

**МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
ЄВРОПЕЙСЬКИЙ УНІВЕРСИТЕТ
АСОЦІАЦІЯ НАВЧАЛЬНИХ ЗАКЛАДІВ УКРАЇНИ ПРИВАТНОЇ
ФОРМИ ВЛАСНОСТІ
ГО «АСОЦІАЦІЯ СПЕЦІАЛІСТІВ КІБЕРБЕЗПЕКИ»
АКАДЕМІЯ WSB У ДОМБРОВІ ГУРНІЧІЙ (ПОЛЬЩА)
УНІВЕРСИТЕТ TURİVA (ЛАТВІЯ)
НАЦІОНАЛЬНА ІНЖЕНЕРНА АКАДЕМІЯ
РЕСПУБЛІКИ КАЗАХСТАН**

**АКТУАЛЬНІ ПИТАННЯ ЗАБЕЗПЕЧЕННЯ КІБЕРБЕЗПЕКИ
ТА ЗАХИСТУ ІНФОРМАЦІЇ**

Матеріали

*XI Міжнародної науково-практичної конференції «Актуальні питання
забезпечення кібербезпеки та захисту інформації»*

24 квітня 2025 р.



Київ-2025
Видавництво Європейського університету

УДК: 004(063)

Редакційна колегія: Тимошенко О. І. - ректор, доктор філософських наук, професор,

Гудзь Ю.Ф. – доктор економічних наук, доцент,

Яровий Р.О. - кандидат технічних наук, доцент,

Скляренко О. В. – кандидат фізико-математичних наук, доцент

Відповідальні секретарі:

к.ф.-м.н., доцент Скляренко О.В., к.т.н., доцент Яровий Р.О.

Актуальні питання забезпечення кібербезпеки та захисту інформації:

матеріали XI Міжнародної науково-практичної конференції, 24 квітня 2025 р., Київ / редкол.: О. І. Тимошенко [та ін.]. – Київ: Вид-во Європейського університету, 2025. – 163 с.

Збірник містить матеріали XI Міжнародної науково-практичної конференції «Актуальні питання забезпечення кібербезпеки та захисту інформації».

Матеріали друкуються за редакцією авторів

© Європейський університет, 2025

ЗМІСТ

Тимошенко О.І.

Актуальні питання забезпечення кібербезпеки та формування цифрової обізнаності у здобувачів освіти в умовах сучасних викликів.....8

Аксьонова І.В., Гаврилова А.А.

Фреймворк аналізу кіберзагроз як інструмент управління корпоративною безпекою.....10

Букатов Д.В., Романенко О.І.

Сучасні підходи до захисту графічного контенту.....13

Бойко М.М.

Кібергігієна як первинний етап захисту в умовах повномасштабної війни.....14

Вах О.Ф., Коцун В.І.

Проблема несанкціонованого доступу до інформації в цілях навчання штучного інтелекту.....16

Ващенко М.В.

Формування цифрової стійкості як критичного елементу довіри до суб'єктів господарювання в умовах ескалації кіберзагроз.....17

Венжик А.О., Ковальчук М.А.

Кібербезпека в Java-додатках: можливості інтеграції штучного інтелекту для захисту систем.....18

Вілянський А.В., Позняк А.А.

Організаційні аспекти впровадження моделі Zero Trust у корпоративному середовищі.....20

Волкова Н.М.

Розвиток навичок критичного мислення та кібербезпеки у школярів як протидія інформаційно-психологічним маніпуляціям та впливам.....22

Галушка П.А., Коробейнікова Т.І.

Аналіз ефективності симетричних та асиметричних методів шифрування у VPN.....25

Гетьман О.І.

Оцінка ефективності захисту API у гібридному середовищі мобільного сервісу.....27

Гордон-Басанський Ю.В.

Автоматизація тестування застосунків за допомогою Selenium в контексті кібербезпеки України..30

Гудаков Д.О.

Аспекти кібербезпеки при інтеграції штучного інтелекту в процеси управління IT-проєктами.....31

Дейнека О.І.

Розподілене прийняття рішень у багатодронових системах з використанням MARL.....33

Денисюк В.М.

Інтеграція моделей валідації вразливостей у SIEM-системи для підвищення ефективності виявлення та реагування на кіберзагрози.....35

Дмитрів Н.С., Коцун В.І.

Використання блокчейн-технологій для забезпечення цифрової безпеки держави в умовах війни..37

Долгополов Н.О.	
Інтеграція машинного навчання та теорії ігор у системи критичної інфраструктури для виявлення кібератак.....	38
Доровський Д.В., Хабеця Є.О.	
Моделі і методи ідентифікації просторово-часових ознак в системах раннього попередження про надзвичайні ситуації.....	40
Доровська І.О., Шабанов Б.Ш. огли	
Моделі і методи об'ємного відтворення динамічних процесів за результатами дистанційних спостережень в інформаційних системах.....	42
Ємельянов В.К.	
Інтелектуальні методи ройового управління групою динамічних об'єктів при уникненні від переслідування.....	45
Ємельянов Н.М.	
Кібербезпека міжнародних фінансових операцій.....	46
Єрмоленко М.В., Бойко М.М.	
Використання ChatGPT для посилення комп'ютерної безпеки організацій.....	48
Захаренков Д.Ю.	
Сучасні загрози для мереж п'ятого покоління 5G.....	49
Іванченко Є.В., Туровський О.Л., Рижаков М.М.	
Розвиток новітніх методів і технологій для забезпечення кібербезпеки та захисту критичної інформаційної інфраструктури України в умовах сучасних глобальних загроз та військового конфлікту.....	51
Ільящук С.С., Милашенко В.М.	
Виявлення фейкових матеріалів, згенерованих штучним інтелектом, у соціальних мережах: виклики та підходи до протидії.....	57
Калмиков А.М., Доровський В.О.	
Метод пошуку безпечних траєкторій руху безпілотних апаратів.....	58
Кіт Ю.А., Коробейнікова Т.І.	
Огляд міжнародних стандартів управління кібербезпекою для малих підприємств.....	60
Ковальов О.С.	
DevSecOps для університетських систем: кейс впровадження у CRM деканату.....	62
Корбут В.В.	
Штучний інтелект в комп'ютерній графіці.....	63
Коцун В.І., Сушинський О.Є.	
Стратегії посилення кібербезпеки корпоративних даних при використанні Salesforce platform: багаторівневий підхід.....	65
Кулагін В.П.	
Напрямки оптимізації мікросервісної архітектури.....	67

Курючкін А.Е. Дослідження алгоритмів штучного інтелекту для виявлення аномальної поведінки гравців в онлайн-іграх.....	69
Левченко С.В., Савченко Я.О. Вплив розвитку інформаційних технологій на методи та засоби навчання студентів.....	71
Литвин А.А. Порівняння архітектур нейронних мереж для виявлення аномалій у часових рядах у контексті кібербезпеки.....	72
Литвиненко Б.В. Методи моніторингу об'єктів захисту інформації на основі технології Data Mining.....	75
Литвиненко Л.О., Бойко М.М. Впровадження штучного інтелекту у системи кіберзахисту: можливості та виклики.....	78
Максимович М.В. Сучасні інструменти автоматизованого тестування безпеки програмного забезпечення.....	80
Мойсеєнко Є.В., Яровий Р.О. Аналіз стану кібербезпеки фінансових установ та методи протидії кібератакам.....	83
Ніколаєвський О.Ю., Ващенко М.В. Стратегічна сліпота у корпоративному управлінні як фактор зростання кібервразливості економіки України в умовах збройного конфлікту.....	85
Овсянко Д.О. Розподілені цифрові двійники в публічних мережах: виклики конфіденційності та підходи на основі блокчейну.....	86
Одольський М.Ю. Вплив штучного інтелекту на безпеку і якість програмних систем: погляд з позиції QA-спеціаліста.....	88
Опольський М.В., Позняк А.А. Огляд кібератаки на енергетичну та телекомунікаційну інфраструктуру України: аналіз інцидентів з «Укренерго» та «Київстар».....	90
Павлів Т.Ю., Коробейнікова Т.І. Методи захисту в мережах WPA2 та WPA3: від Krack до Dragonblood.....	92
Пашорін В.І., Ващенко М.В., Саморай О.К. Між стратегічною індиферентністю та цифровою стійкістю: проблема усвідомлення важливості кібербезпеки керівництвом компаній в умовах війни.....	94
Переверзєв А.М., Подвиженко А.В., Левченко С.В. Безпека хмарних обчислень: захист даних та інфраструктури при використанні SAAS, PAAS, SAAS рішень.....	96
Подвиженко А.В., Переверзєв А.М. Публічні Wi-Fi мережі та захист особистої інформації.....	99

Позняк А.А.	
Підвищення стійкості кібербезпеки України в умовах війни шляхом впровадження ефективних практик автоматизованого тестування безпеки.....	101
Покидько Д.Ю.	
Соціальна інженерія в умовах війни: концепція створення бота для виявлення цифрових загроз у месенджерах.....	102
Рихальський О.Ю.	
Шляхи підвищення кіберстійкості програмного забезпечення: сучасні мови програмування та супутні методології.....	105
Саморай О.К., Ващенко М.В.	
ШІ-оптимізовані SIEM-системи: підвищення швидкості та точності виявлення інцидентів.....	108
Світличний Д.О.	
Аналіз сучасного стану кіберзагроз для об'єктів критичної інфраструктури та існуючих підходів до їх класифікації.....	110
Семенко А.І., Бойко М.М.	
Аналіз аномальної активності в мережах за допомогою машинного навчання.....	112
Семенюк О.О.	
Окремі аспекти використання методів OSINT.....	113
Сенишин Р.В., Коцун В.І.	
Потенційні ризики застосування генеративного штучного інтелекту для автоматизованого створення документації та технічних специфікацій.....	116
Скицький Т.Р., Коцун В.І.	
Захист персональних даних користувачів з інвалідністю у просторових обчислювальних середовищах.....	117
Скляренко О.А., Позняк А.А.	
Безпечне збереження та передача даних у мобільних застосунках.....	120
Скляренко О.В., Ващенко М.В.	
Атаки через непрямі вектори у цифровому середовищі як інструмент посилення кіберзагроз проти підприємств в умовах гібридної війни.....	122
Снігур М.М.	
Актуальні питання забезпечення кібербезпеки: захист критичної інфраструктури.....	124
Стрілець Є.А.	
Технологія забезпечення безпеки документів в системах електронного документообігу.....	125
Сушинський О.Є., Коцун В.І.	
Інноваційні методи захисту конфіденційної інформації в екосистемі Salesforce: виклики та рішення.....	127
Терещук Г.М., Скляренко П.А.	
Використання технології Zero Trust для підвищення кібербезпеки бізнес-процесів в умовах цифрових загроз.....	129

Тритиниченко О.О. Штучний інтелект проти штучного інтелекту: потенціал і ризики використання генеративних моделей у кібератаках та системах захисту.....	131
Троян К.М. Перспективи трансформер-моделей в аналізі фондового ринку.....	132
Федік О.І., Складенко О.В. Питання кібербезпеки при впровадженні цифрових рішень у бізнес-процеси в умовах воєнного стану.....	134
Фролов І.М., Колодінська Я.О. Кіберзахист освітніх стартапів: виклики цифрової безпеки та інноваційні рішення.....	138
Чорненький М.В. Методи безпечного зберігання конфіденційних даних CRM/ERP-систем.....	141
Шаповалов Б.Д., Коробейнікова Т.І. Порівняльний аналіз міжнародних стандартів управління кіберризиками в корпоративному секторі.....	143
Шевчук А.А. Захист цифрового активу в умовах цифровізації та впровадження інструментів штучного інтелекту в управління бізнес процесами.....	144
Шугалій Д.А., Яровий Р.О. Сучасний ландшафт кіберзагроз та методи захисту.....	147
Ямнич А.Б., Коробейнікова Т.І. Інноваційні методи контролю доступу персоналу в інформаційних системах підприємств на основі блокчейн-технологій.....	149
Янішевський В.І. Моделювання загроз для систем управління навчанням, підсилених штучним інтелектом.....	151
Яровий Р.О., Рихальський О.Ю. Вплив вибору мови програмування на стійкість програмного забезпечення до кібератак.....	154
Яровий Я.Р. Кіберзахист критичних об'єктів залізничної інфраструктури.....	156
Bilyi V.O., Kamenev M.S. Analysis of the use of SIEM systems for threat monitoring in critical infrastructure.....	157
Glova A.S., Kotsun V.I. AI-enhanced phishing simulations: opportunities and risks.....	159
Melko T., Kotsun V. Adversarial attacks on cybersecurity with machine learning in defence.....	161

АКТУАЛЬНІ ПИТАННЯ ЗАБЕЗПЕЧЕННЯ КІБЕРБЕЗПЕКИ ТА ФОРМУВАННЯ ЦИФРОВОЇ ОБІЗНАНОСТІ У ЗДОБУВАЧІВ ОСВІТИ В УМОВАХ СУЧАСНИХ ВИКЛИКІВ

Тимошенко О.І.

*доктор філософських наук, професор,
ректор ПВНЗ «Європейський університет», м.Київ*

Проблематика кібербезпеки та захисту інформації набуває особливої актуальності в контексті глобальних геополітичних та технологічних трансформацій, особливо після повномасштабного збройного вторгнення російської федерації в Україну у лютому 2024 року. Відтак, наукові дослідження та прикладні розробки у сфері кіберзахисту посідають чільне місце в системі забезпечення національної безпеки.

Від 2014 року на базі Приватного вищого навчального закладу «Європейський університет» проводиться щорічна Міжнародна науково-практична конференція «Актуальні питання забезпечення кібербезпеки та захисту інформації». Її мета полягає в обміні досвідом, поширенні інноваційних підходів, результатів наукових досліджень і формуванні фахової спільноти. До збірника тез, підготовленого за підсумками даної конференції, увійшли праці, що охоплюють концептуальні, правові та технічні аспекти безпеки цифрового середовища. Зокрема, дослідження присвячені функціонуванню та захисту об'єктів критичної інфраструктури в умовах воєнного стану, кіберзагрозам в економічному та управлінському секторі, вдосконаленню комплексних систем захисту інформації, правовому та адміністративному регулюванню у сфері кібербезпеки, а також розвитку методів захисту інформації в управлінні корпоративною кібербезпекою.

У сучасному інформаційному суспільстві, де цифрові технології відіграють ключову роль у функціонуванні державних, військових і корпоративних структур, зростає потреба у надійних та гнучких інструментах протидії кіберзагрозам. Окрім традиційних атак типу DoS/DDoS, дедалі частіше фіксуються високотехнологічні кампанії, спрямовані на викрадення інформації, маніпуляцію даними, саботаж та втручання в цифрову інфраструктуру [1].

Особливу небезпеку становить неконтрольоване поширення пристроїв, підключених до мережі Інтернет (Інтернет речей, IoT), більшість із яких характеризуються низьким рівнем захисту, що робить їх потенційними векторами проникнення для кібератак.

Глобалізація кіберзлочинності є додатковим фактором, який ускладнює процес виявлення та протидії загрозам: транснаціональні хакерські угруповання діють через кордони, об'єднуються у децентралізовані мережі, використовуючи сучасні засоби шифрування, цифрові валюти та ботнети [2].

Правове регулювання питань кібервійни й досі залишається фрагментарним. Наявні міжнародні правові норми лише частково охоплюють сферу інформаційної безпеки, що ускладнює реагування на масштабні кіберінциденти в межах міждержавних конфліктів.

У цьому контексті надзвичайно важливим є підвищення рівня цифрової обізнаності та кіберграмотності серед усіх категорій користувачів: як фахівців ІТ-сфери, так і пересічних громадян. Саме людина залишається найуразливішим елементом у системі інформаційної безпеки, а тому якісна освіта та підготовка кадрів — ключовий чинник кіберстійкості держави [3, 4].

Розбудова національної системи кібербезпеки потребує багатовекторного підходу, який поєднує такі ключові аспекти як вдосконалення нормативно-правового поля, розвиток технологічної інфраструктури, міжсекторальну координацію (державний, приватний, академічний сектори), міжнародну співпрацю в галузі безпеки кіберпростору, просвітницьку та науково-освітню діяльність.

Отже, комплексне вирішення проблем інформаційної та кібербезпеки можливе лише за умов інтеграції технологічних, політичних, правових та освітніх зусиль. У сучасних реаліях війни та гібридних загроз Україна потребує системного розвитку кіберінфраструктури, підвищення захищеності критичних об'єктів та зміцнення кадрового потенціалу у сфері безпеки даних. Наукові дослідження, які спрямовані на формування міждисциплінарних моделей оцінки ризиків, адаптацію міжнародних стандартів до українського законодавства та створення стійких механізмів протидії високотехнологічним кіберзагрозам є надзвичайно актуальними і затребуваними в умовах сучасних викликів.

Список використаної літератури

1. *Актуальні питання забезпечення кібербезпеки та захисту інформації: колективна монографія / за заг. наук. ред. А.М. Давиденко, Київ: Європейський університет, 2023. – 240 с.*
2. *Тимошенко О. І., Литвиненко Л. О., Колодінська Я. О. Загрози та безпека кіберпростору в умовах сучасних викликів: проблеми, інструменти, рішення. Актуальні питання забезпечення кібербезпеки та захисту інформації: колективна монографія / Київ: Європейський університет, 2023. – С. 10-18.*
3. *Тимошенко О. І., Гудзь Ю. Ф. Передмова: теоретичний огляд напрацювань проблематики безпекового середовища воєнного стану. Наш 2023 рік. Формування безпекового середовища в умовах воєнного стану: колективна монографія / за заг. наук. ред. О. І. Тимошенко. — Київ: Європейський університет, 2023. – С. 5-12.*
4. *Leonid Arsenovych, Oleksandr Nikolaievsky, Olena Skliarenko, Leonid Lytvynenko, Ivan Kydriavskiy. Organization of Training with the Use of Digital Technologies for Ensuring Cybersecurity in the Educational Space // WSEAS Transactions on Computer Research, ISSN / E-ISSN: 1991-8755 / 2415-1521, Volume 12, 2024, p. 524-536, DOI: 10.37394/232018.2024.12.51 , <https://wseas.com/journals/articles.php?id=9988>*

ФРЕЙМВОРК АНАЛІЗУ КІБЕРЗАГРОЗ ЯК ІНСТРУМЕНТ УПРАВЛІННЯ КОРПОРАТИВНОЮ БЕЗПЕКОЮ

*Аксьонова І.В.,
к.е.н., доцент кафедри кібербезпеки,
Гаврилова А.А.,
PhD, доцент кафедри кібербезпеки,
Національний технічний університет
«Харківський політехнічний інститут»*

В сучасному цифровізованому світі однією з найважливіших сфер інформаційної безпеки є кібербезпека. Даний аспект пов'язаний з тим, що цифрове суспільство породжує різке зростання кількості кіберзагроз, інтенсивність та складність кібератак, що вимагає уряд, бізнес та окремих користувачів серйозніше підходити до захисту даних та мереж. Одним із ключових напрямів у забезпеченні кібербезпеки виступає побудова фреймворків аналізу кіберзагроз. Фреймворки дозволяють структурувати та автоматизувати процеси виявлення, аналізу, класифікації та реагування на кіберзагрози.

Велику цінність для ідентифікації кіберзагроз мають різноманітні глобальні дослідження в кіберзахисній сфері, зокрема, звіт ENISA Threat Landscape Report [2]. За результатами даного звіту можна виділити наступні сучасні тренди серед основних загроз кібербезпеки (табл. 1). Загрози в табл. 1. розташовані по рейтингу відповідно до обсягу дослідження інцидентів безпеки.

Таблиця 1

Сучасні тренди серед основних загроз кібербезпеки

Вид кіберзагрози	Тренд (напрямок розвитку за останні роки)
1. Шкідливе програмне забезпечення	Стабільний ↔
2. Веб-атаки та атаки веб-додатків	Має тенденцію до зростання ↑
3. Фішинг, спам, DoS-атаки	Має тенденцію до зростання ↑
4. Програми-вимагачі та ботнети	Має тенденцію до зростання ↑
5. Внутрішні загрози (зловмисні, випадкові)	Стабільний ↔
6. Фізичні маніпуляції /пошкодження/ крадіжка/втрата	Стабільний ↔
7. Порушення даних та крадіжка персональних даних	Має тенденцію до зростання ↑
8. Витік інформації та кібершпигунство	Має тенденцію до зростання ↑

Як видно з табл. 1. в графі «Тренд» показано відносну зміну тяжкості наслідків кожної загрози і напрямом її зміни за останні роки.

Виявлення ознак та видів кіберзагроз не лише підсилює ситуаційну обізнаність, а й є основою архітектури фреймворків, які використовують класифікатор як компонент для логіки розпізнавання загроз. Будь-яку кіберзагрозу можливо ідентифікувати на основі характерних для неї ознак. Така ідентифікація дозволяє розкрити сутність загрози, а, відповідно, визначити релевантні заходи з її нейтралізації: від профілактичних дій (у разі, коли інцидент ще не настав) до

реагування та мінімізації наслідків (якщо загроза вже реалізується у вигляді активного кібервпливу).

Даний підхід формує основу для виділення ключових кіберзагроз у заданому середовищі, а також ознак, за якими ці загрози можуть бути вчасно виявлені. Саме систематизація ознак і відповідних реакцій на загрози є вихідною точкою для побудови фреймворків аналізу кіберзагроз, які є структурованими моделями, що автоматизують процес виявлення, класифікації та реагування на загрози інформаційній безпеці.

