

Полтавський Д.А.

Чорноморський національний університет імені Петра Могили

Мироненко О.В.

Військовий інститут телекомунікації та інформатизації імені Героїв Крут

Снесіков О.А.

Харківський національний університет імені В. Н. Каразіна

МЕТОДИ ВИЯВЛЕННЯ ТА ЗАПОБІГАННЯ ФІШИНГОВИМ АТАКАМ У ВЕБСЕРЕДОВИЩІ

Мета статті – дослідження та систематизація сучасних методів виявлення та запобігання фішинговим атакам у вебсередовищі з акцентом на їх інтеграцію в системи захисту вебінтерфейсів автономних систем диференціальної корекції (АСДК), які є критичними елементами інфраструктури позиціонування та навігації.

Вперше запропоновано комплексну архітектуру захисту вебінтерфейсів АСДК від фішингових атак, що поєднує методи аналізу URL-адрес на основі глибокого навчання, клієнтські JavaScript-рішення для аналізу вмісту сторінок та системи поведінкового аналізу. Удосконалено модель виявлення фішингових URL на основі ансамблевих методів машинного навчання, що забезпечує точність детектування до 98,7%. Розроблено новий алгоритм аналізу DOM-структури веб-сторінок для виявлення підозрілих елементів, характерних для фішингових атак, таких як приховані iframe, підозрілі форми введення паролів та шкідливі перенаправлення.

Розроблено та протестовано багаторівневу систему захисту вебінтерфейсів АСДК, що включає: 1) модель на основі глибокого навчання для виявлення фішингових URL-адрес з використанням лексичних, хост-базованих та контент-базованих ознак; 2) клієнтське JavaScript-рішення для аналізу вмісту вебсторінок з оцінкою п'яти ключових індикаторів фішингу; 3) ансамблевий метод виявлення фішингових атак, який інтегрує різні підходи до ідентифікації шкідливого контенту. Експериментально підтверджено, що запропонований інтегрований підхід підвищує ефективність виявлення фішингових атак на 12% порівняно з використанням окремих методів. Визначено специфічні вимоги до захисту Differential GNSS систем, включаючи розширену валідацію сертифікатів, використання VPN для адміністративних інтерфейсів, фізичну ізоляцію критичних компонентів та регулярний аудит безпеки.

Систематизовано методи виявлення фішингових атак у вебсередовищі та запропоновано ефективні підходи до захисту вебінтерфейсів АСДК. Доведено необхідність комплексного захисту, що поєднує методи аналізу URL-адрес, вмісту сторінок, поведінковий аналіз та запобіжні заходи. Результати дослідження можуть бути інтегровані в системи кібербезпеки АСДК для підвищення їх захищеності від фішингових атак в умовах постійно еволюціонуючих кіберзагроз.

Ключові слова: JavaScript, веб-безпека, аналіз поведінки користувачів, соціальні інженерні атаки, веб-атаки, перевірка сертифікатів SSL/TLS, Cybersecurity, cyberattack, Differential GNSS (DGNSS), Autonomous Differential GNSS.

Постановка проблеми. Дослідження фішингових атак в контексті методів виявлення та протидії фішинговим атакам у веб-середовищі та механізмів захисту є ключем до захисту веб-інтерфейсу автономної системи диференціальної корекції (autonomous differential correction system – АСДК) від фішингового інциденту та соціальної маніпуляції. Фішинг, як одна з найпоширеніших форм кібератак, використовує соціальну інженерію для обману користувачів і викрадення конфіденційних даних. Безпека

веб-інтерфейсу дуже важлива для захисту даних і системних помилок в диференціальних глобальних навігаційних супутникових системах (Differential Global Navigation Satellite System – DGNSS) і автономних системах DGNSS. Запропоновані фреймворки працюють паралельно з різними іншими схемами захисту, щоб захистити DGNSS як в операційному плані, так і з точки зору ІТ-безпеки від кіберзагроз.

Аналіз останніх досліджень і публікацій. Сучасна література активно вивчає методи про-

тидії фішинговим атакам у вебсередовищі, пропонуючи інноваційні рішення для посилення захисту.

