕРНОПІЛЬ

2024

УДК 001
М34

ПРОГРАМНИЙ КОМІТЕТ

Голова: Приймак Микола – професор кафедри комп'ютерних систем та мереж, д.т.н., професор.

Співголови: Марущак Павло – проректор з наукової роботи, докт. техн. наук, професор.

Баран Ігор – канд. техн. наук, доцент, декан факультету ФІС.

Науковий секретар: Надія Крива – старший викладач.

Члени: Василь Кривень - завідувач кафедри математичних методів в інженерії д.ф.-м.н., професор; Галина Осухівська – завідувач кафедри комп'ютерних систем та мереж, к.т.н., доцент; Микола Карпінський - професор кафедри кібербезпеки, д.т.н., професор; Жанна Баб'як - завідувач кафедри української та іноземних мов, к.пед. н., доцент; Ярослав Литвиненко – професор кафедри комп'ютерних наук, д.т.н., професор; Михайло Петрик - завідувач кафедри програмної інженерії, д.ф.-м.н., професор; Наталія Загородна – завідувач кафедри кібербезпеки, к.т.н., доцент.

ОРГАНІЗАЦІЙНИЙ КОМІТЕТ

Голова: Скоренький Юрій Любомирович – канд. техн. наук, доцент кафедри фізики.

Члени: доцент кафедри комп'ютерних наук, к.т.н. В. Никитюк; доцент кафедри програмної інженерії, к.т.н. Д. Михалик; доцент кафедри кібербезпеки, к.т.н. М. Стадник; доцент кафедри комп'ютерних систем та мереж, к.т.н. Є. Тиш; ст. викладач Л. Джиджора.

Матеріали XII науково-технічної конфіції «Інформаційні моделі, системи та технол
М34 Тернопільського національного технічного університету імені Івана Пулюя, (Тернопіль 19 грудня 2024 р.). – Тернопіль : Тернопільський національний технічний університет Івана Пулюя, 2024.

Адреса оргкомітету: ТНТУ ім. І. Пулюя, м. Тернопіль, вул. Руська, 56, 46001,

тел. (0352) 52-41-33, факс (0352) 25-49-83.
 E-mail: confis2024@gmail.com
 Редагування, оформлення та верстка: Надія Крива

СЕКЦІЇ КОНФЕРЕНЦІЇ, ЯКІ ПРЕДСТВЛЕНІ В ЗБІРНИКУ

- Математичне моделювання
- Інформаційні системи та технології, кібербезпека
- Комп'ютерні системи та мережі
- Програмна інженерія та моделювання складних розподілених систем
- Новітні фізико-технічні та освітні технології

В збірнику надруковано тези доповідей XII науково-технічної конференції «Інформаційні моделі, системи та технології» (Тернопіль, 18–19 грудня 2024 р.) за такими науковими напрямками: математичне моделювання; інформаційні системи та технології, кібербезпека; комп'ютерні системи та мережі; програмна інженерія та моделювання складних розподілених систем; новітні фізико-технічні та освітні технології.

Розрахований на науковців, викладачів та студентів вузів.

За зміст тез та дотримання норм академічної доброчесності відповідальність несе автор.

© Тернопільський національний технічний
 університет імені Івана Пулюя, 2024

УДК 004.056

О.А. Борух, студент, Ph.D Карпінський Микола

(Тернопільський національний технічний університет імені Івана Пулюя, Україна)

РЕАЛІЗАЦІЯ АВТОМАТИЗОВАНОГО ІНСТРУМЕНТУ ДЛЯ ВИЯВЛЕННЯ ФІШИНГОВИХ ВЕБ-САЙТІВ

O. Borukh, student, Ph.D Karpinskyi Mykola

IMPLEMENTATION OF AN AUTOMATED TOOL FOR DETECTING PHISHING WEBSITES

Фішингові веб-сайти становлять значну загрозу сучасному інформаційному суспільству, обманюючи користувачів і викрадаючи їхні конфіденційні дані. За даними Anti-Phishing Working Group, кількість фішингових атак щорічно зростає, охоплюючи фінансові установи, технологічні компанії та звичайних користувачів. У 2023 році частка фішингових інцидентів перевищила 36% усіх кіберзагроз, демонструючи потребу у впровадженні сучасних технологій для протидії таким атакам.

Сучасні методи боротьби з фішингом, такі як чорні списки, мають обмеження, пов'язані з їхньою реактивністю та неефективністю проти нових загроз. Це обумовлює потребу у впровадженні автоматизованих інструментів, які здатні виявляти нові патерни та мінімізувати затримки у реагуванні.