В дослідженнях [3,4] пропонуються різні підходи до структури класифікатора кіберзагроз, однак не враховуються можливості розподілу загроз за аспектами безпеки, а саме: кібербезпекою (КБ), інформаційною безпекою (ІБ), захистом інформації (ЗІ) та їх впливом на сервіси безпеки, такі як конфіденційність, цілісність, автентичність, доступність і невідомість. Автори роботи [7] розглядають синергетичний підхід до моделювання загроз, але не враховують потенційний вплив методів соціальної інженерії, що значно посилює реалізацію цільових загроз. У роботі [8] запропоновано загальний підхід до універсалізації класифікаторів загроз, проте не враховується необхідність створення багатоконтурних систем захисту в умовах багатоплатформених середовищ, до яких належать кіберфізичні системи.

Також для побудови ефективного класифікатора кіберзагроз необхідно враховувати низку вимог, які охоплюють технічні, функціональні та організаційні аспекти. Це забезпечує можливість якісного виявлення, моніторингу та аналізу загроз. Основні вимоги включають: гнучкість і масштабованість, інтеграцію з наявними інструментами безпеки, автоматизацію збору та аналізу даних, адаптивність до нових типів загроз, функціональність аналізу вразливостей і ризиків, прозорість і доступність звітності, захист і конфіденційність даних, відповідність стандартам інформаційної безпеки. [5,6]

Враховуючи вище розглянуті недоліки та вимоги до фреймворків аналізаторів кіберзагроз, дослідниками НТУ «ХПІ» удосконалено уніфікований класифікатор загроз з урахуванням гібридності та синергізму цільових (змішаних) атак. При цьому, під гібридністю цільових атак розглядається вплив атаки на одну з послуг безпеки (конфіденційність, цілісність, доступність, автентичність, корисність) для всіх компонентів безпеки (кібербезпека, інформаційна безпека, загальна інформаційна безпека), а під синергізмом таких атак мається на увазі вплив на одну складову безпеки, проте на всі послуги одночасно (конфіденційність, цілісність, доступність, автентичність, корисність).

Крім того, удосконалений класифікатор загроз безпеці інформаційних ресурсів об'єктів критичної інфраструктури має можливість комплексування з методами соціальної інженерії та побудови багатоконтурних систем захисту інформації в соціокіберфізичних системах. [5,6]

Базуючись на поняттях синергізму та гібридності, реалізація фреймворку аналізу кіберзагроз, полягає в здійсненні наступних етапів [1,9]:

1. Формування класифікатора кіберзагроз. На даному етапі, зважаючи на те, що загрози інформаційній безпеці (ІБ), кібербезпеці (КБ), безпеці інформації (БІ) впливають на різні послуги безпеки (конфіденційність, цілісність, доступність,

автентичність, приналежність) з різною інтенсивністю, фахівцями з безпеки (експертами) вирішується завдання щодо нормування метричних коефіцієнтів загроз за послугами безпеки та формування класифікації загроз на основі розробленого класифікатора.

2. Аналіз загроз та ризиків. На даному етапі відбувається визначення реалізації кожної і-ї загрози з урахуванням імовірності прояву атаки її виникнення.

3. Синергія. На даному етапі визначається узагальнений показник рівня захищеності ресурсів об'єктів критичної інфраструктури (ОКІ). Для цього на основі сформованої множини загроз ІБ, КБ, БІ на безпеку інформаційних ресурсів визначається залежність між інформаційними активами ресурсів ОКІ і загрозами ІБ, КБ, БІ за такими напрямками:

- визначення зв'язку між інформаційними активами та послугами безпеки;
- визначення зв'язку між інформаційними активами та рівнем інфраструктури ISO/OSI.

Таким чином, результат всіх етапів дозволяє розробити корпоративні стратегії протидії загрозам та нейтралізації ризиків на основі розрахунку інтегральних показників ефективності за послугами безпеки, інформаційними ресурсами та елементами інфраструктури.

Список використаної літератури

1. Мілевський С. В. Розробка класифікатора загроз у соціокіберфізичних системах // *Ukrainian Scientific Journal of Information Security*, 2023. – vol. 29, issue 3. – pp. 118-123. DOI: 10.18372/2225-5036.29.18070
2. ENISA Threat Landscape 2021. – Режим доступу: <https://www.enisa.europa.eu/publications/enisa-threat-landscape-2021>
3. Hryshchuk R. The synergetic approach for providing bank information security: the problem formulation / R. Hryshchuk, S. Yevseiev // *Безпека інформації*, 2016. – № 22 (1). – С. 64 – 74. DOI: 10.18372/2225-5036.22.10456
4. IoT Security Maturity Model: Description and Intended Use. – Режим доступу: http://www.iiconsortium.org/pdf/SMM_Description_and_Intended_Use_2018-04-09.pdf
5. Milevskiy S., Korol O., Myktyyn G., Lozova I., Solnyshkova S., Husarova I., Hrebeniuk A., Vlasov A., Sukhoteplyi V., & Balagura D. Development of the sociocyberphysical systems' multi-contour security methodology // *Eastern-European Journal of Enterprise Technologies*, 2024. – №1(9). – pp. 34–51. DOI: org/10.15587/1729-4061.2024.298844
6. Models of socio-cyber-physical systems security: monograph / S. Yevseiev, Yu. Khokhlachova, S. Ostapov, O. Laptiev and others. – Kharkiv: PC TECHNOLOGY CENTER, 2023. – 168 p.
7. Pohasii S. and other. Development of conception for building a critical infrastructure facilities security system // *Eastern-European Journal of Enterprise Technologies*, 2024. – № 3/9 (111). – P. 63–83.
8. Shmatko O., Balakireva S., Vlasov A., Zagorodna N., Korol O., Milov O., Petrov O., Pohasii S., Rzayev Kh., Khvostenko V. Development of methodological foundations for a classifier of threats to cyberphysical systems design // *Eastern-European Journal of Enterprise Technologies*, 2020. – № 3/9 (105). – P. 6–19.
9. Synergy of building cybersecurity systems: monograph / S. Yevseiev, V. Ponomarenko, O. Laptiev, O. Milov and others. – Kharkiv: PC TECHNOLOGY CENTER, 2021. – 188 p.

СУЧАСНІ ПІДХОДИ ДО ЗАХИСТУ ГРАФІЧНОГО КОНТЕНТУ

Букатов Д.В.

*викладач циклової комісії математичних дисциплін
та інноваційного проектування*

Романенко О.І.

*викладач циклової комісії комп'ютерних наук
та програмної інженерії
ПВНЗ «Європейський університет»*

У контексті сучасних загроз для графічного контенту, зокрема масового використання зображень для навчання генеративних моделей без згоди авторів, технологія Nightshade стала революційною відповіддю художників та дослідників. Розроблена у 2023-2024 роках науковцями з University of Chicago, ця система реалізує концепцію data poisoning - цілеспрямованого отруєння даних, на яких навчаються моделі на кшталт Stable Diffusion або Midjourney [1].

Nightshade змінює пікселі зображення у спосіб, непомітний для людського ока, але критично важливий для глибоких нейромереж. Після навчання на такому зображенні модель формує хибні асоціації між візуальними ознаками та текстовими описами: наприклад, «кіт» може почати асоціюватися з образом змії або химерної істоти. У великих масштабах це призводить до деформації генеративної логіки ШІ, навіть у разі багаторазового застосування, - до повної деградації здатності створювати осмислені зображення [2].

Розглянемо далі основні особливості та застосування технології Nightshade.

1. Інтеграція з Glaze. Nightshade може поєднуватись з іншим інструментом, Glaze, що захищає саме стиль, тоді як Nightshade впливає на зміст. Nightshade технічно доповнює інший інструмент - Glaze, що був першим масштабним рішенням для маскування авторського стилю. Якщо Glaze викривлює стиль зображення на рівні глибоких ознак (deep features), змушуючи ШІ моделі інтерпретувати акварель як олійний живопис або замість стилю художника бачити «нейтральну шумову текстуру», то Nightshade впливає на змістовну інтерпретацію. Наприклад, kota можна подати як зображення сови або лиса для AI, зберігаючи при цьому зовнішню візуальну достовірність для глядача [3].

Ці два інструменти часто використовуються в парі: Glaze змінює «як виглядає» (стиль), а Nightshade – «що це є» (семантика). В результаті - подвійний рівень захисту: навіть якщо система обходить стилістичну маскування, вона зіштовхується зі спотвореним змістом, що знижує шанси ефективного навчання на даному зображенні [2].

2. Використовується художниками у відкритому доступі, зокрема, через інструменти самооборони цифрового мистецтва. Сотні тисяч користувачів вже застосували Nightshade для захисту своїх портфоліо [1].

3. Має мінімальний поріг впливу, тобто достатньо 50-100 зображень, щоб помітно спотворити результати генерації по певному запиту [3].

Отже, технологія Nightshade постає як форма цифрової «самооборони». Автори захищають своє право на недоторканність контенту, використовуючи

інструмент як превентивний щит проти нелегального збору даних. З юридичної точки зору, змінювати власний контент не заборонено, однак точаться дискусії, чи можна вважати масове отруєння навчальних наборів недобросовісною конкуренцією [2]. Хоча знову таки з іншого боку закон і досі не може однаково захистити авторів із різних країн [1].

Список використаної літератури

1. Sung, M. (2024). *Nightshade, the tool that 'poisons' data, gives artists a fighting chance against AI*. URL: <https://techcrunch.com/2024/01/26/nightshade-the-tool-that-poisons-data-gives-artists-a-fighting-chance-against-ai/>
2. Leffer, L. (2024). *Artists Are Slipping Anti-AI 'Poison' into Their Art. Here's How It Works*. URL: <https://www.scientificamerican.com/article/art-anti-ai-poison-heres-how-it-works/>
3. Miller, S. (2025). *Poisoning the machine*. URL: <https://mag.uchicago.edu/uchi-con>

КІБЕРГІГІЄНА ЯК ПЕРВИННИЙ ЕТАП ЗАХИСТУ В УМОВАХ ПОВНОМАСШТАБНОЇ ВІЙНИ

Бойко М.М.

*викладач кафедри комп'ютерних наук та програмної інженерії,
ПВНЗ «Європейський університет»*

В умовах повномасштабної війни кіберпростір перетворюється на поле бою. Кібератаки спрямовані не лише на військові або критичні інфраструктури, а й на кожного користувача цифрових технологій. У таких обставинах інформаційна безпека набуває стратегічного значення, а кібергігієна стає першим рубежем оборони [2]. Саме поведінка користувачів визначає вразливість систем, тому держава та організації мають не лише впроваджувати технічні засоби захисту, а й розвивати культуру відповідального цифрового поведіння.

Кібергігієна — це сукупність щоденних звичок і дій користувача, які спрямовані на профілактику інцидентів безпеки. За даними ENISA, навіть найкращі захисні технології не дають результату без належного рівня обізнаності користувачів [1]. Основні складові ефективної кібергігієни включають:

- сильні паролі та їх регулярне оновлення;
- використання менеджерів паролів;
- багатофакторну автентифікацію;
- обережність при роботі з електронною поштою, особливо у випадку підозрілих вкладень або посилань;
- своєчасне оновлення програмного забезпечення, включно з операційною системою, антивірусом та браузером;
- розуміння основ соціальної інженерії та потенційних загроз, зокрема фішингових технік, шахрайських повідомлень у месенджерах, використання підроблених сайтів [3].

Важливо не лише знати ці правила, а й дотримуватися їх систематично, перетворивши на звичку. Особливо критично це в умовах війни, коли ворог активно використовує кіберпростір як інструмент психологічного, економічного та інформаційного тиску.

Більшість успішних кібератак не є результатом прориву складних систем захисту, а виникають через помилки користувачів. Наприклад, відкриття шкідливого листа, встановлення невідомого ПЗ, чи передача облікових даних у відповідь на фішинговий запит. За даними аналітики SANS Institute, до 85% інцидентів можна було б уникнути, якби користувачі дотримувалися базових принципів кібергігієни [3].

Розвиток кібергігієни неможливий без освітньої підтримки. До найбільш ефективних форм належать:

- масові інформаційні кампанії, спрямовані на громадян;
- внутрішні корпоративні тренінги для працівників установ;
- симуляції фішингових атак, які допомагають перевірити готовність персоналу;
- використання ігрових підходів (гейміфікація), що підвищують мотивацію до навчання;
- регулярні оновлення знань у зв'язку з появою нових кіберзагроз [2].

Кожен користувач — це потенційна точка входу для кібератаки. Саме тому людська свідомість і пильність такі ж важливі, як і технічні бар'єри.

Таким чином, у період воєнного часу кібергігієна стає не просто інструментом особистої безпеки, а й елементом національної оборони. Її розвиток вимагає синергії держави, бізнесу, освітніх установ і самих громадян. З огляду на характер сучасних кіберзагроз, які часто мають приховану, але потенційно загрозливу природу, саме превентивні заходи виступають найефективнішим інструментом протидії. Тому масова просвіта, інформаційна грамотність і практичні навички кібергігієни мають бути в пріоритеті для всіх рівнів — від школи до держслужби. Тільки так можна зміцнити цифровий фронт України в умовах гібридної війни.

Список використаної літератури

1. ENISA. (2023). *Cyber Hygiene for All – Guidelines for Citizens*.
2. Наumenko, I.B. (2023). Кібергігієна в умовах воєнного стану. *Інформаційна безпека України*, №2.
3. SANS Institute. (2022). *Security Awareness Planning Kit*.

ПРОБЛЕМА НЕСАНКЦІОНОВАНОГО ДОСТУПУ ДО ІНФОРМАЦІЇ В ЦІЛЯХ НАВЧАННЯ ШТУЧНОГО ІНТЕЛЕКТУ

*Вах О.Ф.
аспірант,
Коцун В.І.*

*к. т. н., доц., завідувач кафедри математики та комп'ютерних дисциплін
Львівська філія ПВНЗ «Європейський університет»*

Роль штучного інтелекту в житті сучасної людини стає дедалі вагомішою, проникаючи практично в усі сфери нашої діяльності. На сьогоднішній день штучний інтелект активно інтегрується в побутових пристроях, медицині, моніторингу, навігаційних системах та багато іншого.

Водночас постають значні виклики, які перешкоджають подальшому розвитку штучного інтелекту. Однією з найкритичніших перешкод є дефіцит достатньої кількості даних. Великі мовні моделі (LLM) потребують величезних обсягів для вивчення шаблонів і зв'язків у мові. Без цього LLM не здатні ефективно виконувати різноманітні завдання обробки природної мови. Подолання цієї проблеми має вирішальне значення для подальшого вдосконалення та ширшого впровадження цих потужних моделей.

Можливим рішенням виступає аугментація даних, що полягає в створенні штучних варіацій даних для збільшення ефективного розміру та різноманітності навчального набору. Такий метод вимагає збільшення витрат на обчислювальні операції оскільки процес генерації зображень в достатній кількості є складним процесом [1]. Також генерація даних може призвести до деградації ефективності моделі.

Враховуючи все вище написане, великі корпорації почали змагатись за дані користувачів. Компанії використовують дані клієнтів для навчання своїх інструментів і моделей штучного інтелекту, не отримуючи явного дозволу.

Компанія Slack використовує файли, повідомлення та дані своїх користувачів для навчання своїх функцій штучного інтелекту. При чому, користувачі програмного забезпечення автоматично підключаються до цієї домовленості, без надання їхньої згоди.

Корпоративний гігант Meta навчає модель Llama 3 на основі постів, зображень користувачів на платформах Facebook та Instagram. Компанія не надає можливості користувачам відмовитись, спираючись на те, що використання цієї інформації відповідає законним інтересам Meta.

Stack Overflow нещодавно оголосила, що незабаром OpenAI зможе навчати свої моделі штучного інтелекту, використовуючи знання та відповіді, надані користувачами платформі за останні 15 років. Компанія стверджує, що контент не належить користувачам після публікації на платформі. Простіше кажучи, компанія робить неможливим для своїх користувачів відкликати дозвіл Stack Overflow на публікацію, зберігання, розповсюдження та використання такого вмісту.

Задля збереження приватності та захисту інформації слід дотримуватись наступних рекомендацій:

- не публікувати приватну інформацію в соцмережах та не використовувати її при запитах до моделей штучного інтелекту;
- не надавати дозвіл на використання особистих даних для навчання моделей (для цього слід змінити налаштування профілю);
- власникам веб-сайтів варто скористатись послугами компанії Cloudflare [2], яка запобігає скануванню веб-сайту без дозволу.

Таким чином, приватність користувачів є під загрозою, тому слід дотримуватись перелічених правил, щоб запобігти витоку персональних даних.

Список використаної літератури

1. A. Mikołajczyk and M. Grochowski. (2018). *Data augmentation for improving deep learning in image classification problem. International Interdisciplinary PhD Workshop (IIPHDW), Świnouście, Poland, pp. 117-122, doi: 10.1109/IIPHDW.2018.8388338.*
2. Bajrami, Enes and Idrizi, Florim and RUSHITI, Agim (2024) *ENHANCING WEBSITE SPEED AND SECURITY WITH CLOUDFLARE CDN: A CASE STUDY ON WORDPRESS WEBSITES. JNSM Journal of Natural Sciences and Mathematics of UT, 9 (17-18). pp. 294-304. ISSN 2671-3039.*

ФОРМУВАННЯ ЦИФРОВОЇ СТІЙКОСТІ ЯК КРИТИЧНОГО ЕЛЕМЕНТУ ДОВІРИ ДО СУБ'ЄКТІВ ГОСПОДАРЮВАННЯ В УМОВАХ ЕСКАЛАЦІЇ КІБЕРЗАГРОЗ

*Ващенко М.В.
викладач кафедри кібербезпеки та захисту інформації
ПВНЗ «Європейський університет»*

Ескалація кіберзагроз у світі, спричинена активізацією гібридних атак, суттєво змінила вимоги до стійкості суб'єктів господарювання. У сучасних умовах цифрова стійкість організацій розглядається не лише як технічний аспект захисту, але й як критичний фактор формування довіри серед клієнтів, партнерів та інвесторів [1].

Довіра у цифровому середовищі сьогодні визначається здатністю організації забезпечувати безпеку інформаційних систем, прозоро реагувати на інциденти та активно комунікувати про свою політику безпеки [2]. Згідно з даними World Economic Forum, цифрова стійкість входить до трійки головних критеріїв довіри до бізнесу у 2024 році [3]. Відсутність підтвердженої стійкості стає бар'єром для залучення інвестицій і укладення міжнародних угод.

Формування цифрової стійкості включає не лише дотримання загальноновизнаних стандартів кібербезпеки, таких як принципи, викладені в рекомендаціях OECD щодо цифрової безпеки [4], а й впровадження практик оцінки ризиків у ланцюгах постачання, розвиток культури безперервного вдосконалення безпеки та побудову процесів обміну інформацією про кіберзагрози. Ініціатива CISA «Shields Up» підкреслює важливість активного моніторингу ризиків навіть у «мирний» час [5].

Важливу роль у формуванні довіри відіграє також відкритість у питаннях сертифікації безпеки. Європейський акт про кібербезпеку встановлює рамки для

сертифікації послуг та продуктів, що дозволяє об'єктивно оцінювати цифрову стійкість підприємств [6]. Компанії, які проходять подібну сертифікацію, демонструють готовність відповідати найвищим вимогам безпеки на глобальному ринку.

Аналітика Microsoft у своєму Digital Defense Report підкреслює, що саме репутація у сфері кібербезпеки стала критичним фактором вибору бізнес-партнерів у постпандемічному та воєнному світі [7]. Водночас організаційна стійкість має включати не лише технічні рішення, але й здатність до адаптації, стратегічного управління ризиками та відновлення після інцидентів, як це підкреслюється в Індексі організаційної резильєнтності BSI [8].

Таким чином, цифрова стійкість стає ключовим компонентом довіри до суб'єктів господарювання в умовах ескалації кіберзагроз. Її формування потребує комплексного підходу, що включає технічні, організаційні та комунікаційні аспекти, спрямовані на забезпечення безперервності бізнесу, зміцнення репутації та підтримку інтеграції у глобальну економіку.

Список використаної літератури

1. ENISA. ENISA Threat Landscape 2023. European Union Agency for Cybersecurity. URL: <https://www.enisa.europa.eu/publications/enisa-threat-landscape-2023>
2. PwC. Global Digital Trust Insights 2024. PricewaterhouseCoopers. URL: <https://www.pwc.com/gx/en/services/consulting/cybersecurity/global-digital-trust-insights.html>
3. World Economic Forum. Global Risks Report 2024. URL: <https://www.weforum.org/publications/global-risks-report-2024/>
4. OECD. Recommendation on Digital Security Risk Management for Economic and Social Prosperity, 2021. URL: <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0456>
5. CISA. Shields Up Campaign, 2023. URL: <https://www.cisa.gov/shields-up>
6. ENISA. EU Cybersecurity Act Overview. URL: <https://www.enisa.europa.eu/topics/certification/cybersecurity-act>
7. Microsoft. Digital Defense Report 2023. URL: <https://www.microsoft.com/en-us/security/business/security-intelligence-report>
8. BSI. Organizational Resilience Index 2023. British Standards Institution. URL: <https://www.bsigroup.com/en-GB/our-services/organizational-resilience/organizational-resilience-index/>

КІБЕРБЕЗПЕКА В JAVA-ДОДАТКАХ: МОЖЛИВОСТІ ІНТЕГРАЦІЇ ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ ЗАХИСТУ СИСТЕМ

*Венжик А.О.
магістрант факультету інформаційних систем та технологій,
Ковальчук М.А.,
аспірант,
ПВНЗ «Європейський університет»*

У сучасному цифровому середовищі мова програмування Java посідає ключове місце у розробці масштабованих, надійних і критично важливих програмних систем. За даними компанії Oracle [4], Java експлуатується на понад 3 мільярдах пристроїв у світі, що підкреслює необхідність посиленої уваги до аспектів

її безпечного використання. Така широка поширеність обумовлює потребу у впровадженні сучасних підходів до захисту Java-орієнтованих середовищ, зокрема в контексті мобільних, вбудованих та хмарних рішень.