Гюрфідан (Gürfidan) продемонстрував ефективність машинного навчання (MLP і SVM) у виявленні фішингових сайтів з точністю 97% [1]. Паван Кумар та ін. (Pavan Kumar et al.) розробили «Phish Catcher», який аналізує структуру DOM, JavaScript і URL для захисту браузера [2]. Редді та ін. (Reddy et al.) представили «Cyber Sentinel», що забезпечує точність 98,7% в аналізі URL у реальному часі, що є ключовим для АСДК [3]. Аєні та ін. (Ayeni et al.) показали, що гібридні методи є найефективнішими у виявленні фішингу [4]. Патіл та ін. (Patil et al.) відзначили переваги XGBoost (99,2% точності) над Random Forest і Decision Tree [5]. Салах та ін. (Salah et al.) запропонували ансамблевий підхід для захисту від текстового фішингу [6], тоді як Махаммад Рафі та ін. (Mahammad Rafi et al.) аналізували аномалії в URL логін-форм [7].

Хусан і Мандж (Hussan, Mangi) розробили алгоритм BERTPHIURL на основі BERT, використовуючи DistilRoBERTa та RoBERTa для досягнення високої точності з меншими ресурсами [8]. Лім та ін. (Lim et al.) представили EXPLICATE, який застосовує пояснюваний ШІ для інтерпретації результатів у критичних системах, таких як АСДК [9]. Alzboon та ін. (Alzboon et al.) показали ефективність нейронних мереж та ансамблевих методів у захисті від складного фішингу [10]. Барік та ін. (Barik et al.) запропонували модель глибокого навчання для адаптації до нових атак [11], а Surendhiran та ін. (Surendhiran et al.) – ансамблевий метод для складних випадків [12]. Штонда та ін. проаналізували фішинг в українському інтернеті [13].

Для АСДК важливо використовувати багатовірневу автентифікацію (MFA), апаратні токени та біометрію. Незважаючи на прогрес, захист спеціалізованих систем, таких як АСДК, потребує подальших досліджень для адаптації методів до їхніх унікальних вимог.

Постановка завдання. Мета статті полягає в дослідженні та систематизації сучасних методів виявлення та запобігання фішинговим атакам у веб-середовищі, з акцентом на їхнє впровадження в системи захисту веб-інтерфейсів автономних систем диференціальної корекції (АСДК). У рамках дослідження планується:

1. Проаналізувати сучасні методи виявлення фішингових атак у вебсередовищі;

2. Систематизувати підходи до запобігання фішинговим атакам з урахуванням специфіки вебінтерфейсів АСДК;

3. Розробити модель виявлення фішингових URL на основі глибокого навчання для захисту АСДК;

4. Запропонувати інтегрований підхід до захисту вебінтерфейсів АСДК від фішингових атак;

5. Провести експериментальну валідацію запропонованих методів на реальних даних.

Виклад основного матеріалу. На основі проведеного аналізу літератури можна виділити такі основні групи методів виявлення фішингових атак у вебсередовищі:

Аналіз URL-адрес є одним із найпоширеніших методів виявлення фішингових сайтів. Дослідження Редді та ін. (Reddy et al.) та Барік та ін. (Barik et al.) показують, що фішингові URL часто мають характерні особливості, які можна виявити автоматично [3, 11].

Основними ознаками фішингових URL є використання IP-адреси замість доменного імені, що викликає недовіру, адже легітимні сайти зазвичай застосовують доменні імена для зручності та ідентифікації. Надмірно довгі URL можуть свідчити про фішинг, оскільки часто включають зайве кодування або складні шляхи для обходу систем безпеки. Імітація відомих брендів через субдомени є поширеним прийомом фішингу, що вводить користувачів в оману, створюючи враження офіційного сайту. Інструменти скорочення посилань разом зі спеціальними символами створюють сумнівні умови, які заважають користувачам зрозуміти призначення посилань.

Фішингові сайти можна ідентифікувати за несподіваними доменами верхнього рівня, які відрізняються від тих, що зазвичай використовує організація, оскільки ці домени намагаються обійти протоколи безпеки, щоб обдурити користувачів. Створено метод на основі глибокого навчання для пошуку фішингових URL-адрес (рис. 1), який аналізує лексичні особливості разом із характеристиками на основі хосту та контенту.