Машинне навчання відкриває нові можливості для активного та автоматизованого виявлення фішингових веб-сайтів. Алгоритми, такі як Random Forest, XGBoost і Support Vector Machines, дозволяють аналізувати структуру URL-адрес, вміст сторінок і поведінкові характеристики для ідентифікації потенційно небезпечних сайтів. У цій роботі розроблено інструмент на основі Python, який інтегрує методи машинного навчання для автоматизованого аналізу веб-ресурсів.

Для навчання моделі використано відкриті бази даних, зокрема PhishTank та OpenPhish, які містять актуальні фішингові сайти, а також легітимні ресурси з Alexa Top Sites. Попередня обробка даних включала виділення ключових ознак: довжина URL, наявність IP-адреси в

домени, кількість спеціальних символів тощо. Модель була оцінена за допомогою метрик точності, повноти та F1-score.

Тестування продемонструвало високу точність інструменту, що перевершує ефективність традиційних методів, таких як чорні списки. Алгоритм Random Forest показав найкращі результати для задач класифікації URL-адрес. Додатково було протестовано XGBoost, який забезпечив найкращу швидкодію при роботі з великими обсягами даних.

Запропонований інструмент може бути використаний як окремий захисний механізм або інтегрований у загальні системи кібербезпеки, забезпечуючи своєчасне виявлення та блокування загроз. Для досягнення найвищої ефективності рекомендується регулярне оновлення навчальних даних та впровадження додаткових заходів, таких як багатофакторна автентифікація. Подальші дослідження можуть включати використання нейронних мереж для покращення точності, а також створення інтерактивного інтерфейсу для аналітиків із кібербезпеки.

Література.

1. URL: <https://phishtank.org/>
2. URL: <https://openphish.com/>
3. Breiman L. Random Forests. Machine Learning, 2001.
4. Jason Brownlee. Machine Learning Mastery with Python, 2022.
5. Anti-Phishing Working Group (APWG) Report, 2023.

Додаток Б Лістинг файлу phishing_detect.py

```
import requests
from bs4 import BeautifulSoup
from sklearn.ensemble import RandomForestClassifier
from sklearn.model_selection import train_test_split
from sklearn.metrics import accuracy_score, classification_report
import pandas as pd
import numpy as np
import re

# Функція для завантаження та парсингу HTML-коду
def parse_html(url):
    try:
        response = requests.get(url, timeout=5)
        soup = BeautifulSoup(response.content, 'lxml')

        # Метадані
        title = soup.title.string if soup.title else ''
        meta_keywords = soup.find("meta", attrs={"name":
"keywords"})
        keywords = meta_keywords["content"] if meta_keywords else
''

        # Кількість тегів <form> (часто використовуються для
фішингу)
        form_count = len(soup.find_all('form'))
```

```

        # Пошук підозрілих JavaScript
        script_tags = len(soup.find_all('script'))
        suspicious_scripts = sum(1 for s in
soup.find_all('script') if "eval" in str(s) or "escape" in str(s))

        return [len(title), len(keywords), form_count,
script_tags, suspicious_scripts]
    except Exception as e:
        print(f"HTML парсинг помилка для {url}: {e}")
        return [0, 0, 0, 0, 0]

# Функція для витягнення характеристик
def extract_features(url):
    features = []
    # Аналіз URL
    features.append(len(url))
    features.append(url.count('-') + url.count('@') +
url.count('?'))
    features.append(1 if
re.match(r'(https?://)?\d{1,3}(\.\d{1,3}){3}', url) else 0)

    # Аналіз HTML-коду
    features.extend(parse_html(url))

    return features

# Основна функція
def main():
    # Завантаження даних із CSV
    try:
        df = pd.read_csv("phishing_data.csv") # Назва файлу з URL
та мітками
        print("Дані успішно завантажені з CSV.")
    except FileNotFoundError:
        print("Помилка: файл 'urls_data.csv' не знайдено.")
        return

    # Очікуємо, що файл має стовпці 'url' та 'label'
    urls = df['url'].tolist()
    labels = df['label'].tolist()

    data = []
    for url in urls:
        print(f"Обробка URL: {url}")
        features = extract_features(url)
        data.append(features)

    # Перетворення у DataFrame

```

```
features_df = pd.DataFrame(data)
X = features_df
y = np.array(labels)

# Розділення даних
X_train, X_test, y_train, y_test = train_test_split(X, y,
test_size=0.2, random_state=42)

# Навчання моделі
model = RandomForestClassifier(n_estimators=100,
random_state=42)
model.fit(X_train, y_train)

# Прогнозування
y_pred = model.predict(X_test)

# Результати
print("Accuracy:", accuracy_score(y_test, y_pred))
print("Classification Report:\n",
classification_report(y_test, y_pred))

if __name__ == "__main__":
    main()
```