Кіберзагрози, з якими стикаються Java-додатки, включають не лише традиційні атаки (SQL-ін'єкції, XSS, CSRF), але й нові, пов'язані з маніпуляцією трафіком, викраденням сесій, ін'єкцією через API тощо [2]. У звіті IBM зазначено: «Штучний інтелект здатен скоротити середній час виявлення загрози з 280 до 45 хвилин» [3], що підкреслює актуальність його застосування в системах кіберзахисту.

У Java-середовищі існує низка рішень для інтеграції машинного навчання. Наприклад, Weka — популярна бібліотека для створення моделей класифікації, яка легко вбудовується в проекти через стандартні Java API. Її можна використовувати для виявлення аномалій у вебтрафіку:

$$x_i \in \mathbb{R}^n \rightarrow f(x_i) = \{ 0 \text{ — нормальний; } 1 \text{ — аномальний} \}$$

Такі моделі натреновані на попередніх запитах і можуть в реальному часі відхиляти потенційно шкідливі запити. У роботі Hall et al [5] зазначено, що Weka дозволяє реалізовувати понад 100 алгоритмів класифікації, що відкриває широкі можливості в контексті аналізу безпеки.

Ще один приклад — DeepLearning4j, який підтримує роботу з нейронними мережами безпосередньо у JVM. Така інтеграція дозволяє створювати самонавчальні агенти, що аналізують поведінку користувача та виявляють підозрілі патерни.

Захист пристроїв Інтернету речей (IoT), що функціонують на основі Java ME, є важливим напрямом застосування сучасних засобів кібербезпеки. Згідно з дослідженнями компанії Konduit AI [1], штучний інтелект може бути інтегрований навіть у малопотужні пристрої, що забезпечує можливість локального виявлення та реагування на загрози. Такий підхід є особливо актуальним для критичної інфраструктури в сферах транспорту, енергетики та телекомунікацій.

Особливу увагу слід приділити концепції Zero Trust, яка ґрунтується на принципі повної недовіри до всіх запитів незалежно від їхнього походження. Як зазначено в офіційній документації Google Cloud, модель Zero Trust не спирається на традиційний мережевий периметр, натомість здійснює індивідуальну перевірку кожної сесії, що забезпечує підвищений рівень контролю доступу та зменшує ризики несанкціонованого проникнення [7]. Для Java-додатків це означає потребу в гнучких механізмах автентифікації та поведінковому моніторингу, що знову ж таки може бути реалізовано через AI-моделі [3].

Крім того, Java підтримує розширені механізми шифрування, серед яких AES, RSA, TLS. Проте навіть ці методи можуть бути уразливими без контекстної оцінки ризиків. Штучний інтелект тут застосовується для аналізу активності, виявлення підозрілих спроб доступу до зашифрованих даних.

Слід наголосити, що сучасні інтегровані середовища розробки, зокрема IntelliJ IDEA, уже на сьогодні активно впроваджують інструменти на основі штучного інтелекту. Згідно з офіційним джерелом компанії JetBrains, AI-асистенти здатні автоматично виявляти і усувати потенційні вразливості в програмному коді ще на

етапі розробки, що сприяє підвищенню безпеки програмного забезпечення до його виконання [6].

Отже, поєднання Java з сучасними інструментами штучного інтелекту створює потужну платформу для побудови надійних, адаптивних систем захисту. Такий підхід має особливе значення для кібербезпеки України, яка продовжує боронити свій цифровий фронт в умовах гібридної війни.

Список використаної літератури

1. *Konduit AI. (2023). Deeplearning4j Documentation. URL: <https://deeplearning4j.konduit.ai/>*
2. *Java Security Architecture. Oracle Documentation. URL: <https://docs.oracle.com/javase/8/docs/technotes/guides/security/spec/security-spec.doc1.html>*
3. *IBM. Using AI to enhance cybersecurity. URL: <https://www.ibm.com/security/artificial-intelligence>*
4. *Secure Coding Guidelines for Java SE (Oracle). URL: <https://www.oracle.com/java/technologies/javase/seccodeguide.html>*
5. *Weka. URL: <https://waikato.github.io/weka-wiki/documentation/>*
6. *JetBrains. AI Assistant Features in IntelliJ IDEA. URL: <https://blog.jetbrains.com/idea/2024/01/ai-assistant>*
7. *Google Cloud. Zero Trust Whitepaper. URL: <https://cloud.google.com/zero-trust>*

ОРГАНІЗАЦІЙНІ АСПЕКТИ ВПРОВАДЖЕННЯ МОДЕЛІ ZERO TRUST У КОРПОРАТИВНОМУ СЕРЕДОВИЩІ

Вілянський А.В.,

аспірант,

Позняк А.А.

*викладач Фахового бізнес-коледжу
ПВНЗ «Європейський університет»*

У сучасних умовах постійного зростання кіберзагроз традиційна модель периметрального захисту втрачає свою ефективність [2, 3]. Замість довіри до внутрішніх користувачів або пристроїв, концепція Zero Trust (ZT) передбачає повну відмову від презумпції безпечного середовища та формування стратегії «нікому не довіряй, завжди перевіряй». Впровадження моделі Zero Trust у корпоративному секторі є не лише технічним завданням, але й вимагає адаптації до правових норм та адміністративних процедур. Окреслимо ключові аспекти застосування ZT у корпоративному середовищі з акцентом на нормативно-правові та організаційні виклики.

Модель Zero Trust ґрунтується на кількох базових принципах: багатофакторна автентифікація (MFA), мінімізація прав доступу (least privilege), мікросегментація мережі, постійний моніторинг активності, а також перевірка кожного запиту незалежно від розташування користувача чи пристрою. Zero Trust — це не окремий продукт, а стратегічна рамка для побудови захищеного IT-середовища.

Дана модель стала особливо актуальною після пандемії COVID-19 та масового переходу на віддалену роботу, коли периметр організації розширився далеко за межі офісу. Провідні стандарти, зокрема, NIST SP 800-207, а також рекомендації ENISA та Європейської комісії вже визначають ZT як сучасну основу кіберзахисту [1].

Для впровадження Zero Trust у корпоративному середовищі в Україні важливо враховувати норми таких законодавчих актів, як Закон України «Про основні засади забезпечення кібербезпеки» (2017), Закон «Про захист персональних даних», GDPR (у разі обробки даних громадян ЄС), Директива NIS2 (у перспективі її імплементації до українського законодавства).

Всі ці документи прямо або опосередковано вимагають забезпечення належного рівня інформаційної безпеки, фіксації політик доступу, логування дій користувачів та звітності перед державними регуляторами. Однією з головних вимог, що корелює з філософією Zero Trust, є принцип accountability — необхідність фіксувати, хто саме і коли отримував доступ до конкретних інформаційних ресурсів.

Застосування моделі ZT у компанії передбачає збір і обробку значного обсягу поведінкових та технічних даних користувачів. Це, у свою чергу, вимагає юридичного обґрунтування обробки персональних даних, визначення цілей обробки та відповідності принципам прозорості й пропорційності, особливо при використанні SIEM-систем, аналізу трафіку та поведінкової аналітики.

З точки зору корпоративного управління, впровадження Zero Trust вимагає:

- формування політик контролю доступу на всіх рівнях (IdP, файлові системи, корпоративні застосунки);
- оновлення інструкцій з інформаційної безпеки, підготовки персоналу;
- забезпечення кадрового потенціалу — включно з призначенням CISO або іншої відповідальної особи;
- побудови системи внутрішнього аудиту, яка відповідатиме критеріям ISO/IEC 27001.

Ключовим бар'єром є відсутність у багатьох українських компаніях системного підходу до інформаційної безпеки, що проявляється у фрагментарності впроваджених рішень, недостатній комунікації між IT та юристами, а також у відсутності механізмів звітування про кіберінциденти. Адміністративні механізми, що підтримують ZT, включають ревізію доступів, регулярну оцінку ризиків, побудову системи інцидент-менеджменту та контроль аутсорсингу, особливо у хмарних рішеннях.

Отже, впровадження моделі Zero Trust у корпоративному середовищі — це не лише технологічне оновлення, а й трансформація підходу до управління ризиками, відповідальності та правової відповідності. Успішна реалізація вимагає міждисциплінарної співпраці технічних, юридичних та управлінських команд. З огляду на глобальні тренди, включно з імплементацією NIS2, та зростаючу кількість атак на українські компанії, інтеграція принципів ZT повинна стати пріоритетом для підприємств усіх масштабів.

Список використаної літератури

1. *National Institute of Standards and Technology. Zero Trust Architecture : NIST Special Publication 800-207 / J. Kindervag, S. Rose. – Gaithersburg, MD : NIST, 2020. – 59 p. – Режим доступу: <https://doi.org/10.6028/NIST.SP.800-207>*
2. Яровий, Р., Улічев, О., Скляренко, О., Пашорін, В. (2024). Моделювання мультиагентних систем захисту інформаційних ресурсів. *Herald of Khmelnytskyi National University. Technical Sciences*, 337(3(2)), 278-284. <https://doi.org/10.31891/2307-5732-2024-337-3-42>
3. Lakhno, V.A., Kasatkin, D.Y., Skliarenko, O.V., Kolodinska, Y.O. *Modeling and Optimization of Discrete Evolutionary Systems of Information Security Management in a Random Environment // Machine Learning and Autonomous Systems. Smart Innovation, Systems and Technologies*, vol 269. Springer, Singapore. – 2022 – p. 9-22. https://doi.org/10.1007/978-981-16-7996-4_2

РОЗВИТОК НАВИЧОК КРИТИЧНОГО МИСЛЕННЯ ТА КІБЕРБЕЗПЕКИ У ШКОЛЯРІВ ЯК ПРОТИДІЯ ІНФОРМАЦІЙНО- ПСИХОЛОГІЧНИМ МАНІПУЛЯЦІЯМ ТА ВПЛИВАМ

*Волкова Н.М.,
старший викладач кафедри математичних дисциплін
та інноваційного проектування
ПВНЗ “Європейський університет”*

Сучасний інформаційний простір характеризується стрімким зростанням обсягів інформації, її різноманітністю та доступністю. Водночас, це створює нові виклики для молодого покоління, зокрема, зростання ризиків інформаційно-психологічних маніпуляцій та деструктивних впливів. Школярі, як активні користувачі цифрового середовища, особливо вразливі до подібних загроз. У зв'язку з цим, набуває особливої актуальності питання розвитку у них критичного мислення та навичок кібербезпеки як ключових інструментів протидії цим негативним явищам. Задля їх впровадження необхідно об'єднати зусилля науковців, педагогів, представників державних установ та громадських організацій для обговорення актуальних проблем, обміну досвідом та визначення перспективних шляхів розвитку критичного мислення та кібербезпеки у школярів в контексті протидії інформаційно-психологічним маніпуляціям та впливам.

Розглянемо наступні теоретико-методологічні засади розвитку критичного мислення та кібербезпеки:

- критичне мислення як ключова компетенція інформаційного суспільства, де складовими є аналіз, оцінка, інтерпретація, умовиводи, пояснення, саморегуляція;
- роль критичного мислення у формуванні інформаційної грамотності та медіаграмотності;
- взаємозв'язок критичного мислення з іншими ключовими компетентностями, такими як комунікація, співпраця, креативність.

Щодо представлення кібербезпеки як складової особистої та національної безпеки, можна казати про наступні складові - визначення кібербезпеки та її основні аспекти: конфіденційність, цілісність та доступність інформації.

Основними загрозами кібербезпеки для школярів можна вважати фішинг, шахрайство, кібербулінг, шкідливе програмне забезпечення, небезпечний контент. Також важливо розглядати правові та етичні аспекти кібербезпеки. При визначенні інформаційно-психологічних маніпуляцій та їхніх основних методах можна говорити про фейкові новини, пропаганду, емоційний вплив, когнітивні викривлення. Важливим є розуміння психологічних механізмів сприйняття та протидії маніпуляціям, а також впливу інформаційно-психологічних операцій на особистість, суспільство та національну безпеку.

Розглянемо сучасний стан та виклики розвитку критичного мислення та кібербезпеки у школярів України. При аналізі навчальних програм та освітніх практик можемо казати про важливе місце та роль розвитку критичного мислення та кібербезпеки в чинних освітніх стандартах та навчальних програмах різних предметів, ефективність існуючих методик та технологій навчання, проблеми інтеграції компетентностей критичного мислення та кібербезпеки в освітній процес.

Якщо розглядати рівень сформованості навичок критичного мислення та кібербезпеки серед українських школярів, то можна виділити такі характеристичні показники як:

- результати національних та міжнародних досліджень щодо інформаційної та медіаграмотності, кібербезпеки серед молоді;
- виявлення типових помилок та вразливостей школярів в інформаційному просторі;
- аналіз впливу соціальних мереж та онлайн-платформ на формування інформаційної поведінки учнів.

Якщо казати про виклики та перешкоди на шляху розвитку критичного мислення та кібербезпеки можна перелічити наступні фактори, як:

- недостатня методична забезпеченість педагогів;
- обмеженість навчальних ресурсів та технологій;
- низька мотивація учнів та батьків до розвитку цих навичок;
- вплив дезінформації та пропаганди на освітній процес.

Розглянемо педагогічні стратегії та інноваційні підходи до розвитку критичного мислення та кібербезпеки.

1. Інтеграція критичного мислення в навчальні предмети [2]:

- методи активного навчання: проблемне навчання, кейс-стаді, дебати, дискусії, тощо;
- розвиток навичок аналізу джерел інформації, виявлення упереджень та маніпуляцій;
- формування вміння ставити запитання, генерувати аргументи та робити обґрунтовані висновки.

2. Формування навичок кібербезпечної поведінки:

- ознайомлення з основними загрозами та правилами безпечної роботи в онлайн-середовищі;
- розвиток вміння розпізнавати фішингові атаки, шкідливе програмне забезпечення та недостовірну інформацію;
- навчання безпечному використанню соціальних мереж та онлайн-платформ;

- формування відповідального ставлення до власної цифрової репутації та персональних даних.

3. Використання інформаційно-комунікаційних технологій (ІКТ) для розвитку критичного мислення та кібербезпеки:

- застосування інтерактивних платформ, онлайн-симуляторів, ігор та віртуальних навчальних середовищ [1];
- розвиток цифрової грамотності та медіакомпетентності за допомогою ІКТ.
- використання онлайн-ресурсів для самостійного навчання та підвищення обізнаності.

Важливою є міжсекторальна співпраця та перспективи розвитку та її складові, такі як:

- роль сім'ї, школи, держави та громадських організацій у розвитку критичного мислення та кібербезпеки;
- формування безпечного та сприятливого інформаційного середовища для дітей та молоді;
- підвищення обізнаності батьків щодо ризиків в онлайн-просторі та методів їхньої профілактики;
- розробка та впровадження національних стратегій та програм з розвитку критичного мислення та кібербезпеки;
- залучення громадських організацій та експертів до освітніх ініціатив.

Розглядаючи перспективи подальших досліджень та розробок, можемо відзначити наступні:

- вивчення впливу новітніх технологій (-штучний інтелект, метавсесвіт) на інформаційну безпеку та когнітивні процеси молоді.
- розробка інноваційних методик оцінювання рівня сформованості критичного мислення та кібербезпеки;
- створення ефективних інструментів та ресурсів для навчання педагогів та учнів;
- врахування міжнародного досвіду та можливості його адаптації в Україні.

Підсумовуючи все, що було сказано вище, можемо зробити наступні висновки. Розвиток критичного мислення та навичок кібербезпеки є нагальною потребою сучасного українського суспільства та важливим елементом національної безпеки. Ефективна протидія інформаційно-психологічним маніпуляціям та впливам вимагає комплексного підходу, що включає модернізацію освітніх програм, впровадження інноваційних педагогічних стратегій, активну міжсекторальну співпрацю та постійне наукове дослідження. Результати конференції сприятимуть консолідації зусиль науково-педагогічної спільноти, державних органів та громадських організацій для формування покоління молодих громадян, здатних критично оцінювати інформацію, безпечно діяти в цифровому середовищі та протистояти деструктивним впливам.

Список використаної літератури

1. Склярєнко О.В., Ягодзінський С.М., Ніколаєвський О.Ю., Невзоров А.В. Цифрові інтерактивні технології навчання як невід’ємна складова сучасного освітнього процесу // Науковий журнал «Інноваційна педагогіка», Випуск 68, Т.2. – 2024. - с. 250-257.
2. Дуброва О.М. Розвиток навичок критичного мислення сучасних учнів та студентів – представників покоління Z: інтерактивні практики. УДК 378:37.01 DOI <https://doi.org/10.32782/2663-6085/2024/72.9>
<http://www.innovpedagogy.od.ua/archives/2024/72/11.pdf>
3. Курс із кібербезпеки для школярів і методичні посібники до нього. Медіа “Нова українська школа”. <https://nus.org.ua/2023/04/26/kurs-iz-kiberbezpeky-dlya-shkolyariv-i-metodychni-posibnyky-do-nogo-dos-tupni-dlya-vsih/>
4. Стратегія інформаційної безпеки до 2025 року. Урядовий портал органів виконавчої влади України. <https://www.kmu.gov.ua/news/uryad-shvaliv-strategiyu-informacijnoyi-bezpeki-do-2025-roku>
5. Тезисна інструкція на тему кібербезпеки школярів. Бібліотека порталу «Всеосвіта». <https://vseosvita.ua/library/tezysna-instruktsiya-na-temu-kiberbezpeky-shkolyariv-560631.html>

АНАЛІЗ ЕФЕКТИВНОСТІ СИМЕТРИЧНИХ ТА АСИМЕТРИЧНИХ МЕТОДІВ ШИФРУВАННЯ У VPN

*Галушка П. А.
бакалавр
студент кафедри безпеки інформаційних технологій
Коробейнікова Т. І.
к.т.н., доцент
доцент кафедри безпеки інформаційних технологій
Національний університет «Львівська політехніка»*

Використання мережі для зберігання та передавання конфіденційних даних вимагає захищених каналів зв’язку. У VPN-технологіях для цього застосовують криптографічні алгоритми, які забезпечують конфіденційність і цілісність інформації. Основними методами є симетричне шифрування (DES, 3DES, AES), яке характеризується високою швидкістю, та асиметричне шифрування (RSA, ECC), яке забезпечує безпечний обмін ключами.

Метою дослідження є аналіз ефективності симетричних та асиметричних методів шифрування, які використовуються у VPN, з оцінкою їх продуктивності та рівня безпеки для вибору оптимальних рішень залежно від типу мережевих задач.

Однією з ключових сфер застосування криптографічних алгоритмів є забезпечення безпеки у віртуальних приватних мережах (VPN). VPN використовують криптографічні методи для створення захищених каналів зв’язку між користувачами та серверами через незахищені мережі, такі як Інтернет. У більшості сучасних VPN-рішень використовується комбінація симетричного та асиметричного шифрування: асиметричні алгоритми застосовуються для захищеного

обміну ключами, після чого основне шифрування трафіку відбувається за допомогою симетричних методів.

Ефективність криптографічного захисту у VPN залежить від вибору алгоритму. Симетричні методи забезпечують швидке шифрування і підходять для передачі великих обсягів даних, тоді як асиметричні гарантують безпечну автентифікацію та розподіл ключів. Для кращого розуміння відмінностей між цими підходами у контексті VPN доцільно розглянути порівняльну характеристику їх основних властивостей у таблиці 1.

Таблиця 1.

Порівняння симетричного та асиметричного шифрування

Симетричне шифрування	Асиметричне шифрування
Один ключ використовується для шифрування і дешифрування даних.	Пара ключів використовується для шифрування і дешифрування. Ці ключі відомі як «відкритий ключ» і «закритий ключ».
Простий метод шифрування, так як використовується тільки один ключ.	У зв'язку з тим, що використовується пара ключів – процес складний.
Використовується для шифрування великих об'ємів даних.	Забезпечує автентифікацію.
Забезпечує високу продуктивність і вимагає менше обчислювальної потужності.	Складні процеси протікають повільніше і вимагають більшої обчислювальної потужності.
Для шифрування даних використовується менша довжина ключа (128-256 біт).	Використовуються довші ключі шифрування (1024-4096 біт).
Ідеально підходить для шифрування великої кількості даних.	Використовується при шифруванні невеликого об'єму даних.
Стандартні алгоритми: RC4, AES, DES, 3DES і QUAD.	Стандартні алгоритми: RSA, Diffie-Hellman, ECC, El Gamal і DSA.

Джерело: авторська розробка на основі [1-3]

Авторами досліджено застосування симетричних та асиметричних алгоритмів шифрування у VPN для забезпечення конфіденційності та цілісності переданих даних. Визначено, що поєднання цих методів дозволяє досягти оптимального балансу між продуктивністю та безпекою. Отримані результати сприяють подальшому розвитку технологій VPN та вдосконаленню методів криптографічного захисту інформації у сучасних мережах.

Список використаної літератури

1. H25.io. (2024). *RSA vs AES Encryption: Understanding Symmetric vs Asymmetric Choices*. H25 Cybersecurity.
2. Kuchuk, N., Kovalenko, A., Tkachov, V., Rosinskiy, D., & Kuchuk, H. (2021). *Predicting traffic anomalies in container virtualization*. Computer And Information Systems And Technologies.
3. TechTarget. (2024). *Symmetric vs. asymmetric encryption: What's the difference?* TechTarget.