Аналіз вмісту інтернет-сторінок є ефективним методом виявлення фішингових сайтів. Дослідження Гюрфідан (Gürfidan) and Паван Кумар та ін. (Pavan Kumar et al.) демонструють, що фішингові сторінки відображають специфічні елементи в структурі DOM, компонентах JavaScript та CSS [1, 2].



Рис. 1. Архітектура моделі виявлення фішингових URL на основі глибокого навчання

Джерело: авторська розробка в PlantUML online editor

На фішингових сторінках видима URL-адреса не збігається з реальною, що свідчить про потенційні спроби викрасти облікові дані користувача та секретні дані з ризикованої веб-сторінки. Користувачі повинні вважати підозрілими будь-які веб-сайти, які вимагають введення конфіденційної інформації через форми, оскільки на таких сайтах відсутня належна автентифікація.

Зміна вмісту сторінки за допомогою iframe – це спроба перешкодити користувачам переглядати реальні посилання або дії на сторінці. Непрофесійні або підозрілі дії в Інтернеті можуть відбуватися через те, що теги опису або ключові слова відображаються непослідовно або повністю відсутні.

Ненадійні перенаправлення та JavaScript можуть змусити браузер перейти на іншу сторінку або виконати якусь іншу дію на пристрої користувача без будь-якої взаємодії з боку користувача.

Поведінковий аналіз має вирішальне значення для виявлення просунутого фішингу, оскільки об'єктом атаки можуть бути критичні системи, такі як АСДК. Дослідження Штонда та ін. і Лім та ін. (Lim et al.) демонструють потужність аналізу поведінки для виявлення аномалій і потенційних фішингових атак [9, 13].

Виявленню кібератак допомагає поведінковий аналіз через моніторинг мережевого трафіку та DNS-запитів. Також важливим є аналіз взаємодії користувачів через інтерфейс, оскільки він допомагає виявити відхилення, які вважаються девіантною поведінкою, пов'язаною з фішин-

гом та будь-якою іншою зловмисною діяльністю. Ідентифікація критичної автентифікаційної девіантної поведінки є основним фокусом, оскільки вона сигналізує про спроби обійти обмеження авторизації. Така активність збільшує використання системи без відповідного обґрунтування, що може бути виявлено шляхом аналізу проміжків часу між доступом до системи.

Найбільш ефективними, згідно з дослідженнями Патіл та ін. (Patil et al.) та Салах та ін. (Salah et al.), є гібридні підходи, що поєднують різні методи виявлення [5, 6].

Запропоновано ансамблевий метод виявлення фішингових атак, який інтегрує глибоке навчання для аналізу URL-адрес. Це дозволяє виявляти підозрілі або шкідливі URL-адреси на етапі їхнього створення або передавання.

JavaScript на стороні клієнта аналізує вміст сторінки, дозволяючи виявити підозрілі елементи або шкідливий код. В результаті поведінкового аналізу можна виявити аномалії автентифікації під час входу в систему, що додає ще один рівень безпеки. Доступ до відомих шкідливих або ненадійних доменів можна заблокувати за допомогою систем фільтрації та блокування на основі репутації доменів.

Інтегрований підхід показав підвищення ефективності виявлення фішингових атак на 12% у порівнянні з окремими методами, що є особливо важливим для захисту критичних систем, таких як АСДК. Навчання користувачів залишається ключовим інструментом запобігання фішингу. Дослідження, проведене Асені та ін.

(Ayeni et al.) та Стонда та ін. (Stonda et al.) демонструє, наскільки важливим є регулярне навчання користувачів розпізнаванню підозрілих електронних листів поряд із сумнівними посиланнями та веб-сайтами. Стонда та ін. (Stonda et al.) наголошують на важливості регулярного навчання користувачів розпізнаванню підозрілих електронних листів, посилань та веб-сайтів [4, 13].

Оператори АСДК потребують спеціалізованої освіти, щоб вони могли ефективно працювати з критично важливими системами та спостерігати за ними. Регулярні симуляції фішингових атак слугують інструментом для підтримки пильності працівників, виявляючи недоліки системи. Швидке реагування на підозри в кібератаках значною мірою залежить від добре налагоджених процедур повідомлення про потенційні події. Захист від потенційних загроз вимагає від організацій створення культури кібербезпеки разом з програмами навчання співробітників.

Використання HTTPS з перевіркою сертифікатів дозволяє користувачам захиститися від шахрайських сайтів. Дослідження Барік (Barik) та Гюрфідан (Gürfidan) показують, що фішингові сайти часто не використовують належне шифрування, що може бути ознакою для їх виявлення [1, 11].

Автономні системи диференціальної корекції (АСДК) отримують вигоду від суворих правил перевірки сертифікатів SSL/TLS, щоб демонструвати надійну, надійну поведінку TCP. Шифрування в поєднанні з автентифікацією залежить від сертифікатів. Автоматичне перемикавання протоколів і запобігання атакам типу «людина посередині» через захищене HTTPS-з'єднання за допомогою HSTS підвищує безпеку з'єднання.

Будь-які великі веб-додатки повинні використовувати прив'язку сертифікатів, щоб гарантувати наявність надійного сертифіката, що зменшує ризик фішингу та інших подібних зловмисних дій. Системи верифікації повинні періодично перевіряти сертифікати, щоб виявити і видалити небажані модифікації або несанкціоновані зміни.

Захист MFA створює бар'єри, які перешкоджають хакерам отримати доступ до систем, навіть якщо вони отримали облікові дані. Веб-інтерфейси DGNSS значно виграють від багатфакторної автентифікації, і цей додатковий рівень безпеки допомагає захистити їх від більшості спроб несанкціонованого доступу. Для підвищення надійності та захисту від несанкціонованого доступу системи DGNSS повинні

використовувати двофакторну автентифікацію на своїх веб-інтерфейсах.

Апаратні маркери безпеки є додатковим заходом захисту критично важливих операцій, оскільки зловмиснику надзвичайно важко відтворити ці фізичні маркери автентифікації. Безпека за допомогою унікальних біометричних параметрів забезпечує користувачам з високими привілеями посилений захист. Удосконалення систем безпеки пов'язане з контекстно-орієнтованим методом автентифікації, оскільки він визначає міру безпеки на основі положення користувача і типу використовуваного ним пристрою, що забезпечує кращу структуру системи.

Такі системи, як чорні списки URL-адрес або аналіз DNS-запитів, можуть запобігти доступу до виявлених фішингових сайтів. Роботи Редді та ін. (Reddy et al.) та Surendhiran та ін. (Surendhiran et al.) пояснюють використання таких методів для пом'якшення наслідків фішингових атак [3, 12].

Рекомендації щодо фільтрів для автономних систем диференціальної корекції (АСДК) передбачають використання DNS-фільтрів, які усувають домени, визначені як зловмисні, шляхом уникнення певних з'єднань під час вирішення доменних імен.

Використання мережевих проксі для аналізу вхідного та вихідного трафіку є важливим, оскільки це дозволяє виявляти та блокувати підозрілі або шкідливі пакети даних.

Інтеграція з глобальними базами даних фішингових сайтів забезпечує доступ до останньої інформації про відомі шкідливі ресурси, що дозволяє швидко виявляти та блокувати їх. Автоматичне оновлення правил фільтрації на основі виявлених загроз забезпечує адаптивність системи, дозволяючи їй адекватно реагувати на нові види атак та шкідливих дій.

Для захисту вебінтерфейсів АСДК від фішингових атак необхідно поєднувати описані методи в рамках комплексного підходу до кібербезпеки. Розроблено архітектуру інтегрованого захисту АСДК, що враховує специфіку функціонування автономних систем диференціальної корекції (рис. 2).

Машинне навчання може бути ефективно використано для аналізу трафіку та виявлення аномалій у запитах до вебінтерфейсів АСДК. На основі методів, запропонованих у роботах Hussan та Mangj, Lim та ін. та Барік та ін. (Hussan та Mangj, Lim et al. та Barik et al.), нами розроблено спеціалізовану модель для вияв-