ОЦІНКА ЕФЕКТИВНОСТІ ЗАХИСТУ API У ГІБРИДНОМУ СЕРЕДОВИЩІ МОБІЛЬНОГО СЕРВІСУ

*Гетьман О. Л.,
аспірант,*

ПВНЗ «Європейський університет»

За умов активного поширення мобільних сервісів питання захисту прикладного програмного інтерфейсу (Application Programming Interface, API) набуває критичної актуальності. На сьогоднішній день відповідні інтерфейси у мобільному середовищі виступають не лише засобом інтеграції компонентів, але й ціллю для кібератак, що ускладнює задачу забезпечення конфіденційності, цілісності та доступності даних [1]. У рамках даного дослідження пропонується розглядати API як формалізовану модель взаємодії, яка включає набір протоколів та інструментів розробки. Особливу увагу при цьому має бути приділено ключовим аспектам безпеки, як то контроль доступу, аутентифікація та авторизація, моніторинг, захист від кібератак, а також впровадження засобів аудиту та автоматизації захисних функцій. Такий підхід дозволяє знизити ризики несанкціонованого доступу та підвищити стійкість мобільних додатків у гібридному середовищі.

Першим етапом дослідження стало узагальнення та систематизація ключових напрямків захисту API, що дозволяє не лише визначити вразливі ділянки мобільного сервісу, але й сформулювати основу для побудови ефективної архітектури системи захисту. У межах запропонованої комплексної методики виділено наступні категорії [1, 2].

1. Методи аутентифікації та авторизації, що включають у себе керування доступом на основі ролей і атрибутів, а також захист облікових даних з використанням додатку типу «CAPTCHA».
2. Засоби захищеної передачі даних, що передбачають використання шифрування, протоколів TLS і SSL, перевірку сертифікатів та піннінг сертифікатів, а також захист від атак типу MitM (Man-in-the-Middle).
3. Захист інформаційного сховища, що охоплює шифрування даних на клієнтському та серверному боці, використання серверів керування ключами, ротацію ключів та апаратні модулі безпеки.
4. Протидія кібератакам, що включає у себе лімітування запитів, налаштування фаєрволів веб додатків (Web Application Firewall, WAF), балансування

навантаження, а також фільтрацію шкідливого трафіку за допомогою інтелектуальних систем моніторингу.

5. Визначення специфічних аспектів функціонування мобільних додатків, що передбачає обфускацію програмного коду, контейнеризацію та засоби автоматизованого реагування на інциденти.

6. Моніторинг та аудит, що включає у себе ведення журналів подій, автоматизовану аналітику аномалій, перевірку відповідності стандартам та вимогам.

7. Тестування системи безпеки, що охоплює тестування на проникнення (Penetration Testing, PT), автоматизоване сканування вразливостей і регресійне тестування після оновлень.

8. Оновлення та управління версіями, що передбачає контроль життєвого циклу API, впровадження оновлень і патчів, а також забезпечення зворотної сумісності між версіями ключових складових API.

Комплексний аналіз засобів захисту API, проведений у рамках першого етапу дослідження, став основою для подальшого моделювання та виявлення загроз у реальному часі. Ефективне виявлення зовнішніх загроз потребує створення навчального набору даних на основі актуальних типів кібератак і формування класифікації загроз як основи знань для ідентифікації характерних ознак. У відповідності до цього можуть бути сформовані методичні рекомендації щодо впровадження нейромережових алгоритмів із урахуванням специфіки задач, протоколів безпеки та обмежень на обчислювальний ресурс апаратно-програмної платформи. Класифікація, побудована за результатами аналізу даних, дозволила структурувати інформацію про поширені типи кібератак та типові для API вразливості [2, 3] таким чином.

1. Шкідливе програмне забезпечення (Malware, MW). Зазначена категорія охоплює віруси, троянські програми і шпигунські програмні модулі, що використовуються для несанкціонованого доступу, порушення роботи системи або прихованого використання ресурсів. Основними заходами протидії є застосування антивірусних рішень, оновлення програмного забезпечення та впровадження систем виявлення і запобігання вторгненням (IDS/IPS).

2. Отримання доступу до облікових записів (Account Hijacking, AH). Зазначена категорія включає у себе фішинг, підбір паролів та інші методи отримання несанкціонованого доступу до акаунтів користувачів. Засобами протидії є багатофакторна аутентифікація, налаштування політики управління паролями та підвищення обізнаності користувачів.

3. Цільові атаки (Targeted Attacks, TA). Зазначена категорія включає у себе клас атак, що спрямовані на конкретну програмну систему або платформу з метою викрадення даних чи порушення функціональності. Протидія передбачає регулярні аудити безпеки, сегментацію даних, а також активний моніторинг мережевої активності.

4. Ін'єкції шкідливих скриптів (Malicious Script Injection, MSI). Зазначена категорія включає у себе зловмисне впровадження коду, що виконується в межах додатків або операційної системи (ОС). Захист забезпечується шляхом валідації введення, оновлення компонентів серверної інфраструктури та застосування політик контролю контенту (Content Security Policy, CSP).

5. Атаки типу DDoS (Distributed Denial of Service). Зазначена категорія включає у себе клас атак, що спрямовані на перевантаження системи запитами. Ефективна протидія базується на обмеженні частоти запитів, балансуванні навантаження, а також резервному копіюванні даних і контролю доступу до адміністративних інтерфейсів.



Рис. 1. Динаміка процентного внеску загроз API протягом 2014-2022 рр.

Проведений аналіз динаміки зовнішніх загроз API протягом останніх двох десятиріч дозволив виявити стійке зростання частки атак, пов'язаних із впровадженням зловмисного програмного забезпечення (рис. 1). Така тенденція обумовлена підвищенням пропускну здатності мережевих інфраструктур, що сприяє швидкому розповсюдженню шкідливих програм, а також низьким рівнем цифрової грамотності значної частини користувачів. Останнє зумовлює недостатній рівень впровадження базових заходів кіберзахисту, зокрема у масовому сегменті споживачів мобільних сервісів. Натомість спостерігається зниження частки таких загроз, як цільові атаки, впровадження кодових послідовностей, атаки розподіленого доступу та «дефейс». Це може свідчити не стільки про зменшення активності зловмисників у цих напрямках, скільки про зростання частки атак, природа яких не вкладається у відомі класифікації. Така ситуація вказує на актуальність задачі ідентифікації невідомих загроз, які можуть проявлятися як у вигляді нових категорій, що не належать до жодної з усталених типологій, так і у вигляді гібридних атак, що поєднують ознаки декількох відомих класів, але не представлені у навчальних вибірках. Розробка багаторівневої класифікації кібератак дозволяє стандартизувати підходи до оцінки ефективності методів виявлення аномалій у мережевому трафіку.

Список використаної літератури

1. Ronghua, S, Qianxun, W, Liang G. (2022). Cyber Security: Research Towards Key Issues of API Security. Communications in Computer and Information Science. DOI:10.1007/978-981-16-9229-1_11.

2. Phuc, T. S., & Lee, C. (2017). New security approaches for SSL/TLS attacks resistance in practice. The Journal of Society for E-Business Studies, 22(2), 169–185. DOI:10.7838/jsebs.2017.22.2.169.

3. Al-Zewairi, M., Almajali, S. & Ayyash, M. (2020). Unknown security attack detection using shallow and deep ANN classifiers. Electronics, 2020 (9), 1-27, DOI:10.3390/electronics 9122006.

АВТОМАТИЗАЦІЯ ТЕСТУВАННЯ ЗАСТОСУНКІВ ЗА ДОПОМОГОЮ SELENIUM В КОНТЕКСТІ КІБЕРБЕЗПЕКИ УКРАЇНИ

*Гордон-Басанський Ю.В.,
магістрант факультету інформаційних систем та технологій,
ПВНЗ «Європейський університет»*

У сучасних умовах, коли Україна стикається з масштабними кіберзагрозами в період повномасштабної війни, забезпечення надійності та безпеки програмного забезпечення стає критично важливим. Традиційні методи ручного тестування не встигають за швидким темпом розробки й викриття нових вразливостей, тому автоматизація тестового процесу, зокрема за допомогою фреймворку Selenium, здобуває особливу актуальність.

Selenium — це відкритий інструмент для автоматизації веб-браузерів із підтримкою Java, Python, C# та інших мов програмування. Він дозволяє створювати скрипти, що імітують дії користувача (клік, введення тексту, навігацію) і перевіряти функціональність веб-застосунків у різних браузерах. Завдяки модульній архітектурі та підтримці WebDriver, тести можна запускати в розподіленому середовищі, наприклад у Docker-контейнерах чи хмарних CI/CD-пайплайнах (GitLab CI, Jenkins), що пришвидшує регресійне тестування [1].

У контексті кібербезпеки інтеграція Selenium з OWASP ZAP (Zed Attack Proxy) відкриває можливість одночасного виконання функціональних і безпекових перевірок. Достатньо налаштувати проксі-сервер ZAP у тестових скриптах Selenium, щоб перехоплювати HTTP/HTTPS-трафік і запускати сканування на вразливості: SQL-ін'єкції, міжсайтове скриптування (XSS), недостатню перевірку вхідних даних тощо. Такий підхід дає змогу виявити проблемні місця на ранніх етапах розробки й миттєво повідомити команді про критичні дефекти [2].

Останні роки показали, що застосування штучного інтелекту суттєво підвищує ефективність автоматизованого тестування. Інструменти на основі AI можуть аналізувати історію виконання тестів, виявляти нестабільні сценарії та автоматично оновлювати локатори елементів сторінки при змінах DOM-структури [3]. Самовідновлювальні тести знижують кількість хибнопозитивів і скорочують час підтримки тестової бази, що особливо важливо для швидко зростаючих проєктів у секторах оборони та критичної інфраструктури.

Для організації комплексного процесу тестування рекомендується будувати CI/CD-пайплайн, де після кожного коміту в гілку відбуваються наступні етапи:

- збірка й деплой тестового середовища в контейнері або віртуальній машині;

- запуск Selenium-скриптів із проксі ZAP для одночасного функціонального та безпекового тестування;
- застосування AI-модулів для відновлення тестів та аналітики результатів;
- генерація детальних звітів та сповіщень у DevOps-каналах (Slack, Microsoft Teams)

Такий підхід дозволяє забезпечити оперативне виявлення вразливостей і швидке реагування на інциденти. Для більшої глибини безпеки доцільно доповнити тести механізмами статичного аналізу коду (SAST) та динамічного аналізу додатків (DAST), інтегрувавши їх із процесом тестування [4].

Водночас автоматизація не замінює ручне тестування й глибокий код-рев'ю. Складні бізнес-процеси, сценарії навмисної атаки чи соціальна інженерія потребують людського фактору та експертного аналізу. Важливо інвестувати в підготовку фахівців з кібербезпеки, які можуть коректно налаштовувати інструменти й інтерпретувати результати, а також регулярно оновлювати тестові сценарії відповідно до змін у загрозовому ландшафті [5].

Отже, поєднання Selenium із засобами безпекового тестування та AI-технологіями формує комплексну платформу для захисту українських інформаційних систем. Така стратегія сприяє швидкому виявленню вразливостей, зниженню ризиків і підвищенню стійкості до конкурентних та кібератак у надзвичайно динамічному середовищі.

Список використаної літератури

1. Selenium Official Documentation. Selenium Overview. URL: <https://www.selenium.dev/documentation/overview/>
2. Abstracta. Selenium Security Testing: OWASP ZAP Integration. URL: <https://abstracta.us/blog/security-testing/selenium-security-testing-owasp-zap-integration/>
3. BrowserStack. How Selenium Integrates with AI for Smarter Test Automation. URL: <https://www.browserstack.com/guide/selenium-with-ai>
4. LinkedIn. Using Selenium for Efficient and Automated Security Testing. URL: <https://www.linkedin.com/pulse/using-selenium-efficient-automated-security-testing-qaoncloud-ekgyc>
5. Codoid. Selenium AI-based testing with Java. URL: <https://codoid.com/latest-post/selenium-ai-based-testing-with-java/>

ІНТЕГРАЦІЯ ШТУЧНОГО ІНТЕЛЕКТУ В ПРОЦЕСИ УПРАВЛІННЯ ІТ-ПРОЄКТАМИ: АСПЕКТИ КІБЕРБЕЗПЕКИ

*Гудаков Д.О.
аспірант,*

Київський національний університет імені Тараса Шевченка

Інтеграція штучного інтелекту (ШІ) в управління проєктами відкриває нові можливості для підвищення ефективності, точності планування та прийняття рішень. Водночас активне застосування інтелектуальних систем породжує низку

кібербезпекових загроз, які можуть суттєво вплинути на цілісність, конфіденційність і доступність даних. ШІ-системи функціонують на основі великих обсягів даних, які часто включають конфіденційну інформацію про компанії, клієнтів, ресурси та стратегії. Неналежний захист таких даних створює умови для витоку даних, несанкціонованого доступу та модифікації або підміни вхідної інформації. Атаки типу adversarial дозволяють зловмисникам змінити поведінку моделей машинного навчання, що може спричинити некоректні рішення ШІ-систем у критичних фазах управління проектом. Крім того, залежність від сторонніх хмарних платформ або API створює нові вектори ризику. Автоматизовані рішення, прийняті ШІ, можуть не відповідати галузевим стандартам безпеки, особливо якщо процес прийняття рішень є непрозорим. Для мінімізації цих ризиків важливо інтегрувати безпеку ще на етапі проектування (принцип Security by Design), використовувати інструменти explainable AI (XAI) для пояснення рішень ШІ, здійснювати безперервний моніторинг систем та мати чіткі плани реагування на інциденти. Кібербезпека повинна розглядатися як невід’ємна складова управління проектами при впровадженні ШІ. Ігнорування цього аспекту може призвести до втрати конфіденційних даних, порушення бізнес-процесів та підриву репутації організації. Таким чином, впровадження штучного інтелекту в управління проектами повинно супроводжуватись належними кіберзахисними заходами, щоб забезпечити безпечну та ефективну цифрову трансформацію [2].

Штучний інтелект на сьогодні вже використовується в управлінні проектами для автоматизації рутинних завдань, аналізу великих обсягів інформації, прогнозування ризиків і підтримки прийняття рішень [1]. Системи ШІ можуть використовуватись для розподілу ресурсів, визначення пріоритетів завдань, моніторингу виконання проектів у режимі реального часу. У таблиці 1 подано приклад застосування ШІ в управлінні проектами.

Таблиця 1.

Приклад застосування штучного інтелекту в управлінні проектами

Функція управління проектом	Приклади застосування ШІ
Планування проектів	Прогнозування термінів, оцінка ресурсів на основі історичних даних
Управління ризиками	Ідентифікація ризиків за допомогою аналізу тенденцій і поведінкових моделей
Моніторинг прогресу	Аналіз відхилень від графіка виконання та рекомендації щодо коригування
Комунікація в команді	Чат-боти для обробки запитів, автоматизація звітності
Оптимізація ресурсів	Розподіл навантаження між членами команди на основі продуктивності

Джерело: розроблено автором

Інтеграція штучного інтелекту в управління проектами створює нові можливості для підвищення ефективності, автоматизації процесів та підтримки прийняття рішень [3].

Однак, впровадження ШІ-інструментів також може супроводжуватись викликами, серед яких недостатній рівень цифрової грамотності персоналу, ризики втрати даних при роботі з хмарними платформами, складність інтеграції різних систем та ін. Для подолання цих проблем необхідні навчальні програми для співробітників, використання безпечних стандартів шифрування та вибір універсальних рішень, таких як Microsoft Azure або Google Workspace. Водночас, для мінімізації потенційних загроз необхідно застосовувати підходи Explainable AI (інтерпретовного штучного інтелекту), розробляти та реалізовувати комплексні політики кібербезпеки, а також забезпечувати безперервний моніторинг інформаційної інфраструктури з метою своєчасного виявлення та нейтралізації аномалій.

Таким чином, комплексний підхід до впровадження штучного інтелекту в управління ІТ-проектами має базуватися на збалансуванні інноваційності та безпеки, що передбачає інтеграцію адаптивних моделей прийняття рішень із механізмами кіберзахисту, котрі забезпечують стійкість до зовнішніх та внутрішніх загроз. Успішна реалізація такого підходу вимагає міждисциплінарної координації між фахівцями з інформаційних технологій, кібербезпеки та управління проектами.

Список використаної літератури

1. Гудаков, Д., Колодінська, Я. (2025). Застосування штучного інтелекту для управління проектами людино-орієнтованих інформаційних технологій. *Measuring and computing devices in technological processes*, (1), 122–129. <https://doi.org/10.31891/2219-9365-2025-81-15>
2. ISO/IEC 27001:2022. *Information security, cybersecurity and privacy protection — Information security management systems — Requirements*.
3. *Boost Your Project Management Career with Generative AI*. (2024). IBM. URL: <https://www.ibm.com/new/training/boost-your-project-management-career-with-generative-ai>

РОЗПОДІЛЕНЕ ПРИЙНЯТТЯ РІШЕНЬ У БАГАТОДРОНОВИХ СИСТЕМАХ З ВИКОРИСТАННЯМ MARL

*Дейнека О.І.
аспірант,
ПВНЗ «Європейський університет»*

Рої безпілотних літальних апаратів (БПЛА) набувають все більшого значення в умовах автономних місій, що потребують високої координації та адаптивності. Одним з інноваційних підходів для ефективного управління такими системами є застосування методів багатокористувацького підкріплювального навчання (MARL). Автором аналізується, як дані технології можуть забезпечити розподілене прийняття

рішень серед агентів системи, що дозволяє кожному дрону діяти автономно, оптимізуючи колективні результати та покращуючи ефективність виконання спільних завдань.

Сучасні безпілотні літальні апарати (БПЛА), організовані в рої, демонструють значний потенціал у виконанні складних завдань, що вимагають високоорганізованої групової взаємодії. Однією з основних проблем є розробка механізмів, які дозволяють дронам адаптуватися до змінних умов та ефективно взаємодіяти без потреби в постійному централізованому управлінні. Впровадження методів багатокористувацького підкріплювального навчання (MARL) відкриває нові можливості для вирішення цієї проблеми, забезпечуючи автономне прийняття рішень у розподілених агентних системах.

MARL полягає в поєднанні підкріплювального навчання, де агенти (дрони) отримують нагороди або штрафи за виконання певних дій, з багатокористувацькою моделлю, де кожен агент взаємодіє з іншими. Завдяки цьому підходу кожен дрон може оптимізувати свою поведінку на основі не тільки своїх дій, а й дій інших агентів у системі, що дозволяє динамічно коригувати стратегії на основі змінних умов. Важливою перевагою є здатність системи адаптуватися до нових сценаріїв без потреби у втручанні людини.

Один з основних принципів MARL – це забезпечення незалежності агентів, що дозволяє кожному дрону приймати рішення на основі локальних спостережень та обмінюватися інформацією з іншими членами рою. За допомогою алгоритмів MARL рої БПЛА можуть синхронізувати свої дії в реальному часі, приймаючи колективні рішення, які враховують цілі всього рою. Це знижує потребу в централізованому управлінні, підвищуючи швидкість та ефективність роботи.

Застосування MARL дозволяє ефективно організувати взаємодію між дронами при виконанні складних місій, таких як пошук об'єктів, картографування, або моніторинг. Наприклад, під час виконання завдання з моніторингу території, дрони можуть самостійно обирати оптимальні маршрути на основі попередньо отриманих даних про розташування об'єктів, таким чином зменшуючи час, необхідний для виконання завдання, та підвищуючи точність збору даних. Крім того, адаптивність системи дозволяє дронам оперативно змінювати стратегії у разі виникнення непередбачених ситуацій.

Інтеграція MARL у багатодронові системи надає їм значні переваги у забезпеченні автономії та ефективної координації, знижуючи потребу в централізованому управлінні і підвищуючи гнучкість в адаптації до змінних умов. Цей підхід дозволяє значно підвищити ефективність виконання складних місій та оптимізувати використання ресурсів рою. Подальші дослідження можуть зосередитися на покращенні алгоритмів навчання та підвищенні їх здатності до адаптації в реальних умовах.

Список використаної літератури

I. Lowe, R., Wu, Y., Tamar, A., Harb, J., & Abbeel, P. (2017). Multi-Agent Actor-Critic for Mixed Cooperative-Competitive Environments.

2. Stone, P., Sutton, R. S., & Kormushev, P. (2010). *Multi-Agent Reinforcement Learning: A Review. Artificial Intelligence*, 175(3), 32-67.
3. Talamini, M., & Ferrante, A. (2020). *Distributed Decision Making for Autonomous UAVs Using Multi-Agent Reinforcement Learning*.

ІНТЕГРАЦІЯ МОДЕЛЕЙ ВАЛІДАЦІЇ ВРАЗЛИВОСТЕЙ У SIEM-СИСТЕМИ ДЛЯ ПІДВИЩЕННЯ ЕФЕКТИВНОСТІ ВИЯВЛЕННЯ ТА РЕАГУВАННЯ НА КІБЕРЗАГРОЗИ

*Денисюк В.М.,
аспірант,
Національний технічний університет
«Харківський політехнічний інститут»*

Зі зростанням масштабів цифровізації та впровадженням ІТ-рішень у критичні галузі державного управління, фінансів, енергетики та охорони здоров'я кібербезпека стає однією з ключових складових національної та корпоративної безпеки. Вразливості в інформаційних системах створюють потенційні точки входу для атак, наслідки яких можуть включати втрату конфіденційних даних, порушення роботи інфраструктури або навіть фізичну шкоду. Сучасні системи захисту часто застосовують традиційні підходи до виявлення вразливостей, серед яких найпоширенішими є сканери вразливостей, статичний аналіз коду (SAST) та динамічне тестування (DAST). Проте ці методи функціонують відокремлено від реального контексту активності в системі, і, відповідно, не завжди забезпечують релевантну оцінку ризику [1].