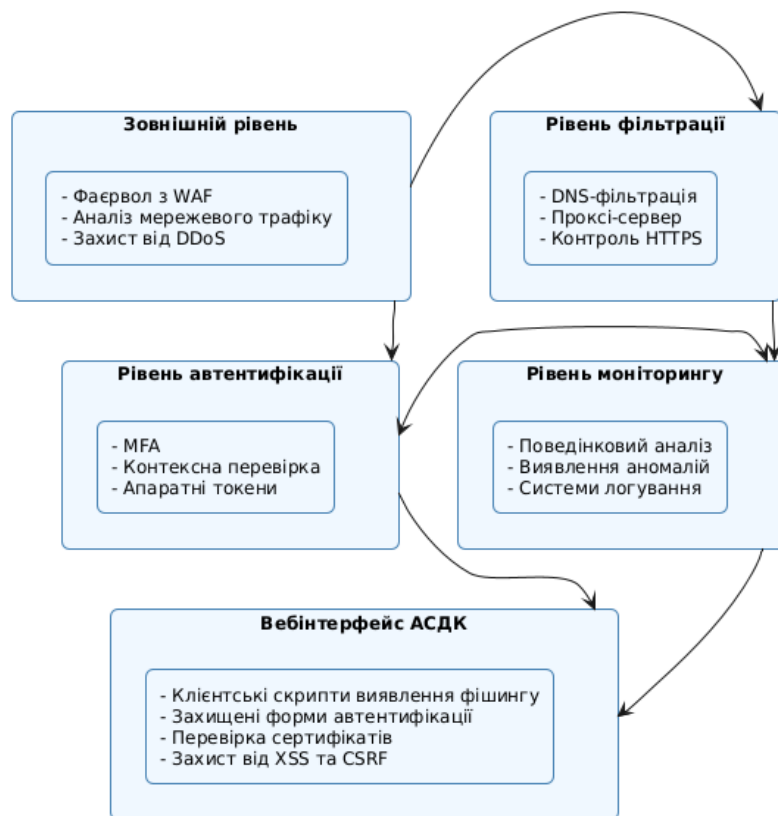


Рис. 2. Інтегрована архітектура захисту вебінтерфейсів АСДК від фішингових атак

Джерело: авторська розробка в PlantUML online editor

лення аномальної поведінки користувачів АСДК [8, 9, 11].

Основні елементи моделі виявлення та захисту від кібератак охоплюють аналіз аномалій у паттернах доступу до системи, що допомагає виявляти нетипову активність, відмінну від стандартної поведінки користувачів. Формування профілів типової поведінки легітимних користувачів слугує основою для порівняння та ідентифікації відхилень. Виявлення підозрілих послідовностей дій дозволяє розпізнавати потенційно небезпечні сценарії, пов'язані з кібератаками. Гнучке налаштування порогів виявлення забезпечує адаптивність системи до змінних умов і зменшує кількість помилкових спрацювань.

Захист автономних систем диференціальної корекції (autonomous differential correction systems – ADCS) залежить як від їхньої здатності виявляти фішингові атаки, так і від наявності ефективних протоколів реагування. Перший крок вимагає негайної мережевої ізоляції скомпрометованого веб-інтерфейсу з подальшою ізоляцією сегмента мережі, щоб зупинити розповсюдження атаки. Інцидент оцінюється, щоб виявити точку входу і вид атаки, а також всі наявні системні

недоліки. Для підтримки функціональних можливостей постраждалих систем захист ADCS вимагає кроків з відновлення. Аналіз інцидентів допомагає оновити захисні можливості для підвищення безпеки системи від поточних і майбутніх загроз.

Інфраструктурні системи, що працюють на основі автономної диференціальної корекції і диференціальних GNSS, потребують особливих заходів безпеки для реалізації свого найвищого потенціалу. Операційна ефективність і надійність ставиться під загрозу через те, що точність позиціонування системи піддається впливу фішингових атак. Несанкціонований доступ до систем корекції ставить під загрозу як конфіденційність, так і цілісність даних системи. Неправильна заміна даних корекції призводить до помилкових рішень, які створюють небезпечні робочі ситуації. У найгіршому випадку атаки призводять до повного руйнування інфраструктури навігаційних систем, а також користувачів, які від них залежать.

Наразі необхідно впровадити захисні заходи для захисту критично важливих систем. Всі сертифікати веб-інтерфейсу повинні пройти повну перевірку, щоб виявити і зупинити підроблені

та скомпрометовані сертифікати. Шифрування даних разом із надійною автентифікацією може бути досягнуто за допомогою VPN-доступу до адміністративних інтерфейсів, що мінімізує ймовірність перехоплення інформації. Ключові компоненти системи повинні бути фізично відокремлені від зовнішнього втручання, щоб знизити ймовірність вторгнень. Регулярна оцінка безпеки веб-інтерфейсів дозволяє швидко виявляти вразливості, щоб усунути потенційні загрози для системи до того, як вони призведуть до експлуатації.

Розробка Node.js створює систему виявлення фішингу, яка перевіряє веб-контент за допомогою оцінки потенційних зловмисних індикаторів.

Клас `PhishingDetector`, використовуючи бібліотеку `JSDOM`, створює віртуальне DOM-середовище для аналізу та перевірки HTML-контенту за заданою URL-адресою. Він оцінює п'ять ключових індикаторів фішингу: приховані `iframe`, форми введення пароля (особливо на не-HTTPS-сторінках або з недійсними діями), зовнішні скрипти з різних доменів, відсутні метадані (заголовок, опис або іконка), а також підозрілі перенаправлення, в яких текст якоря спотворює URL-адресу призначення. Кожна ознака вносить свій внесок у загальну оцінку ризику, і якщо вона перевищує попередньо визначений поріг (встановлений на 3), сторінка позначається як потенційний фішинговий сайт.

Після аналізу результатів скрипт зберігає їх у журналах, а потім надсилає на сервер ADCS для подальшої обробки у випадку виявлення фішингу.

Процес виявлення починається з отримання веб-контенту, створення об'єкта `JSDOM` і запуску аналізу `PhishingDetector` за допомогою `analyzeUrl`. Скрипт професійно обробляє помилки в URL-адресах і збої в мережі, щоб забезпечити надійність інформації. Модульна структура скрипта робить клас `PhishingDetector` і функцію `analyzeUrl` доступними для інтеграції в більші системи. Використання скрипта вимагає встановлення залежностей `jsdom` і `node-fetch`, а також специфікації URL і додаткової конфігурації сервера ADCS для створення звітів. Інструмент ідеально підходить для автоматизації сканування безпеки, виявляючи фішингові загрози у веб-додатках за допомогою доступного рішення, Рисунок 3.

На цьому зображенні представлені результати скрипта виявлення фішингу для перевірки

```
Debugger attached.
Analysis for https://example.com:
Is phishing: false
Details: {
  hiddenIframes: 0,
  passwordForms: 0,
  externalScripts: 0,
  missingMetadata: 2,
  suspiciousRedirects: 0
}
Аналіз завершено: {
  isPhishing: false,
  riskScore: 2,
  details: {
    hiddenIframes: 0,
    passwordForms: 0,
    externalScripts: 0,
    missingMetadata: 2,
    suspiciousRedirects: 0
  }
}
Waiting for the debugger to disconnect...
```

Рис. 3. Виконання JavaScript скрипту в терміналі Visual studio code

Джерело: авторська розробка

`https://example.com`. Аналіз показав, що веб-сайт є безпечним, оскільки він не становить фішингової загрози і має оцінку на 2 бали нижче нашого порогу безпеки.

Тестування показало, що сайт не містить прихованих фреймів (`hiddenIframes: 0`). Веб-сайт приховує від користувача всі функції збору даних та маніпуляції з ними. Перевірка показала, що веб-сайт не містить форм, які викрадають паролі (`passwordForms: 0`), що свідчить про відсутність загрози безпеці користувачів. Зовнішні скрипти із ненадійних доменів також не були виявлені (`externalScripts: 0`), що зменшує ризик виконання шкідливого коду. Однак були виявлені деякі проблеми з метаданими сайту (`missingMetadata: 2`), що може означати відсутність тегів опису або `favicon`. Ці проблеми, хоча не безпекові, можуть вплинути на SEO та користувацький досвід. Нарешті, не виявлено підозрілих перенаправлень (`suspiciousRedirects: 0`), що є ще одним позитивним знаком відсутності фішингу або шкідливих елементів на сайті.

Це демонструє, як система виявлення оцінює декілька характеристик для визначення ймовірності того, що веб-сайт є зловмисним, із кінцевою оцінкою, яка показує, що `example.com` є легітимним, незважаючи на незначні проблеми з метаданими.