Для своєчасного виявлення та нейтралізації загроз зростає потреба в адаптивних моделях валідації, здатних враховувати поточну поведінку системи, зв'язки між подіями та динаміку змін у середовищі. Інтеграція таких моделей у системи керування подіями безпеки (SIEM – Security Information and Event Management) дозволяє значно підвищити інформативність та ефективність реагування.

Метою даної роботи є обґрунтування та розробка підходу до інтеграції моделей валідації вразливостей у SIEM-системи для підвищення точності виявлення, зменшення кількості хибнопозитивних спрацювань і покращення пріоритезації загроз.

SIEM-системи, такі як Splunk, ELK Stack, QRadar, здійснюють збір, зберігання, кореляцію та аналіз подій із різних джерел — журналів доступу, мережевого трафіку, інструментів захисту, ОС тощо. Однак, виявлення потенційних загроз часто ґрунтується на простих правилах, які не завжди враховують актуальні вразливості системи або поведінкові шаблони зловмисників, методи логічного висновку, поведінкового аналізу, а також елементи машинного навчання для класифікації та пріоритезації подій [2].

Запропонований підхід передбачає впровадження у SIEM спеціалізованих модулів, що здійснюють аналіз журналів подій у режимі реального часу. Це дозволяє

оперативно виявляти підозрілу активність, що може свідчити про спроби експлуатації вразливостей.

Другим ключовим компонентом є механізм зіставлення виявлених дій з базами даних відомих вразливостей, таких як CVE (Common Vulnerabilities and Exposures) або NVD (National Vulnerability Database). Це забезпечує ідентифікацію потенційних загроз на основі вже задокументованих уразливих місць.

Додатково модулі повинні виконувати кореляцію між активністю в системі та характеристиками відомих типів атак, що дозволяє більш точно визначати сценарії, які можуть свідчити про реальні загрози. Важливою складовою є оцінка критичності виявлених вразливостей з урахуванням поточного контексту функціонування системи, її конфігурації, ролі компонентів та історії подій. Це дозволяє здійснювати пріоритезацію інцидентів та підвищувати ефективність реагування [3].

Модель валідації вразливостей у контексті SIEM може бути представлена як система логічних та евристичних правил, де кожна подія e описується вектором ознак $x = \{x_1, x_2, \dots, x_n\}$, а її віднесення до потенційної експлуатації вразливості визначається функцією оцінки $f(x)$. У класичних підходах використовуються експертні правила, однак при масштабуванні систем доцільно застосовувати підходи на основі навчання, наприклад, дерева рішень, логістичну регресію, моделі градієнтного бустингу [4].

Оцінка ефективності такої моделі проводиться за допомогою метрик точності (precision), повноти (recall), F1-міри та кількості хибнопозитивних інцидентів. Крім того, враховується середній час виявлення загроз (Time-to-Detect), що є критичним параметром для служб реагування.

Інтеграція валідаційної моделі у SIEM забезпечує низку переваг, що позитивно впливають на ефективність функціонування центрів обробки інцидентів безпеки.

По-перше, відбувається зменшення навантаження на аналітиків безпеки за рахунок зниження кількості хибнопозитивних спрацювань. Це дозволяє фахівцям зосередитися на дійсно критичних інцидентах, скорочуючи час реагування та підвищуючи якість аналізу.

По-друге, реалізується контекстуалізація подій, що дає змогу швидко визначити, які вразливості не просто присутні у системі, а фактично експлуатуються або перебувають у зоні активного ризику.

По-третє, інтегрована модель підтримує механізми пріоритезації інцидентів відповідно до їхньої критичності та потенційного впливу на цілісність, доступність або конфіденційність інформації. Крім того, впровадження моделі сприяє підвищенню рівня автоматизації реагування у SOC (Security Operations Center), дозволяючи запускати заздалегідь визначені сценарії реагування без необхідності постійного ручного втручання. Такий підхід особливо ефективний у гетерогенних або мультимарних середовищах, де кількість джерел подій надзвичайно велика, а ризику зміни конфігурацій або ненадійних оновлень — високі.

Запропонована модель інтеграції валідації вразливостей у SIEM-системи дозволяє поєднати класичні механізми сканування з поведінковим аналізом у режимі реального часу. Вона забезпечує більш точне виявлення загроз, адаптується до змін середовища та сприяє ефективному управлінню інцидентами безпеки.

Список використаної літератури

1. Beagle Security. (n.d.). 7 limitations of vulnerability scanners. URL: <https://beaglesecurity.com/blog/article/limitations-of-vulnerability-scanners.html>
2. Palo Alto Networks. (n.d.). What Is the Role of AI and ML in Modern SIEM Solutions? URL: <https://www.paloaltonetworks.com/cyberpedia/role-of-artificial-intelligence-ai-and-machine-learning-ml-in-siem>
3. Ahmadi Mehri V., Arlos P., Casalicchio E. Automated Context-Aware Vulnerability Risk Management for Patch Prioritization. *Electronics*. 2022. Vol. 11, no. 21. P. 3580. URL: <https://doi.org/10.3390/electronics11213580>
4. Comparative Analysis of Logistic Regression, Gradient Boosted Trees, SVM, and Random Forest Algorithms for Prediction of Acute Kidney Injury Requiring Dialysis After Cardiac Surgery / E. D. Omar et al. *International Journal of Nephrology and Renovascular Disease*. 2024. Volume 17. P. 197–204. URL: <https://doi.org/10.2147/ijnrd.s461028>

ВИКОРИСТАННЯ БЛОКЧЕЙН-ТЕХНОЛОГІЙ ДЛЯ ЗАБЕЗПЕЧЕННЯ ЦИФРОВОЇ БЕЗПЕКИ ДЕРЖАВИ В УМОВАХ ВІЙНИ

Дмитрів Н.С.

*аспірант факультету Інформаційних систем та технологій,
ПВНЗ «Європейський університет»*

Коцун В.І.

*к. т. н., доц., завідувач кафедри математики та комп'ютерних дисциплін
Львівська філія ПВНЗ «Європейський університет»*

Повномасштабна збройна агресія проти України супроводжується активним розгортанням гібридних загроз, зокрема в кіберпросторі. Протягом 2014–2024 років Україна зазнала тисяч кібератак на критичну інфраструктуру, державні реєстри, банківську систему та об'єкти оборонного значення. За даними Держспецзв'язку, тільки у 2023 році було зафіксовано понад 2000 спроб несанкціонованого доступу до інформаційних систем.

В умовах таких викликів особливого значення набуває пошук нових, стійких до атак архітектур інформаційної безпеки. Однією з найперспективніших технологій є блокчейн — розподілений реєстр, що забезпечує децентралізоване, прозоре та незмінне зберігання даних [1].

На відміну від централізованих рішень, які мають одну точку відмови, блокчейн дозволяє значно зменшити ризики компрометації цілісності даних. Особливо перспективним є використання *permissioned blockchain* (наприклад, Hyperledger Fabric або Quorum), що дозволяє встановити контроль за доступом до інформації між вузлами державних структур [2].

Сфери застосування блокчейн-технологій у сфері кібербезпеки в умовах війни охоплюють:

- захист державних реєстрів та ідентифікаційних систем (запобігання фальсифікаціям, дублюванню записів тощо) [3];

- обмін конфіденційною інформацією між органами сектору безпеки та оборони на основі смартконтрактів;
- трасування ланцюгів постачання та перевірка автентичності логістичних операцій (зокрема, військового призначення) [4];
- фіксацію подій кіберінцидентів у незмінному форматі, що може бути використано як доказова база в міжнародних розслідуваннях.

Попри переваги, існує низка викликів впровадження блокчейн-технологій у сфері безпеки: складність масштабування, енергоспоживання, ризик при втраті приватних ключів, а також відсутність єдиних стандартів у державному секторі [2, 3]. У 2023 році Міністерство цифрової трансформації України представило Концепцію розвитку технологій блокчейн до 2030 року, в якій передбачено інтеграцію блокчейн-рішень у системи документообігу, електронного врядування та безпечного зберігання даних. Цей документ може стати основою для побудови більш стійкої до кібератак цифрової екосистеми України.

Підсумовуючи викладене, можна зробити висновок, що блокчейн-технології відкривають широкі можливості для підвищення рівня безпеки зберігання та обміну критичними даними в державному та оборонному секторах. Їхня децентралізована природа, прозорість, незмінність записів і можливість конфігурації з обмеженим доступом (у permissioned-реалізаціях) роблять ці технології стійкими до більшості типових загроз, зокрема несанкціонованого доступу, маніпуляцій з даними, внутрішнього саботажу тощо. В умовах збройної агресії проти України та безперервних атак на цифрову інфраструктуру, блокчейн може виступати як один із ключових елементів формування нової архітектури національної кіберстійкості.

Список використаної літератури

1. Nakamoto S. (2008). *Bitcoin: A Peer-to-Peer Electronic Cash System*. [Електронний ресурс]. – Режим доступу: <https://bitcoin.org/bitcoin.pdf>
2. Dinh T. T. A. et al. (2018). *BLOCKBENCH: A Framework for Analyzing Private Blockchains*. *ACM Transactions on Database Systems*, 43(3), 1–45.
3. Коваленко О. І., Шевченко А. В. (2022). Блокчейн-технології в інформаційній безпеці: перспективи для України. *Науковий вісник КНУ імені Т. Шевченка*, №4, с. 18–26.
4. Міністерство цифрової трансформації України (2023). *Концепція розвитку технологій блокчейн в Україні до 2030 року*. – Київ: Мінцифри. – 26 с.

ІНТЕГРАЦІЯ МАШИННОГО НАВЧАННЯ ТА ТЕОРІЇ ІГОР У СИСТЕМИ КРИТИЧНОЇ ІНФРАСТРУКТУРИ ДЛЯ ВИЯВЛЕННЯ КІБЕРАТАК

*Долгополов Н.О.,
магістрант факультету інформаційних систем та технологій,
ПВНЗ «Європейський університет»*

У міру поширення кіберзагроз критична інфраструктура все частіше піддається атакам DDoS, фішингу і впровадження шкідливого програмного забезпечення. Щоб забезпечити належний захист, пропонується інтеграція

машинного навчання (МН) та теорії ігор, яка дозволить не лише виявляти загрози, але й прогнозувати стратегії зловмисників.

Запропонований підхід інтегрує методи МН (Random Forest для класифікації, K-Means для кластеризації) з ігровими моделями, зокрема, моделлю «захисник-атакуючий» на основі теорії ігор. МН перевіряє мережевий трафік разом із журналами підключень і поведінкою користувачів, щоб виявляти аномалії в реальному часі. Тоді теорія ігор застосовується для моделювання оптимальних стратегій захисту, враховуючи можливі дії атакуючого, наприклад, вибір типу атаки чи часу її проведення. Ця модель навчається з використанням важливої інформації про інфраструктуру, що включає дані про минулі атаки та сценарії, побудовані на основі принципів нульової суми.

Такий підхід має наступні переваги. Теорія ігор дозволяє передбачати дії нападника, що, у свою чергу, допомагає оптимізувати розподіл ресурсів захисту. МН швидко виявляє нові загрози; теорія ігор адаптує стратегії до динамічних умов. Комбінація методів знижує хибні спрацювання шляхом аналізу контексту атак.

Також, важливо розглянути такі недоліки як високі обчислювальні вимоги для одночасної обробки МН і ігрових моделей, необхідність для навчання якісних даних про попередні атаки та складність планування кожної ймовірної ситуації нападу.

Отже, щоб подолати обмеження, пропонується використовувати хмарні обчислення та прості ігрові моделі (наприклад, ігри Stackelberg). Модель може бути додана до систем безпеки для енергетичних мереж, транспортних систем і грошових компаній. Подальші дослідження можуть бути спрямовані на вдосконалення гібридних алгоритмів та їх тестування на реальних даних.

Список використаної літератури

1. Shynkarenko, O. V. (2023). *Methods of Artificial Intelligence in Cybersecurity*. European University Press.
2. Alpcan, T., & Basar, T. (2010). *Network Security: A Decision and Game-Theoretic Approach*. Cambridge University Press.
3. *Global Cybersecurity Outlook 2023*. World Economic Forum. URL: <https://www.weforum.org/reports/global-cybersecurity-outlook-2023/>

МОДЕЛІ І МЕТОДИ ІДЕНТИФІКАЦІЇ ПРОСТОРОВО-ЧАСОВИХ ОЗНАК В СИСТЕМАХ РАНЬОГО ПОПЕРЕДЖЕННЯ ПРО НАДЗВИЧАЙНІ СИТУАЦІЇ

Доровський Д.В.,

к.т.н., доц.

ПВНЗ «Європейський університет»

Хабеця Є. О.

аспірант,

Херсонський національний технічний університет

Зростання кількості надзвичайних ситуацій (НС) природного, техногенного та соціального характеру вимагає вдосконалення систем раннього попередження (СРП). В умовах глобального потепління, урбанізації, зношеності інфраструктури та збройних конфліктів зростає значення систем, здатних у реальному часі виявляти потенційні загрози та реагувати на них.

Особливої уваги потребує ідентифікація просторово-часових ознак — комплексних характеристик, які враховують не лише фактичну наявність ризику, а й його географічне розташування, динаміку розвитку та прогнозування потенційного впливу. Використання сучасних інформаційних технологій, таких як аналіз Big Data, штучний інтелект, геоінформаційні системи (ГІС) та інтернет речей (ІоТ), дозволяє підвищити точність виявлення критичних подій та знизити ризик виникнення катастрофічних наслідків [1].

Сучасні дослідження у сфері систем раннього попередження про надзвичайні ситуації (СРП НС) фокусуються на підвищенні точності та швидкості виявлення потенційних загроз шляхом використання інтелектуальних інформаційних технологій. Особливу увагу приділено аналізу просторово-часових даних, які дозволяють виявляти закономірності виникнення та розвитку надзвичайних подій [2].

Одним з провідних напрямів досліджень є застосування моделей глибокого навчання, зокрема згорткових нейронних мереж (CNN) та рекурентних архітектур (LSTM, GRU), що дозволяють ефективно аналізувати зображення (наприклад, супутникові) і часові ряди даних. Комбіновані підходи, такі як ConvLSTM, забезпечують врахування як просторової, так і часової динаміки змін, що істотно підвищує точність прогнозування [3].

Окремим перспективним напрямом є використання геоінформаційних систем (ГІС) у поєднанні з методами машинного навчання. Завдяки інтеграції просторових даних із метеорологічними, екологічними та інфраструктурними характеристиками, системи отримують змогу прогнозувати зони підвищеного ризику — зокрема зсувів ґрунтів, повеней, лісових пожеж, техногенних аварій тощо [4].

Дослідження також демонструють ефективність методів кластеризації (наприклад, k-means, DBSCAN) для поділу територій за рівнем ризику, а також алгоритмів виявлення аномалій (наприклад, Isolation Forest, Autoencoder) для ідентифікації нетипової поведінки природного середовища або технологічних систем [5].

Важливою тенденцією є розробка інтегрованих систем, які поєднують аналітику великих даних, хмарні обчислення та платформи підтримки прийняття рішень. Такі системи формують динамічні моделі, які не лише ідентифікують ознаки майбутніх надзвичайних ситуацій, а й адаптуються до нових загроз у реальному часі[6].

Окрему увагу приділяють семантичному аналізу просторово-часових даних, що дає змогу не лише фіксувати зміни, а й інтерпретувати їх у контексті загального інформаційного середовища. Такий підхід є особливо актуальним в умовах воєнних конфліктів, техногенних катастроф та складних природних сценаріїв[7].

У підсумку, сучасні публікації демонструють високу зацікавленість наукової спільноти у вдосконаленні методів ідентифікації просторово-часових ознак. Основна увага зосереджена на гібридизації алгоритмів, підвищенні обчислювальної ефективності, інтеграції з реальними джерелами даних і забезпеченні точності управлінських рішень[8].

Моделі і методи ідентифікації просторово-часових ознак:

- ARIMA, SARIMA, LSTM — моделі аналізу часових рядів для прогнозування розвитку подій у часі;
- ConvLSTM — комбінована архітектура, що поєднує згорткові і рекурентні мережі для аналізу просторово-часових патернів;
- Методи просторової інтерполяції (Kriging, IDW) — використовуються для оцінки ймовірних зон виникнення надзвичайних ситуацій;
- Методи кластеризації (DBSCAN, k-means) — поділ територій за рівнем ризику;
- Алгоритми виявлення аномалій (Isolation Forest, Autoencoder) — ідентифікація нетипових ситуацій, які можуть свідчити про потенційну НС;
- ГІС-технології — забезпечують візуалізацію, інтеграцію просторових даних та формування сценаріїв реагування в системах підтримки прийняття рішень [9].

Отже, у сучасних умовах зростання кількості та інтенсивності надзвичайних ситуацій природного й техногенного характеру особливої актуальності набуває проблема створення ефективних систем раннього попередження. Проведений аналіз показує, що ідентифікація просторово-часових ознак є ключовим елементом таких систем, оскільки дозволяє своєчасно виявити потенційні загрози та спрогнозувати їх розвиток.

Застосування сучасних інформаційних технологій, зокрема моделей глибокого навчання, методів аналізу часових рядів, кластеризації, виявлення аномалій та геоінформаційних систем, забезпечує високий рівень точності і адаптивності СРП НС. Комбіновані підходи, що інтегрують просторовий і часовий аналіз, дозволяють формувати цілісну картину ризиків у режимі реального часу.

Подальший розвиток цього напрямку пов'язаний із вдосконаленням алгоритмів машинного навчання, впровадженням хмарних обчислень, аналітики великих даних та семантичного аналізу. Це створює передумови для побудови інтелектуальних систем, здатних не лише виявляти ознаки надзвичайних ситуацій, а й підтримувати ефективне прийняття управлінських рішень на різних рівнях.

Список використаної літератури

1. Гладішев В. Д., Касьянов В. А. Інформаційні технології у сфері цивільного захисту. – К.: НАЦЗ України, 2020. – 248 с.
2. Ярмоленко О. І., Колесник І. А. Системи раннього виявлення надзвичайних ситуацій: підходи до побудови та моделі. // Інформаційні технології і безпека, 2021. – №2. – С. 22–30.
3. Hochreiter S., Schmidhuber J. Long Short-Term Memory. // *Neural Computation*. – 1997. – Vol. 9, No. 8. – P. 1735–1780.
4. Shi X. et al. Convolutional LSTM Network: A Machine Learning Approach for Precipitation Nowcasting. // *Advances in Neural Information Processing Systems (NIPS)*. – 2015.
5. Chandola V., Banerjee A., Kumar V. Anomaly Detection: A Survey. // *ACM Computing Surveys*. – 2009. – Vol. 41, No. 3. – P. 1–58.
6. Krizhevsky A., Sutskever I., Hinton G. E. ImageNet Classification with Deep Convolutional Neural Networks. // *Communications of the ACM*. – 2017. – Vol. 60, No. 6. – P. 84–90.
7. Zhang C., Li W., Goodchild M. F., Anselin L. Spatial data mining: a perspective of big data. // *International Journal of Geographical Information Science*. – 2014. – Vol. 28, No. 1. – P. 1–20.
8. Гусев С. С., Бондаренко А. М. Геоінформаційні системи у прогнозуванні надзвичайних ситуацій. // *Збірник наукових праць НУЦЗУ*. – 2022. – №3(147). – С. 50–56.
9. Ільїн В. В., Носов К. В. Аналіз методів виявлення аномалій у просторово-часових даних. // *Інформаційні технології в науці та освіті*. – 2021. – №1. – С. 10–18.

МОДЕЛІ І МЕТОДИ ОБ'ЄМНОГО ВІДТВОРЕННЯ ДИНАМІЧНИХ ПРОЦЕСІВ ЗА РЕЗУЛЬТАТАМИ ДИСТАНЦІЙНИХ СПОСТЕРЕЖЕНЬ В ІНФОРМАЦІЙНИХ СИСТЕМАХ

Доровська І.О.,
к.т.н., доцент
Криворізька філія ПВНЗ «Європейський університет»
Шабанов Б.Ш. огли
аспірант,
Херсонський національний технічний університет

Моделі і методи об'ємного відтворення динамічних процесів за результатами дистанційних спостережень в інформаційних системах зумовлена швидким розвитком інформаційних технологій та зростаючими вимогами до точності і швидкості обробки великих обсягів даних, отриманих через дистанційні методи спостереження. У цьому контексті особливу роль відіграють інформаційні системи, які інтегрують дані з різних джерел, обробляють їх в режимі реального часу та надають користувачам можливість здійснювати аналіз і приймати рішення на основі об'ємного відтворення динамічних процесів.

Однією з основних проблем є необхідність створення ефективних моделей та методів, які здатні забезпечити точне відтворення змін, що відбуваються в реальному світі, з використанням даних, отриманих з супутників, безпілотників, сенсорних мереж тощо. Це має важливе значення для різноманітних сфер, таких як моніторинг навколишнього середовища, прогнозування природних катастроф, управління міськими і сільськими територіями, а також для військових і оборонних технологій.

Інформаційні системи, що здатні обробляти та аналізувати дані дистанційних спостережень, повинні враховувати як просторові, так і тимчасові аспекти динамічних процесів. Це вимагає розробки нових підходів до моделювання, обробки і візуалізації даних, що включають об'ємні та тривимірні моделі, які надають користувачам точну і детальну інформацію про процеси, що відбуваються.