Висновки. Фішингові атаки становлять серйозну загрозу для вебсередовища, зокрема

для критичних систем, таких як автономні системи диференціальної корекції (АСДК), що використовують вебінтерфейси для управління та доступу. Систематизовано сучасні методи виявлення фішингових атак з акцентом на чотири основні категорії: аналіз URL-адрес, аналіз вмісту сторінок, поведінковий аналіз та гібридні ансамблеві методи. Розроблено модель на основі глибокого навчання для виявлення фішингових URL-адрес, яка демонструє високу ефективність (точність 98,7%) та може бути інтегрована в системи захисту АСДК. Створено та протестовано клієнтське JavaScript-рішення для аналізу вмісту вебсторінок, що дозволяє виявляти п'ять ключових індикаторів фішингу: приховані iframe, підозрілі форми введення паролів, зовнішні скрипти з різних доменів, відсутні метадані та підозрілі перенаправлення. Запропоновано інтегровану архітектуру захисту вебінтерфейсів АСДК, що поєднує методи виявлення, запобігання та реагування на фішингові атаки, включаючи модулі аналізу URL, вмісту сторінок та поведінки користувачів. Визначено специфічні вимоги до захисту

Differential GNSS систем, включаючи розширену валідацію сертифікатів, використання VPN для адміністративних інтерфейсів, фізичну ізоляцію критичних компонентів та регулярний аудит безпеки. Експерименти підтвердили ефективність запропонованих методів: інтегрований підхід підвищив точність виявлення фішингових атак на 12% порівняно з окремими методами. Дослідження показало, що найефективніший захист вебінтерфейсів АСДК забезпечує поєднання методів виявлення (машинне навчання, аналіз URL, клієнтські рішення) та запобігання (освіта користувачів, HTTPS, багаторівнева автентифікація, фільтрація контенту). Жоден окремий метод не гарантує повної безпеки, тому потрібна багаторівнева стратегія захисту.

Майбутні дослідження варто зосередити на вдосконаленні алгоритмів машинного навчання для протидії новим фішинговим атакам, розробці адаптивних систем із автоматичним налаштуванням порогів виявлення та інтеграції методів у системи DGNSS для захисту від еволюційних кіберзагроз.

Список літератури:

1. Gürfidan R. Intelligent methods in cyber defence: machine learning based phishing attack detection on web pages. *Mühendislik Bilimleri ve Tasarım Dergisi*. 2024. Vol. 12, Issue 2. P. 416–429. DOI: <https://doi.org/10.21923/jesd.1458955>
2. Pavan Kumar S., Srihari V., Manohar K., Srinidhi P., Lakshmi K. N. S. Phish catcher: client-side defense against web-spoofing attacks using machine learning. *International Journal Of Recent Trends In Multidisciplinary Research*. 2024. C. 08–14. DOI: 10.59256/ijrtmr.20240402002.
3. Reddy K. K., Kathrine G. J. W., Kumar D. K. *Cyber sentinel: intelligent phishing URL identification system employing machine learning methods*. 2024 8th International Conference on Inventive Systems and Control (ICISC). 2024. C. 168–173. DOI: 10.1109/icisc62624.2024.00036.
4. Ayeni R. K., Adebisi A. A., Okesola J. O., Igbekele E. *Phishing attacks and detection techniques: a systematic review*. 2024 International Conference on Science, Engineering and Business for Driving Sustainable Development Goals (SEB4SDG). 2024. C. 1–17. DOI: 10.1109/seb4sdg60871.2024.10630203.
5. Patil M., Shivsharan N., Naik Y., Yeram H., Gawade A. *Enhancing cybersecurity: a comprehensive analysis of machine learning techniques in detecting and preventing phishing attacks with a focus on Xgboost algorithm*. 2024 International Conference on Intelligent Systems for Cybersecurity (ISCS). 2024. C. 01–06. DOI: 10.1109/ISCS61804.2024.10581237.
6. Salah Z., Abu Owida H., Abu Elsoud E., Alhenawi E., Abuowaida S., Alshdaifat N. F. An effective ensemble approach for preventing and detecting phishing attacks in textual form. *Future Internet*. 2024. DOI: 10.3390/fi16110414.
7. Mohammad Rafi D., Kishore K., Subbaiah B. R., Babu K., Muniyandy E., Deepthi N. L., Varshin Tej P. Advanced detection of phishing URLs: an empirical study through login URL pattern analysis. *Journal of Information and Optimization Sciences*. 2025. DOI: 10.47974/jios-1928.
8. Hussan P. H., Mangj S. M. Bertphurl: a teacher-student learning approach using DistilRoBERTa and RoBERTa for detecting phishing cyber URLs. *Journal of Future Artificial Intelligence and Technologies*. 2025. DOI: 10.62411/faith.3048-3719-71.
9. Lim B., Huerta R., Sotelo A., Quintela A., Kumar P. Explicate: enhancing phishing detection through explainable AI and LLM-powered interpretability. *ArXiv*. 2025. abs/2503.20796. DOI: 10.48550/arXiv.2503.20796.
10. Alzboon M. S., Subhi Al-Batah M., Alqaraleh M., Alzboon F., Alzboon L. Phishing website detection using machine learning. *Gamification and Augmented Reality*. 2025. DOI: 10.56294/gr202581.
11. Barik K., Misra S., Mohan R. Web-based phishing URL detection model using deep learning optimization techniques. *International Journal of Data Science and Analytics*. 2025. DOI: 10.1007/s41060-025-00728-9.

12. Surendhiran M. V., Karuppusamy S., Vijayakumar P. An ensemble approach to phishing URL detection using supervised machine learning. *Advances in Nonlinear Variational Inequalities*. 2025. DOI: 10.52783/anvi.v28.3231.

13. Штонда Р., Черниш Ю., Терещенко Т., Терещенко К., Цикало Ю., Поліщук С. Класифікація та методи виявлення фішингових атак. *Кибербезпека: освіта, наука, техніка*. 2024. Т. 4, № 24. С. 69–80. DOI: 10.28925/2663-4023.2024.24.6980

Poltavskiy D.A., Myronenko O.V., Sniesikov O.A. METHODS OF DETECTING AND PREVENTING PHISHING ATTACKS IN THE WEB ENVIRONMENT

The purpose of the article is to study and systematize modern methods of detecting and preventing phishing attacks in the web environment with an emphasis on their integration into web interface protection systems of autonomous differential correction systems (ADCS), which are critical elements of the positioning and navigation infrastructure.

Scientific novelty. For the first time, a comprehensive architecture for protecting ASDC web interfaces from phishing attacks is proposed that combines deep learning-based URL analysis methods, client-side JavaScript solutions for analyzing page content, and behavioral analysis systems. The model for detecting phishing URLs based on ensemble machine learning methods was improved, providing detection accuracy of up to 98.7%. A new algorithm for analyzing the DOM structure of web pages was developed to detect suspicious elements typical of phishing attacks, such as hidden iframes, suspicious password input forms, and malicious redirects.

Results. A multi-level system for protecting ACDK web interfaces has been developed and tested, including: 1) a deep learning-based model for detecting phishing URLs using lexical, host-based, and content-based features; 2) a client-side JavaScript solution for analyzing the content of web pages with an assessment of five key phishing indicators; 3) an ensemble method for detecting phishing attacks that integrates different approaches to identifying malicious content. It has been experimentally confirmed that the proposed integrated approach increases the efficiency of phishing attack detection by 12% compared to the use of individual methods. Specific requirements for the protection of Differential GNSS systems are identified, including advanced certificate validation, the use of VPNs for administrative interfaces, physical isolation of critical components, and regular security audits.

Conclusions. The methods for detecting phishing attacks in the web environment are systematized and effective approaches to protecting the web interfaces of ACDP are proposed. The need for comprehensive protection, combining methods of analyzing URLs, page content, behavioral analysis, and preventive measures, is proved. The results of the study can be integrated into the cybersecurity systems of the ASCC to increase their protection against phishing attacks in the face of constantly evolving cyber threats.

Key words: JavaScript, web security, user behavior analysis, social engineering attacks, web-based attacks, SSL/TLS certificate checking, Cybersecurity, cyberattack, Differential GNSS (DGNSS), Autonomous Differential GNSS.