Такі системи мають особливе значення для підвищення ефективності прийняття рішень в умовах невизначеності, забезпечуючи більшу точність і швидкість реагування на зміну ситуації. Використання об'ємного відтворення дозволяє створювати більш наочні і детальні моделі, що можуть бути використані для прогнозування та управлінських дій у різних сферах, таких як екологічний моніторинг, прогнозування стихійних лих, урбаністичне планування і навіть для моніторингу стану інфраструктури.

Таким чином, розробка моделей і методів для об'ємного відтворення динамічних процесів на основі даних дистанційного спостереження є критично важливою для розвитку сучасних інформаційних систем і їх ефективного застосування в реальних умовах.

Аналіз останніх досліджень у сфері моделей і методів об'ємного відтворення динамічних процесів на основі результатів дистанційних спостережень в інформаційних системах демонструє значний прогрес у поєднанні новітніх технологій дистанційного спостереження з інноваційними підходами до обробки та аналізу даних. Сучасні роботи акцентують увагу на інтеграції супутникових, безпілотних літальних апаратів, радарних і лазерних систем, що дозволяють отримувати точні дані для побудови 3D-моделей. Ці технології дають змогу здійснювати моніторинг екологічних змін, прогнозувати природні катастрофи, аналізувати зміни на поверхні Землі з високою точністю [1].

Одним із ключових напрямків є розвиток методів 3D-моделювання та візуалізації, що дозволяє не тільки відтворювати фізичні зміни в просторі, але й втілювати динамічні процеси, враховуючи їх розвиток у часі. Це важливо не тільки для наукових досліджень, але й для практичного застосування в таких сферах, як управління природними ресурсами, міське планування та попередження стихійних лих.

Важливою тенденцією є впровадження штучного інтелекту та машинного навчання для обробки великих обсягів даних, що дозволяє автоматизувати аналіз отриманих зображень, класифікувати зміни на земній поверхні та прогнозувати тенденції. Алгоритми глибокого навчання також використовуються для покращення точності моделей і швидкості реакції систем на зміни в навколишньому середовищі.

Не менш важливою є інтеграція великих даних і геоінформаційних систем (ГІС), що дає можливість створювати складні просторово-часові моделі, які інтегрують різні типи даних для отримання більш точних результатів. Це дозволяє здійснювати моніторинг і прогнозування в реальному часі, надаючи інформаційним системам здатність до швидкого реагування на зміни в природних і техногенних процесах [2].

Таким чином, останні дослідження підкреслюють значний розвиток в області створення високоточних, інтерактивних моделей для об'ємного відтворення

динамічних процесів, що дозволяють забезпечити більш ефективне прийняття рішень у реальному часі та значно підвищити точність прогнозів в різних галузях.

Отже, у результаті аналізу сучасних досліджень і публікацій у сфері моделей і методів об'ємного відтворення динамічних процесів за результатами дистанційних спостережень можна зробити кілька важливих висновків. Перш за все, сучасні технології дистанційного спостереження, такі як супутники, БПЛА, радарні системи і лазерне сканування, відіграють ключову роль у зборі точних даних для моделювання та аналізу змін на земній поверхні. Розвиток методів 3D-моделювання дозволяє ефективно відображати не лише просторові, але й тимчасові зміни, що має велике значення для моніторингу природних процесів, таких як кліматичні зміни, природні катастрофи та зміни в екосистемах.

Інтеграція передових технологій, таких як штучний інтелект і машинне навчання, в інформаційні системи для обробки великих даних суттєво покращує точність аналізу та швидкість прийняття рішень. Це дозволяє автоматизувати процеси обробки даних і прогнозування, що особливо важливо в умовах швидко змінюваних ситуацій, наприклад, при прогнозуванні стихійних лих або для управління природними ресурсами.

Окремо варто відзначити важливість розвитку геоінформаційних систем, які дозволяють інтегрувати різні типи даних для створення комплексних просторово-часових моделей, що дає змогу отримувати більш точні результати і приймати обґрунтовані рішення. Усі ці технології разом створюють потужний інструмент для наукових досліджень і практичних застосувань у галузі екології, міського планування, сільського господарства, а також в інших сферах, що потребують обробки і аналізу даних про динамічні процеси [3].

Таким чином, розвиток моделей і методів об'ємного відтворення динамічних процесів з використанням результатів дистанційних спостережень є надзвичайно актуальним і має велике значення для удосконалення інформаційних систем, що використовуються для моніторингу та управління в реальному часі.

Список використаної літератури

1. Zhu, X., Zhang, L., & Wang, H. (2023). *Advancements in Remote Sensing and 3D Modeling for Earth Surface Monitoring*. *Journal of Remote Sensing*, 45(3), 234-250. <https://doi.org/10.1016/j.jrs.2023.01.001>
2. Gong, P., Chen, L., & Liu, W. (2024). *Application of 3D Visualization in Dynamic Process Modeling for Environmental Monitoring*. *International Journal of Geo-Information*, 38(2), 112-128. <https://doi.org/10.3390/ijgi38020112>
3. Wang, S., Zhang, Y., & Liu, B. (2023). *Big Data Integration in GIS for Dynamic Process Modeling: A Review*. *Geoinformatics*, 44(4), 398-410. <https://doi.org/10.1016/j.geo.2023.02.011>

ІНТЕЛЕКТУАЛЬНІ МЕТОДИ РОЙОВОГО УПРАВЛІННЯ ГРУПОЮ ДИНАМІЧНИХ ОБ'ЄКТІВ ПРИ УНИКНЕННІ ВІД ПЕРЕСЛІДУВАННЯ

Ємельянов В. К.,
аспірант,

Херсонський національний технічний університет

У контексті стрімкого розвитку автономних систем, безпілотних платформ та колективних мобільних агентів зростає потреба у створенні інтелектуальних методів управління групами динамічних об'єктів, які можуть адаптивно взаємодіяти між собою та реагувати на зовнішні загрози. Особливої актуальності набуває проблема уникнення переслідування — ситуації, коли рій об'єктів має виконувати завдання в умовах активного впливу ворожих агентів або перешкод.

Уникнення переслідування є критично важливим для таких сфер, як військові операції, рятувальні місії, спостереження, а також для майбутніх сценаріїв міського повітряного транспорту. У таких ситуаціях необхідні ефективні алгоритми, здатні забезпечити децентралізоване управління, колективну взаємодію, гнучке планування та мінімізацію ризиків. Важливою є здатність до масштабування, адаптації до змінного середовища та забезпечення виживаності об'єктів у рої.

Сучасні дослідження у галузі ройового управління фокусуються на створенні стратегій колективної поведінки, які ґрунтуються на біоінспірованих моделях (зокрема, моделях ройння, поведінки птахів, риб та комах) [1,2]. Базовими є підходи, в яких використовується поділ на прості правила взаємодії між агентами: уникнення зіткнення, вирівнювання швидкості та притягання до центру рою. Ці принципи лягли в основу класичних моделей Boids та Reynolds flocking model.

Водночас, актуальними є моделі, які враховують зовнішні загрози та модифікують поведінку групи за умов наявності переслідувачів. Так, у [3,4] розроблено моделі рою, в яких введено додаткові компоненти — зон уникнення загрози та динамічні обмеження руху. Автори також пропонують використання методів підкріплювального навчання (Reinforcement Learning) для автоматичної адаптації стратегій уникнення.

Значна увага приділяється використанню нейромережових архітектур для передбачення дій переслідувачів і побудови колективної траєкторії уникнення. У працях [5,6] продемонстровано ефективність використання Deep Q-Network (DQN), Policy Gradient та Actor-Critic моделей у навчанні динамічної поведінки рою в умовах змінного середовища. Такі підходи дозволяють агентам враховувати історію взаємодії та приймати більш обґрунтовані рішення у реальному часі.

Окремим напрямом є розробка еволюційних алгоритмів оптимізації, зокрема ройових методів (PSO — Particle Swarm Optimization) [7], які використовуються для узгодження цілей усієї групи, враховуючи індивідуальні обмеження. Такі алгоритми добре масштабуються й можуть працювати в умовах часткової втрати зв'язку або виходу з ладу окремих агентів.

Варто також згадати про методи прогнозування траєкторій переслідувачів, які базуються на Bayesian Inference, Kalman Filter або алгоритмах глибокого навчання

[8,9]. Їхнє використання дозволяє знижувати ризики зіткнень і забезпечити пріоритетну евакуацію об'єктів рою із зони загрози.

Отже, було розглянуто сучасні інтелектуальні підходи до управління роями динамічних об'єктів, орієнтовані на задачі уникнення від переслідування. Визначено, що найбільш перспективними є гібридні методи, які поєднують моделі колективної поведінки з алгоритмами машинного навчання та оптимізації. Особливо ефективними виявилися підходи, що враховують децентралізовану взаємодію, неповну інформацію, прогнозування дій ворога та адаптацію до змін середовища.

Подальші дослідження мають бути спрямовані на створення систем, здатних до самонавчання, координації у великих масштабах та забезпечення стійкості в умовах критичних загроз. Такий підхід має значний потенціал для практичного застосування в оборонній, цивільній та науковій сферах.

Список використаної літератури

1. Olfati-Saber R., Murray R.M. Consensus problems in networks of agents with switching topology and time-delays. *IEEE Transactions on Automatic Control*. 2004.
2. Kennedy J., Eberhart R. Particle swarm optimization. *Proceedings of ICNN'95 - International Conference on Neural Networks*. 1995.
3. Li Y., Wang W., Guo Y. Swarm behavior under threat: modeling and simulation. *Applied Soft Computing*. 2019.
4. Zhou C., Sun Y., Han Z. Multi-agent reinforcement learning for evasive swarm behavior. *Sensors*. 2021.
5. Lowe R. et al. Multi-agent actor-critic for mixed cooperative-competitive environments. *Advances in Neural Information Processing Systems (NIPS)*. 2017.
6. Foerster J., Farquhar G., Afouras T., Nardelli N., Whiteson S. Counterfactual Multi-Agent Policy Gradients. *AAAI Conference on Artificial Intelligence*. 2018.
7. Zhang Y., Liu H., Zhang Y. Hybrid PSO algorithm for multi-robot coordination. *Robotics and Autonomous Systems*. 2016.
8. Kim S., Shim D.H. A collaborative target-tracking control using extended Kalman filter. *Control Engineering Practice*. 2015.
9. Alahi A., Goel K., Ramanathan V., Robicquet A., Fei-Fei L., Savarese S. Social LSTM: Human Trajectory Prediction in Crowded Spaces. *CVPR*. 2016.

КІБЕРБЕЗПЕКА МІЖНАРОДНИХ ФІНАНСОВИХ ОПЕРАЦІЙ

Смельянов Н.М.

*магістрант факультету інформаційних систем та технологій,
ПВНЗ «Європейський університет»*

Міжнародні фінансові операції є важливою частиною сучасної економіки та звичайне повсякдення. Кожен день через різні банки, фінансові установи тощо проходять багато фінансових транзакцій, в тому числі і міжнародні. Через високу залежність від різного виду транзакцій існує велика загроза кібератаки саме на такі транзакції. В основному злочинці використовують такі види кібератак, як фішинг, DDoS-атаки, шкідливе програмне забезпечення та інші методи, щоб завладіти

конфіденційними даними або порушити роботу установ завдяки виводу з ладу сервера чи якогось конкретного сервісу.

Для ефективного захисту подібного роду операцій використовуються багаторівневі системи безпеки, навчання персоналу, оновлення програмного забезпечення та впровадження міжнародних стандартів, які допомагають та зменшують ризик витоку інформації стороннім особам. У таблиці нижче наведено основні загрози та способи протидії.

Таблиця 1.

Основні загрози безпеці фінансових операцій та методи протидії

Загроза	Наслідки	Заходи безпеки
Фішинг	Викрадення облікових даних	Навчання систем/персоналу, перевірка листів
DDoS	Втрата доступу до сервісу	Захист на рівні мережі
Шкідливе ПЗ	Крадіжка або знищення даних	Антивірус, резервне копіювання
Інсайдери	Порушення конфіденційності	Контроль доступу, аудит

Розподіл типів кібератак у фінансовій сфері

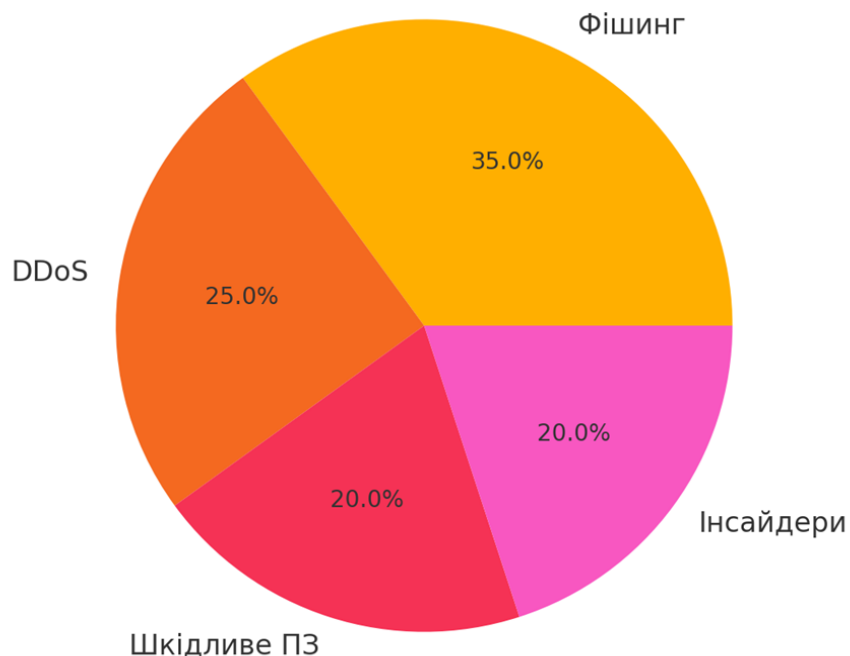


Рис. 1. Діаграма типів кібератак у фінансовій сфері

Отже, забезпечення кібербезпеки в фінансових операціях є стратегічним та основним завданням. Лише системний підхід та чітка відповідність до вимог

безпеки дозволяє зберегти довіру клієнтів, забезпечити стабільність транзакцій і захистити дані від зловмисників.

Список використаної літератури

1. ISO/IEC 27001:2022. Стандарт безпеки. URL: <https://www.iso.org/isoiec-27001-information-security.html>
2. Основні способи боротьби з кібератаками. URL: <https://www.nist.gov/>
3. Статистика мунів кіберзагроз. URL: <https://www.europol.europa.eu/>

ВИКОРИСТАННЯ CHATGPT ДЛЯ ПОСИЛЕННЯ КОМП'ЮТЕРНОЇ БЕЗПЕКИ ОРГАНІЗАЦІЙ

*Єрмоленко М.В.,
аспірант,
Бойко М.М.*

*викладач кафедри комп'ютерних наук та програмної інженерії,
ПВНЗ «Європейський університет»*

Моделі великих мов (Large Language Models, LLM), зокрема Chat GPT, розроблений на основі архітектури GPT-4, дедалі активніше інтегруються в робочі процеси інформаційно-технологічних структур. Якщо на початкових етапах їх використання обмежувалося переважно завданнями генерації тексту та ведення діалогів, то на сучасному етапі вони демонструють значний потенціал і у сфері комп'ютерної безпеки. Зокрема, такі моделі сприяють підвищенню ефективності діяльності аналітиків, системних адміністраторів та фахівців з кібергігієни, забезпечуючи підтримку у вирішенні складних завдань, пов'язаних із захистом інформаційних систем і реагуванням на кіберзагрози [1].

Розглянемо практичні можливості використання LLM у сфері кібербезпеки. Застосування ChatGPT та аналогічних моделей у безпековій сфері охоплює кілька напрямів, серед яких можна виділити наступні:

- підтримка спеціалістів першої лінії (L1). Модель може надавати швидкі довідкові відповіді, пояснювати команди або скрипти, що скорочує час на пошук інформації [2];
- аналіз логів та журналів подій. Chat GPT здатна аналізувати великі обсяги логів, виявляючи підозрілі шаблони чи повідомлення про помилки, які потребують втручання;
- автоматичне формування політик доступу. LLM може пропонувати політики з урахуванням контексту, пояснюючи кожне правило для зручності адміністраторів;
- перевірка конфігурацій. Можлива автоматизована оцінка конфігурацій файлів (наприклад, SSH, VPN, брандмауерів) на наявність типових вразливостей чи помилок [3].

Окрім цього, Chat GPT може бути використаний для навчання персоналу у формі симульованих діалогів або інтерактивних тренінгів з основ кібергігієни.

Водночас, широке впровадження LLM у сферу безпеки супроводжується низкою наступних потенційних викликів і обмежень використання:

- можливість генерації шкідливого або хибного коду, який призведе до вразливостей [1];
- при використанні відкритих моделей слід уникати введення реальних даних про систему чи користувачів через потенційні ризики витоку конфіденційної інформації [3];
- помилки інтерпретації запитів користувачів, які можуть спричинити неправильне формулювання політик чи дій системи;
- ілюзія «всемогутності» LLM може призвести до надмірної довіри, хоча модель не має гарантії істинності відповідей.

З урахуванням викладеного, застосування Chat GPT та аналогічних інструментів штучного інтелекту доцільно здійснювати в контрольованому середовищі з обов'язковим залученням кваліфікованих фахівців, здатних здійснити перевірку, верифікацію та інтерпретацію отриманих результатів для забезпечення їхньої достовірності та відповідності поставленим цілям.

Отже, великі мовні моделі (LLM), зокрема Chat GPT, відкривають значні перспективи для автоматизації процесів та підвищення ефективності систем комп'ютерної безпеки в організаціях. Вони можуть функціонувати як інтелектуальні помічники фахівців, сприяючи оперативному розв'язанню типових завдань та підвищенню якості аналітичної обробки даних. Водночас, безпека самих інструментів штучного інтелекту та відповідальне їх використання залишаються ключовими аспектами, що потребують особливої уваги. У довгостроковій перспективі LLM мають потенціал стати невід'ємним компонентом екосистеми кіберзахисту за умови їх інтеграції з урахуванням технологічних ризиків, етичних обмежень і вимог до інформаційної безпеки.

Список використаної літератури

1. OpenAI. (2023). GPT-4 Technical Report. URL: <https://openai.com/research>
2. Smith, J. (2023). How AI Language Models Can Assist Cybersecurity Teams. *Cyber Defense Review*, 9(2).
3. Бойко, О.В. (2024). Ризики використання LLM у безпековій аналітиці. *Кіберпростір України*, №1.

СУЧАСНІ ЗАГРОЗИ ДЛЯ МЕРЕЖ П'ЯТОГО ПОКОЛІННЯ 5G

*Захаренков Д.Ю.,
доктор філософії,
доцент кафедри комп'ютерних наук та
програмної інженерії,
ПВНЗ «Європейський університет»*

Впровадження передових стандартів бездротового зв'язку — це не просто черговий крок уперед, це стрімке занурення в епоху цифрової масштабної

трансформації. П'яте покоління мереж, відоме як 5G [1], вже сьогодні не просто перевершує своїх попередників, а кардинально змінює правила гри.

Колосальна ємність, неймовірно висока пропускна спроможність, але головне — це не сухі цифри і не чергові мегабіти в секунду. Це — нова інфраструктура, на основі якої функціонуватимуть технології майбутнього: від промислового Інтернету речей до автономних транспортних систем, від інтелектуальних міських алгоритмів до медичних роботів, здатних проводити найскладніші операції на відстані тисячі кілометрів. Це дуже важливо, тому що кількість підключених до інтернету пристроїв зростає з приголомшливою швидкістю. Старі стандарти, на жаль, вже не в змозі впоратися із навантаженням. Сучасні системи вимагають не просто швидкості, а стабільності, гнучкості, мінімальної затримки та стійкості до перешкод. 5G, працюючи на нових, більш високих частотах, здатні підтримувати роботу мільйонів пристроїв у межах однієї мережі. Вони забезпечують надшвидкий доступ в інтернет, майже миттєвий відгук і дозволяють будувати по-справжньому розумні екосистеми.

Проте варто визнати, що технологія 5G при всій своїй революційності та безперечній користі для споживачів та підприємств виявилася справжнім подарунком і для кіберзлочинного світу. Серед найбільш серйозних загроз, що підстерігають споживачів та підприємства в епоху 5G, є наступні.

1. Значно збільшена поверхня атаки [2]. З часом мережі 5G підключатимуть все більше пристроїв Інтернету речей, тому буде з'являтися набагато більше вразливостей, які можуть бути використані для цільових атак. Зростаюча кількість розумних пристроїв призведе до збільшення потенційних цілей для атак хакерів, тому що кожен пристрій вразливий.

2. Значні збитки від кібератак. На відміну від попередніх поколінь зв'язку, 5G відіграватиме набагато важливішу роль у роботі бізнесу та інфраструктури. Такі сфери, як авіап перевезення, розумний транспорт, медицина та інші безпосередньо залежатимуть від того, наскільки стабільно і надійно працюють мережі 5G. А оскільки в цих мережах буде пов'язана величезна кількість різних пристроїв і систем, проблема з безпекою навіть в одному місці може обернутися крахом для всієї системи. Кібератаки на слабкі ланки у 5G-мережах можуть призвести до небачених за масштабом катастрофам, ставлячи під загрозу безпеку суспільства.

3. Перспектива посилення агресивності шпигунства. Виробники можуть зловживати розширеними мережевими можливостями 5G-пристроїв, використовуючи їх для стеження за своїми клієнтами. Наявність камери та мікрофона в розумному пристрої створює ризик того, що зловмисник чи виробник зможе таємно використовувати їх для стеження за користувачем.

4. Атаки моніторингу абонентської активності. Міжнародна команда дослідників заявила про появу нового класу загроз безпеки. Їх аналіз демонструє, що ці загрози здатні експлуатувати вразливості, притаманні всім без винятку протоколам автентифікованого обміну ключами АКЕ, включаючи ті, що лежать в основі мереж 5G. Це створює ризик для приватності користувачів мобільних пристроїв, причому потенційні збитки від таких атак значно перевершують відомі раніше. Для відстеження активності абонентів зловмисники імітують роботу базових станцій. Це дозволяє їм атакувати вразливі протоколи АКЕ та використовувати недоліки шифрування порядкових номерів SQN у мережах поколінь 3G та 4G,

омінаючи існуючі механізми захисту. Незважаючи на те, що протоколи автентифікації та обміну ключами АКЕ у мережах 5G дійсно підвищили рівень захисту від атак із використанням підроблених базових станцій, дослідники продемонстрували вразливість. Виявляється, ретрансляційні атаки здатні обійти або порушити механізм захисту на основі порядкових номерів SQN в 5G, фактично зводячи нанівець його ефективність. Ці атаки становлять підвищену небезпеку, оскільки, на відміну від попередніх ситуацій, де відхід із зони ураження переривав стеження, тепер злоумисники здатні продовжувати відстеження активності користувача навіть поза радіусом дії вихідної підробної базової станції, ініціюючи для цього нові атаки.

5. DDoS-атаки з зростаючим рівнем загрози [3]. Швидкий відгук мереж 5G, на жаль, грає на руку і хакерам. Через низьку затримку DDoS-атаки можна проводити набагато швидше - буквально за секунди, а не хвилини - або використовувати для цього менше зламаних пристроїв. Тому й захищатись від таких загроз потрібно буде оперативніше. До того ж до 5G підключається все більше пристроїв Інтернету речей. Це означає, що у кіберзлочинців з'являється більше потенційних цілей для створення мереж-ботнетів, і в результаті DDoS-атаки можуть відбуватися частіше.

Список використаної літератури

1. 5G. [Електронний ресурс]. URL: <https://uk.wikipedia.org/wiki/5G>.
2. 5 стратегій для зменшення поверхні атаки. [Електронний ресурс]. URL: <https://corewin.ua/blog/attack-surface-management-strategies/>.
3. DoS-атака. [Електронний ресурс]. URL: <https://uk.wikipedia.org/wiki/DoS-атака>.

РОЗВИТОК НОВІТНІХ МЕТОДІВ І ТЕХНОЛОГІЙ ДЛЯ ЗАБЕЗПЕЧЕННЯ КІБЕРБЕЗПЕКИ ТА ЗАХИСТУ КРИТИЧНОЇ ІНФОРМАЦІЙНОЇ ІНФРАСТРУКТУРИ УКРАЇНИ В УМОВАХ СУЧАСНИХ ГЛОБАЛЬНИХ ЗАГРОЗ ТА ВІЙСЬКОВОГО КОНФЛІКТУ

Іванченко Є.В.

*доктор технічних наук, професор,
директор інституту ННІКЗІ,*

Туровський О.Л.

*доктор технічних наук, професор,
завідувач кафедри ТСК,*

Рижаков М.М.

старший викладач, кафедри ТСК,

Державний університет інформаційно-комунікаційних технологій

Сучасний етап розвитку інформаційного суспільства характеризується безпрецедентною залежністю ключових сфер від цифрових технологій. Критична інформаційна інфраструктура (енергетика, зв'язок, транспорт, фінанси тощо) стала невід'ємною основою функціонування держави і суспільства. Водночас, глобальні

кіберзагрози та гібридні виклики стрімко зростають. Зокрема, гібридна війна дедалі частіше включає кібернетичний компонент – цілеспрямовані атаки на об'єкти критичної інфраструктури, кампанії дезінформації та інші засоби дестабілізації. Яскравим прикладом є війна Росії проти України, яка продемонструвала новий вимір сучасного воєнного конфлікту: поєднання традиційних бойових дій з масованими кібератаками на урядові системи, енергомережі, телекомунікації і бізнес-структури. У таких умовах питання забезпечення кібербезпеки набуває особливої актуальності, оскільки кіберпростір став полем бою, де атаки можуть призвести до фізичних наслідків.

Нинішні глобальні загрози характеризуються високим рівнем складності та новизною. Держави-агресори та організовані хакерські угруповання використовують найсучасніші засоби, включно зі штучним інтелектом, для пошуку вразливостей і реалізації атак. Наприклад, аналітичні дані вказують, що противник активно задіює інструменти на основі штучного інтелекту для автоматизації атак та вдосконалення фішингових кампаній. Відомо, що під час російсько-української війни кіберудари нерідко координуються з реальними ударами: атаки на енергосистему України супроводжувалися ракетними обстрілами, аби максимально порушити функціонування критичних служб. Впродовж 2022–2023 рр. зафіксовано серію масштабних кіберінцидентів: злам найбільшого телеком-оператора призвів до виведення з ладу 40% його інфраструктури, атаковано провідні підприємства паливно-енергетичного та фінансового секторів. Ці факти підкреслюють нагальну потребу у впровадженні новітніх методів кіберзахисту, здатних ефективно протидіяти як відомим, так і невідомим раніше загрозам.

Для України, що наразі протидіє воєнній агресії та гібридним атакам, питання захисту критичної інформаційної інфраструктури є питанням національної безпеки. Кібератаки на об'єкти енергетики, зв'язку, транспорту та державного управління можуть спричинити катастрофічні наслідки – від масштабних відключень електроенергії до порушення роботи екстрених служб і витоку конфіденційних даних. Агресор неодноразово демонстрував здатність використовувати кіберзброю для підриву стабільності: від атак типу «BlackEnergy» на енергомережі (2015 р.) до запуску руйнівного вірусу NotPetya (2017 р.), який уразив критичні сервіси. З початком повномасштабної війни у 2022 році масштаби та частота таких атак лише зросли. Таким чином, актуальність розробки новітніх методів кібербезпеки обумовлена необхідністю протистояти *сучасним глобальним загрозам*, серед яких домінують: цілеспрямовані атаки АРТ (Advanced Persistent Threat), кібершпигунство, деструктивні віруси-вимагачі та атаки на ланцюги поставок. Традиційні засоби кіберзахисту (сигнатурні антивіруси, класичні мережеві екрани, системи виявлення вторгнень на основі правил) більше не гарантують належного рівня захисту, особливо проти нульових вразливостей та новітніх технік. Причина полягає в їх реактивній суті: класичні системи покладаються на заздалегідь відомі шаблони атак і потребують частого оновлення вручну. В умовах динамічного кіберландшафту та постійно змінних тактик нападників це створює «вікно уразливості», коли нова атака може залишатися невиявленою. Натомість, застосування штучного інтелекту (ШІ) відкриває якісно нові можливості для забезпечення кібербезпеки. Алгоритми машинного навчання здатні аналізувати великі масиви даних у реальному часі,

виявляючи аномальні патерни, притаманні зловмисній активності, та миттєво реагувати на інциденти. Важливо, що такі системи працюють проактивно: вони не просто шукають відомі ознаки атаки, а прогнозують потенційно небезпечну поведінку, виходячи із виявлених відхилень від норми. В умовах війни і безперервних кібератак це дозволяє випереджати дії противника та знижувати навантаження на фахівців з безпеки. Зокрема, AI-асистовані системи кіберзахисту можуть автоматизувати рутинні процеси моніторингу, розпізнавати складні багатовекторні атаки, виявляти приховані вразливості і навіть частково брати на себе функції реагування (блокування трафіку, ізоляція вузлів тощо) без втручання людини. Відповідно до досліджень, інструменти на базі ШІ забезпечують більш високу ефективність виявлення загроз порівняно з традиційними: вони здатні розпізнати раніше невідомі атаки (такі як нові види шкідливого ПЗ або нетипову поведінку зловмисника), чого не досягають методи, основані лише на статичних правилах. Все це робить тему дослідження надзвичайно актуальною для зміцнення кіберстійкості України у сучасних умовах.

Основним напрямом розвитку новітніх технологій кібербезпеки є впровадження інтелектуальних систем виявлення та реагування. У рамках даного дослідження пропонується модель кіберзахисту, що поєднує аналіз поведінки користувачів/системи та технології машинного навчання для прогнозування навантаження і виявлення аномалій. Центральним компонентом є автоенкодер – глибока нейронна мережа, здатна навчитися компресувати та відновлювати дані, виділяючи в них основні приховані характеристики. Автоенкодер складається з двох частин: *кодера* (encoder), що згортає вхідну інформацію у стислу внутрішню репрезентацію, та *декодера* (decoder), що відновлює вихідні дані з цієї стислої представлення (рис. 1). Мережу тренують на нормальних (легітимних) даних таким чином, щоб вона мінімізувала різницю між оригіналом і відновленим сигналом. Після навчання автоенкодер здатний відтворювати типову поведінку системи; якщо ж в нових даних присутні незвичні відхилення, мережа не зможе точно їх реконструювати, і помилка відновлення зросте – це сигнал про аномалію. Таким чином, автоенкодер виконує роль датчика, що постійно контролює стан мережевого трафіку чи дій користувачів і підсвічує нетипові патерни для подальшого аналізу або автоматичної протидії.

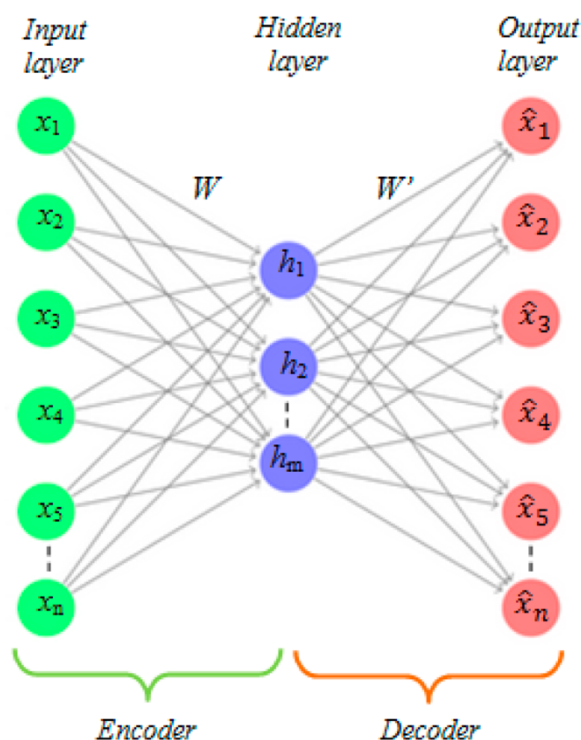


Рис. 1. Спрощена схема автоенкодера, що складається з кодеру (звужує багатовимірні вхідні дані до «прихованого» представлення) та декодеру (відновлює дані з цього представлення).

Мета навчання – досягти, щоб вихід $\hat{x}$ максимально відповідав вхідним даним x . У застосуванні до кіберзахисту - це дозволяє навчити модель на нормальній поведінці та виявляти відхилення, як потенційні загрози.

Запропонована модель користувача передбачає побудову профілю нормальної активності на основі даних про навантаження системи і дії користувачів, з використанням автоенкодера як інструмента для виявлення аномалій. Зокрема, враховується динаміка змін у часі: до архітектури автоенкодера інтегровано рекурентні шари (наприклад, мережу типу LSTM або GRU), які дозволяють моделі враховувати попередній контекст і прогнозувати очікуване навантаження на систему. Це означає, що модель не просто оцінює поточний стан, а й передбачає, яким він має бути, виходячи з минулих тенденцій (добових коливань трафіку, типових робочих циклів тощо). Якщо фактичні показники (інтенсивність запитів, обсяг мережевого трафіку, дії користувача) істотно відхиляються від прогнозованих, спрацьовує маркер аномалії. Такий підхід поєднує прогнозування навантаження зі своєчасним виявленням аномальних ситуацій, що особливо важливо для критичних систем, де навіть короточасні збої можуть мати серйозні наслідки.

Технічно модель функціонує наступним чином. На етапі навчання автоенкодер (з можливим включенням рекурентних блоків для часових рядів) тренується на масивах журналів подій, показників навантаження мережі та інших телеметричних даних, зібраних у нормальних умовах (без інцидентів безпеки). В ході оптимізації мережа автоматично виявляє найважливіші особливості нормальної активності та стискає їх у прихованих шарах. Важливо підкреслити, що навчання відбувається без учителя, тобто не потребує розмічених прикладів атак – це критична перевага, адже

в реальних умовах далеко не всі можливі атаки відомі заздалегідь або присутні у навчальних даних. Етап виявлення (експлуатації моделі) полягає у безперервному моніторингу: нові дані пропускаються через натренований автоенкодер, і обчислюється метрика помилки відновлення. Якщо ця помилка перевищує певний динамічний поріг – відповідний інцидент класифікується як потенційна атака чи аномалія, й ініціюються заходи реагування. Алгоритм може бути доповнений компонентом реагування в реальному часі: наприклад, при виявленні аномального сплеску трафіку система автоматично може перемкнути мережу в захищений режим, розширити пропускну здатність для поглинання можливого DDoS або ізолювати підозрілі хости. Така автоматизація значно підвищує швидкість і скоординованість кіберзахисту.

Запропонований підхід із використанням ШІ-моделі на базі автоенкодера демонструє високу ефективність у підвищенні рівня кібербезпеки критичних систем. По-перше, досягається покращення виявлення загроз. Автоматичний аналіз поведінкових аномалій дає змогу фіксувати навіть ті атаки, які не мають відомих сигнатур або маскуються під легітимну активність. Дослідження підтверджують, що автоенкодерні системи здатні успішно виявляти різноманітні типи інцидентів – від вторгнень у мережу до відмовостійких атак на сервери. Зокрема, у галузі енергетики було показано, що поєднання рекурентного автоенкодера з методами класифікації аномалій (локальний фактор сусідства LOF тощо) перевершує за точністю класичні алгоритми виявлення (такі як One-Class SVM, Isolation Forest). Іншими словами, глибока модель, навчена на характеристиках нормального трафіку, розпізнає кіберудари з вищою чутливістю, зменшуючи ймовірність пропущення атаки. По-друге, використання ШІ дозволяє значно прискорити реакцію на інциденти. Вбудовані механізми автоматичного реагування, інтегровані в нашу модель, дають змогу виконати первинні дії (блокування IP-адреси нападника, відключення ураженого вузла, перемикання на резервні системи) майже миттєво, ще до того як оператор вручну ідентифікує проблему. Така швидкодія критична при масованих або комбінованих атаках, коли рахунок йде на секунди. Важливою перевагою описаного підходу є його адаптивність та масштабованість. На відміну від традиційних рішень, що потребують регулярного оновлення сигнатур і правил, інтелектуальна система з автоенкодером здатна навчатися на нових даних постійно. Це означає, що по мірі розвитку інфраструктури, появи нових типів легітимного трафіку чи сервісів модель може бути донавчена або перебудована, щоб і надалі точно описувати нормальну поведінку. Внаслідок цього зберігається низький рівень хибних позитивних спрацювань, оскільки система підлаштовується під зміни у профілі роботи мережі. Одночасно, така гнучкість ускладнює життя нападникам: атаки, що могли б залишатися непоміченими завдяки зміні тактики, все одно будуть виявлені, адже модель реагує на будь-яку невідповідність очікуваному стану. Автоенкодер виконує роль інструмента аналізу великих даних без суттєвого залучення людини: він згортає багатомірні журнали подій до компактних ознак, що дозволяє фахівцям зосередитися на інтерпретації знайдених аномалій і стратегічному реагуванні, а не на ручному перегляді логів.

Варто зауважити, що перехід на такі інтелектуальні системи потребує врахування низки факторів. По-перше, необхідно забезпечити якість даних для

навчання моделі: «отруєння» даних (англ. data poisoning), коли зломисник навмисно вносить аномальні дані у тренувальний набір, може вплинути на точність автоенкодера. Тому при розгортанні рішення важливо впровадити механізми захисту навчального контенту та періодичної верифікації моделі. По-друге, слід комбінувати розроблену модель з іншими засобами кібербезпеки у багаторівневу оборону. Автоенкодерна система ефективно виявляє невідомі атаки, тоді як відомі типові загрози можуть швидше блокуватися традиційними засобами (фаєрволами, системами запобігання вторгнень). Отже, найбільшого ефекту можна досягти, інтегрувавши AI-модель в існуючу інфраструктуру безпеки, щоб отримати переваги обох підходів.

В умовах сучасних викликів, що постали перед Україною та світом, кібербезпека критичної інформаційної інфраструктури вимагає проактивних та інноваційних рішень. Проведений аналіз показав, що застосування штучного інтелекту, зокрема методів глибокого навчання на основі автоенкодерів, є перспективним напрямом для підсилення кіберзахисту. Розроблена модель, яка поєднує прогнозування навантаження і виявлення аномалій, демонструє високу ефективність у швидкому виявленні навіть нових загроз, автоматизації реагування та підтримці стабільної роботи систем під час атак. У науковій та практичній площині це означає суттєву перевагу над традиційними підходами: більш високий рівень виявлення, менше затримок в реагуванні, вища гнучкість. Для України впровадження таких технологій є важливою складовою зміцнення кіберстійкості держави та захисту стратегічних об'єктів від ворожих посягань. В подальших дослідженнях планується удосконалення запропонованої моделі, зокрема інтеграція технологій розподіленого штучного інтелекту та федеративного навчання для захисту від спроб отруєння даних, а також розгортання системи на реальних сегментах критичної інфраструктури. Таким чином, розвиток новітніх методів і технологій кібербезпеки на базі ШІ дозволить випереджати сучасні глобальні загрози та забезпечити надійний захист критичної інформаційної інфраструктури України навіть в умовах агресивного кібер-протистояння.

Список використаної літератури

1. Buczak, A. L., & Guven, E. (2016). *A Survey of Data Mining and Machine Learning Methods for Cyber Security Intrusion Detection*. *IEEE Communications Surveys & Tutorials*, 18(2), 1153–1176. <https://doi.org/10.1109/COMST.2015.2494502>
2. Chandola, V., Banerjee, A., & Kumar, V. (2009). *Anomaly Detection: A Survey*. *ACM Computing Surveys (CSUR)*, 41(3), 1–58. <https://doi.org/10.1145/1541880.1541882>
3. Kolosok, I. O., & Mazur, V. G. (2023). Застосування глибоких автоенкодерів для виявлення аномалій у трафіку промислових систем. *Науковий вісник НУ «Львівська політехніка». Комп'ютерні науки*, №4, 31–39.
4. Національний координаційний центр кібербезпеки при РНБО України. Аналітичний звіт про кіберзагрози 2022–2023 рр. <https://ncscc.gov.ua/>

ВИЯВЛЕННЯ ФЕЙКОВИХ МАТЕРІАЛІВ, ЗГЕНЕРОВАНИХ ШТУЧНИМ ІНТЕЛЕКТОМ, У СОЦІАЛЬНИХ МЕРЕЖАХ: ВИКЛИКИ ТА ПІДХОДИ ДО ПРОТИДІЇ

*Ільяшук С. С.
магістрант факультету інформаційних систем та технологій,
Милашенко В.М.
старший викладач кафедри кібербезпеки та захисту інформації,
ПВНЗ «Європейський університет»*

У контексті повномасштабної гібридної війни проти України особливу загрозу становить поширення дезінформації через соціальні медіа, яка набуває нових форм завдяки стрімкому розвитку генеративних моделей штучного інтелекту (ШІ), зокрема таких як GPT, DALL·E, Stable Diffusion та інших. Ці технології дозволяють створювати фотореалістичні зображення, переконливі тексти та синтезоване відео, які важко відрізнити від справжнього контенту неозброєним оком, що значно ускладнює традиційні методи верифікації інформації.

Системна протидія подібним викликам вимагає впровадження автоматизованих рішень, які спираються на складні алгоритми класифікації, стилOMETричний аналіз, методи виявлення ознак штучного походження контенту (наприклад, характерні артефакти генерації зображень або шаблонні мовні конструкції), а також використання нейромережевих детекторів, натренованих на контрастних вибірках синтетичних і справжніх матеріалів [1,2].

Особливої уваги потребують дослідження, спрямовані на вдосконалення методів семантичного аналізу та контекстного зіставлення повідомлень у медіапросторі, оскільки генеративні моделі можуть адаптуватися до тематичної й стилістичної специфіки цільових аудиторій, що ускладнює їх виявлення. Доцільним є також розширення напрямів досліджень у сфері мультимодального аналізу (аналіз одночасно тексту, зображення та звуку), що дозволяє виявляти фейкові повідомлення комплексно, на всіх рівнях їхнього подання [3].

На відміну від ручного фільтрування чи фактчекінгу, які обмежені масштабом і залежать від суб'єктивного сприйняття, пріоритетним є розвиток системних технічних рішень, зокрема інтеграції інструментів виявлення фейків безпосередньо в алгоритми соціальних платформ. Такий підхід передбачає тісну співпрацю між технічними командами, регуляторними органами та незалежними дослідницькими центрами. Зокрема, актуальним є запровадження маркування контенту, згенерованого ШІ, на рівні платформи, а також розробка відкритих API для взаємодії з незалежними верифікаторами [4].

У подальших дослідженнях доцільно зосередитися на створенні національних або регіональних баз даних фейкових прикладів, які могли б використовуватися для тренування спеціалізованих моделей в умовах локального контексту, адже більшість наявних рішень розробляються для англomовного середовища. Важливим також є удосконалення криптографічних технологій автентифікації оригінального контенту, зокрема шляхом імплементації стандартів цифрових водяних знаків та систем підтвердження джерела публікації [5].

Таким чином, подолання загроз, пов'язаних із поширенням фейкових матеріалів, створених за допомогою ШІ, вимагає комплексного підходу, що передбачає розвиток інфраструктури виявлення на технічному рівні, підтримку наукових досліджень у сфері аналізу даних і кіберзахисту, а також активну участь держави у формуванні політик цифрової безпеки.

Список використаної літератури

1. Adobe, Microsoft, & others. Content Authenticity Initiative. URL: <https://contentauthenticity.org>
2. Guarnera, L., Giudice, O., & Battiato, S. DeepFake Detection by Analyzing Convolutional Traces. *IEEE Conference on Computer Vision*.
3. Jawahar, G., Sagot, B., & Seddah, D. What Does BERT Learn about the Structure of Language? *ACL Anthology*. URL: <https://aclanthology.org/2020.acl-main.664>
4. OpenAI. GPT-4 System Card. <https://openai.com/research/gpt-4-system-card>
5. Wang, Z., She, J., & Dong, J. Multimodal Fake News Detection via Hierarchical Fusion with Knowledge Enhancement. *Information Processing & Management*.

МЕТОД ПОШУКУ БЕЗПЕЧНИХ ТРАЄКТОРІЙ РУХУ БЕЗПІЛОТНИХ АПАРАТІВ

*Калмиков А.М.
аспірант,
Доровський В.О.,
доктор технічних наук, професор
ПВНЗ «Європейський університет»*

Сучасні інформаційні системи відіграють ключову роль у забезпеченні автономної навігації та управління безпілотними літальними апаратами (БПЛА). Інтеграція методів пошуку безпечних траєкторій у такі системи дозволяє автоматизувати процес планування руху, підвищити надійність та безпечність експлуатації апаратів у різноманітних умовах — від міських територій до зон бойових дій.

В умовах швидкої цифровізації та розвитку технологій Інтернету речей (IoT) інформаційні системи стають все більш складними, інтегруючи у собі великі обсяги даних з сенсорів, супутників, радарів та інших джерел. Ефективна обробка цієї інформації для побудови безпечних траєкторій вимагає застосування новітніх алгоритмів штучного інтелекту, обчислювальної геометрії та теорії прийняття рішень [1].

Водночас зростає потреба в системах, які можуть працювати в режимі реального часу, адаптуватися до змін у навколишньому середовищі та виконувати складні завдання в автономному режимі. Тому розробка методів пошуку безпечних траєкторій у рамках інформаційних систем є вкрай актуальним завданням для забезпечення безпеки, ефективності та надійності застосування БПЛА в різних сферах.

Проблема динамічного планування безпечних траєкторій руху безпілотних апаратів (БА) є предметом активного дослідження у наукових колах, як в Україні, так і за кордоном. Основну увагу приділяють створенню адаптивних алгоритмів, які здатні забезпечити безпечний та узгоджений рух групи БА в динамічному та частково спостережуваному середовищі.

Значна кількість досліджень зосереджується на глобальних методах планування траєкторій, які базуються на апіорній інформації про середовище, зокрема про статичні перешкоди. Для цього часто застосовуються алгоритми Дейкстри, A^* , D^* та інші евристичні методи [2]. Однак, вони мають високу обчислювальну складність і не здатні ефективно працювати в режимі реального часу при змінних умовах.

Для вирішення задач реактивного перепланування все більшої популярності набувають локальні методи, серед яких варто виділити метод штучних потенційних полів (APF) та методи випадкової вибірки, такі як PRM (Probabilistic RoadMap) та RRT (Rapidly-exploring Random Tree). Попри свою ефективність, ці методи мають певні обмеження. Наприклад, APF-методи не завжди гарантують вихід з локальних мінімумів, а PRM вимагає значних попередніх обчислень.

На сьогодні дослідники все частіше звертаються до гібридних рішень, які поєднують сильні сторони різних підходів. Зокрема, поєднання RRT з потенційними полями дозволяє зменшити обчислювальну складність та підвищити гнучкість системи до змін середовища. До перспективних напрямів також належить використання нечіткої логіки, топологічних моделей простору, а також адаптивних механізмів на основі штучного інтелекту [4].

У результаті аналізу наукової літератури можна зробити висновок, що існуючі методи динамічного планування мають значний потенціал, однак ще не здатні повною мірою забезпечити узгоджене та безпечне управління спільним рухом БА в умовах високої динамічності та невизначеності середовища. Це визначає актуальність подальших досліджень у напрямку вдосконалення моделей і алгоритмів планування.

Проведене дослідження підтверджує високу актуальність задачі динамічного планування безпечних траєкторій спільного руху безпілотних апаратів (БА) в умовах часткової спостережуваності, динамічних змін середовища та численних обмежень. Використання великих гетерогенних груп БА для виконання складних місій, таких як промислове рибальство, вимагає високої точності, узгодженості руху та здатності системи до реактивного перепланування траєкторій в реальному часі.

Аналіз існуючих підходів засвідчив, що класичні методи глобального планування не задовольняють вимогам високої швидкодії та адаптивності. Натомість локальні та гібридні методи, зокрема поєднання RRT з методами штучних потенційних полів, демонструють перспективність у розв'язанні задачі в складних динамічних умовах. Однак, і ці методи потребують подальшого вдосконалення для зменшення обчислювальної складності та підвищення надійності результатів.

Таким чином, подальші дослідження мають бути зосереджені на розробці нових ефективних гібридних алгоритмів, здатних забезпечити швидко, точно та безпечно планування траєкторій у режимі реального часу із врахуванням просторової конфігурації групи БА та зовнішніх збурень.

Список використаної літератури

1. Білоус В.І., Іваненко І.Є. Дослідження безпеки руху безпілотних літальних апаратів у динамічному середовищі // Збірник наукових праць НА ДПСУ. – 2022. – № 3(59). – С. 56–64.
2. Кир'янов А.М. Децентралізоване управління груповим рухом безпілотних апаратів. – Дис. канд. техн. наук: 05.13.06. – КІП ім. Ігоря Сікорського. – Київ, 2020. – 180 с.
3. Федоров О.В., Мельник В.М. Моделі безпечного руху безпілотних апаратів в умовах невизначеності // Вісник технічного університету. – 2021. – № 2(58). – С. 112–119.
4. Рогальський В.М., Грінчишин О.С. Адаптивні методи планування траєкторій для групи безпілотних літальних апаратів в умовах реального часу // Науковий вісник Одеської політехніки. – 2020. – № 2(56). – С. 45–53.

ОГЛЯД МІЖНАРОДНИХ СТАНДАРТІВ УПРАВЛІННЯ КІБЕРБЕЗПЕКОЮ ДЛЯ МАЛИХ ПІДПРИЄМСТВ

Kim Ю.А.

студент кафедри безпеки інформаційних технологій

Коробейнікова Т.І.

к.т.н., доцент

доцент кафедри безпеки інформаційних технологій

Національний університет «Львівська політехніка»

Сучасні реалії цифровізації бізнесу висувають нові вимоги до управління інформаційною безпекою, зокрема в малих підприємствах. У цьому контексті актуальним є вибір ефективного, гнучкого та економічно доцільного підходу до впровадження систем кіберзахисту. Найбільш поширеними рамковими моделями в галузі є стандарт ISO/IEC 27001 та NIST Cybersecurity Framework.

ISO/IEC 27001 встановлює вимоги до побудови системи управління інформаційною безпекою (ISMS), орієнтованої на формалізацію процесів, документування процедур, проходження аудиту та сертифікації. Натомість NIST Cybersecurity Framework надає більш гнучку структуру, зосереджену на практичному управлінні ризиками та адаптації до змін.

Основними етапами впровадження систем управління кібербезпекою є наступні.

1. Оцінка ризиків: виявлення активів, загроз, вразливостей та розрахунок можливих наслідків.
2. Аудит безпеки: перевірка відповідності політик і процедур вимогам обраного стандарту.
3. Тестування на проникнення (penetration testing): моделювання атак з метою виявлення слабких місць.

Порівняльний аналіз показує, що ISO/IEC 27001 є більш формалізованим та ресурсозатратним, тоді як NIST орієнтований на адаптивність і практичність. Нижче у таблиці узагальнено ключові відмінності між підходами.

Таблиця 1.

Порівняльний аналіз підходів до управління інформаційною безпекою

Критерій	Загальний підхід	Визнання	Орієнтація	Документація	Вартість	Гнучкість	Підходить для МСП
ISO/IEC 27001	Формалізована ISMS	Міжнародний стандарт	Процесна модель	Необхідна для сертифікації	Висока (сертифікація, аудит)	Обмежена формальними вимогами	Обмежено через ресурси
NIST Cybersecurity Framework	Гнучка структура управління ризиками	Розроблений у США, поширюється глобально	Практичні функції: ідентифікація, захист, виявлення, реагування, відновлення	Менший акцент на документацію	Доступніший варіант для МСП	Легко адаптується до змін	Оптимальний вибір

Особливу увагу слід приділити впровадженню NIST Cybersecurity Framework у зв'язку з його адаптивністю до нових технологій. Наприклад, інтеграція з мережевими контролерами, такими як Omada, вимагає динамічного підходу до управління безпекою. NIST дозволяє організаціям гнучко реагувати на оновлення програмного забезпечення, появу нових загроз та масштабування інфраструктури без зайвої бюрократії.

Таким чином, у випадку малого підприємства з обмеженими ресурсами, NIST Cybersecurity Framework є доцільним вибором. Його структура дозволяє забезпечити належний рівень кіберзахисту при мінімізації витрат на формалізацію та сертифікацію. Це дає можливість зосередитися на впровадженні практичних заходів безпеки, орієнтованих на конкретні потреби організації.

Список використаної літератури

1. ISO/IEC 27001:2013 Information technology – Security techniques – Information security management systems – Requirements.
2. NIST Framework for Improving Critical Infrastructure Cybersecurity, Version 1.1, 2018.
3. Whitman, M. E., Mattord, H. J. (2018). Principles of Information Security. Cengage Learning.
4. Von Solms, R., Van Niekerk, J. (2013). From information security to cyber security. Computers & Security.
5. Tipton, H. F., Krause, M. (2007). Information Security Management Handbook. Auerbach Publications.

DEVSECOPS ДЛЯ УНІВЕРСИТЕТСЬКИХ СИСТЕМ: KEYС ВПРОВАДЖЕННЯ У CRM ДЕКАНАТУ

*Ковальов О.С.,
магістрант факультету інформаційних систем та технологій,
ПВНЗ «Європейський університет»*

В умовах військової агресії, яка триває в Україні, питання кібербезпеки набуває надзвичайної актуальності. Освітні установи, зокрема університети, все частіше потрапляють у поле зору кіберзлочинців. Ці інституції обробляють значний обсяг конфіденційних даних: особисту інформацію студентів і викладачів, результати навчання, фінансові документи, а також технічну документацію. Однією з найвразливіших ланок у цій системі є CRM-системи деканатів, які, будучи інтегрованими з іншими сервісами ЗВО, можуть виступати воротами до всієї IT-інфраструктури університету.

Традиційні методи забезпечення інформаційної безпеки часто виявляються неефективними в умовах динамічного розвитку загроз та високих темпів зміни коду. У зв'язку з цим, підходи типу DevSecOps – інтеграція безпеки на всіх етапах життєвого циклу розробки – виглядають більш перспективними та доцільними для застосування у вищих навчальних закладах [2].

Метою цієї роботи є розробка та демонстрація практичного підходу до впровадження DevSecOps у CRM-систему деканату Європейського університету як кейс для захисту освітньої IT-інфраструктури.

У рамках дослідження було визначено основні загрози, що виникають у процесі розробки, тестування та експлуатації CRM-систем. Найпоширенішими з них є:

- витік конфіденційної інформації через hardcoded credentials у коді;
- вразливості у контейнеризованому середовищі;
- неправильне керування ролями доступу (RBAC);
- відсутність або недостатність журналювання та моніторингу змін;
- відсутність інструментів політик доступу на основі принципів найменших привілеїв.

Для мінімізації вказаних ризиків було розгорнуто експериментальний пайплайн CI/CD із реалізацією наступних безпекових етапів:

- статичний аналіз коду Python за допомогою Semgrep;
- сканування контейнерів Docker інструментами Trivy і Dockle;
- виявлення секретів і ключів доступу за допомогою Gitleaks;
- перевірка інфраструктурного коду на основі Terraform через tfsec;
- впровадження LDAP-автентифікації та аудит ролей доступу;
- моніторинг подій і журналювання через Grafana, Loki та Falco;
- підписування артефактів та використання принципу Zero Trust.

Інтеграція всіх компонентів виконувалась через Jenkins та GitHub Actions із подальшою автоматизованою перевіркою на кожному етапі розгортання змін. Окрім цього, застосовано систему перевірки оновлень open-source-бібліотек, що використовуються у проєкті, на предмет вразливостей (Dependabot, safety-checks).

У процесі реалізації було встановлено, що DevSecOps дозволяє досягнути не лише високого рівня автоматизації виявлення загроз, а й сприяє формуванню культури відповідального ставлення до безпеки серед студентів, які залучені до розробки та супроводу систем.

Проведені експерименти довели ефективність запропонованої моделі: зменшилась кількість інцидентів, зросла швидкість і якість відгуку на вразливості, покращилась прозорість змін у кодовій базі. Впроваджені інструменти забезпечили постійний зворотній зв'язок між командами розробки, адміністрування та безпеки. Також позитивним результатом є скорочення кількості ручних перевірок, що унеможливорює людські помилки.

Запропонована архітектура може бути адаптована не лише для CRM, а й для інших освітніх сервісів: електронних журналів, порталів вступника, розкладу занять, бібліотечних систем тощо. Університети, що інтегрують DevSecOps у свою інфраструктуру, отримують змогу швидко реагувати на кіберзагрози, відповідати вимогам нормативно-правових актів та підвищити загальний рівень довіри до власної ІТ-системи.

Список використаної літератури

1. *Aqua Security. Trivy: Vulnerability Scanner for Containers.* URL: <https://github.com/aquasecurity/trivy>
2. *Leonid Arsenovych, Oleksandr Nikolaievsky, Olena Skliarenko, Leonid Lytvynenko, Ivan Kydriavskyi. Organization of Training with the Use of Digital Technologies for Ensuring Cybersecurity in the Educational Space // WSEAS Transactions on Computer Research, ISSN / E-ISSN: 1991-8755 / 2415-1521, Volume 12, 2024, p. 524-536, DOI: 10.37394/232018.2024.12.51 , https://wseas.com/journals/articles.php?id=9988*
3. *NIST SP 800-53 Rev. 5. Security and Privacy Controls for Information Systems and Organizations. NIST, 2020.*
4. *OWASP DevSecOps Guideline. OWASP Foundation.* URL: <https://owasp.org/www-project-devsecops-guideline/>

ШТУЧНИЙ ІНТЕЛЕКТ В КОМП'ЮТЕРНІЙ ГРАФІЦІ

*Корбут В.В.
магістрант факультету інформаційних систем та технологій,
ПВНЗ «Європейський університет»*

Розвиток технологій штучного інтелекту (ШІ) кардинально змінює підходи до генерації, обробки та візуалізації графічного контенту. На стику комп'ютерної графіки та глибинного навчання формується новий напрям — інтелектуальна візуалізація, що охоплює автоматизацію творчих процесів, генеративне моделювання та реалістичну симуляцію середовищ. У цій роботі розглянуто ключові напрями інтеграції ШІ в комп'ютерну графіку, оцінено перспективи застосування та ідентифіковано потенційні загрози.

Одним із найзначущих проривів є застосування генеративно-змагальних мереж (GANs). Ці моделі навчаються генерувати зображення, текстури, обличчя або цілі сцени, які складно відрізнити від справжніх. Зокрема, відомі моделі StyleGAN від компанії NVIDIA використовуються для створення реалістичних віртуальних персонажів у відеоіграх, рекламі, кінематографі. Вони також використовуються в проектах реконструкції історичних облич або синтезу медичних зображень для діагностики.

Іншим важливим напрямом є дискретизація та синтез графіки з допомогою трансформерних моделей і дифузійних алгоритмів, зокрема Latent Diffusion Models (LDM) та Stable Diffusion. Ці системи дозволяють створювати зображення на основі текстових запитів (text-to-image generation), відкриваючи нові горизонти у візуальному сторітелінгу, дизайні та генерації графіки на замовлення. Наприклад, художники більше не змушені створювати сотні ітерацій вручну — ШІ може надати варіанти дизайну за лічені секунди.

У сфері анімації та тривимірного моделювання ШІ використовується для автоматизованого створення рухів персонажів, прогнозування положення тіла в динамічних сценах, а також для оптимізації поведінки персонажів на основі навчання з підкріпленням. Наприклад, ігрові студії активно застосовують системи на основі Motion Capture + AI, які дозволяють адаптивно анімувати персонажів залежно від середовища або реакції користувача.

Не менш важливим є використання ШІ у візуалізації фізично-коректних сцен (Physically Based Rendering). Алгоритми машинного навчання дозволяють оптимізувати трасування променів (ray tracing), передбачаючи шляхи відбиття світла в сценах із великою кількістю об'єктів або складною геометрією. Це значно знижує обчислювальні витрати при збереженні високої якості.

Разом із цим виникає проблема зловживань. ШІ також застосовується для створення deepfake-контенту, що сприяє поширенню дезінформації, фальсифікації відео- або фотофактів. Це зумовлює розвиток технологій детекції ШІ-згенерованого контенту та захисту цифрової автентичності. Розробка алгоритмів, здатних розрізнити штучний і справжній контент, стає ключовим завданням для інформаційної безпеки.

У тезах проаналізовано сучасний стан застосування штучного інтелекту в комп'ютерній графіці, розглянуто основні технології, їхні переваги та обмеження. Також зроблено акцент на потенційні виклики в етичному, правовому та інформаційному аспектах.

Штучний інтелект істотно трансформує комп'ютерну графіку, відкриваючи нові можливості для автоматизації творчих процесів, генерації візуального контенту та оптимізації рендерингу. Генеративні моделі, дифузійні алгоритми та нейронні мережі дозволяють досягати високої якості графіки при зниженні трудомісткості.

Разом з перевагами виникають ризики: використання ШІ для створення маніпулятивного контенту, проблема автентичності та етичні виклики. Подальший розвиток потребує балансу між інноваціями та безпекою в цифровому середовищі.

Перспективи подальших досліджень пов'язані з інтеграцією ШІ у кросдисциплінарні проекти, розвитком інтерфейсів «людина–машина», адаптивного дизайну, індустріального візуалізування і збереження цифрової спадщини.

Список використаної літератури

1. Rombach R. et al. High-Resolution Image Synthesis with Latent Diffusion Models // *arXiv:2112.10752* – 2022.
2. Wang Y. et al. Motion Diffusion Model: Human Motion Generation with Diffusion Model // *arXiv:2209.14916* – 2022.
3. Величко, С. В. Застосування алгоритмів машинного навчання у генерації графічного контенту // *Системи обробки інформації*. – 2023. – Вип. 1 (175). – С. 45–49.

СТРАТЕГІЇ ПОСИЛЕННЯ КІБЕРБЕЗПЕКИ КОРПОРАТИВНИХ ДАНИХ ПРИ ВИКОРИСТАННІ SALESFORCE PLATFORM: БАГАТОРІВНЕВИЙ ПІДХІД

*Коцун В.І.,
кандидат технічних наук, доцент,
доцент кафедри комп'ютерних наук та програмної інженерії
Сушинський О.Є.,
доктор технічних наук, доцент,
завідувач кафедри комп'ютерних наук та програмної інженерії
ПВНЗ «Європейський університет»*

У сучасному цифровому середовищі захист корпоративних даних став критично важливим завданням для організацій, що використовують Salesforce Platform [1]. Багаторівневий підхід до кібербезпеки забезпечує комплексний захист інформації та мінімізує ризики несанкціонованого доступу. Впровадження ефективних стратегій захисту вимагає розуміння всіх рівнів безпеки та їх взаємодії в рамках екосистеми Salesforce (Рис.1).

Фундаментальний рівень безпеки включає базові механізми захисту, які формують основу системи кібербезпеки. На цьому рівні впроваджуються надійні паролі, регулярне їх оновлення та використання генераторів випадкових паролів. Важливим елементом є налаштування базових параметрів безпеки організації, включаючи політики паролів, обмеження IP-діапазонів та налаштування сесій користувачів. Також на цьому рівні встановлюються основні правила розмежування доступу та визначаються ролі користувачів.

Впровадження багатофакторної автентифікації (MFA) є обов'язковим елементом сучасної системи безпеки. Salesforce пропонує різні методи другого фактору автентифікації, включаючи SMS-коди, автентифікатори та біометричні дані. На цьому рівні також налаштовується Single Sign-On (SSO) для централізованого управління доступом та інтеграції з корпоративними системами ідентифікації. Важливим аспектом є впровадження адаптивної автентифікації, яка враховує контекст доступу та рівень ризику [2].

Шифрування даних забезпечує додатковий рівень захисту конфіденційної інформації. Salesforce Shield Platform Encryption дозволяє шифрувати дані як під час зберігання, так і при передачі. На цьому рівні впроваджуються політики управління ключами шифрування, включаючи ротацію ключів та їх безпечне зберігання.

Особлива увага приділяється шифруванню критично важливих полів та документів, а також налаштуванню правил доступу до зашифрованих даних [3].

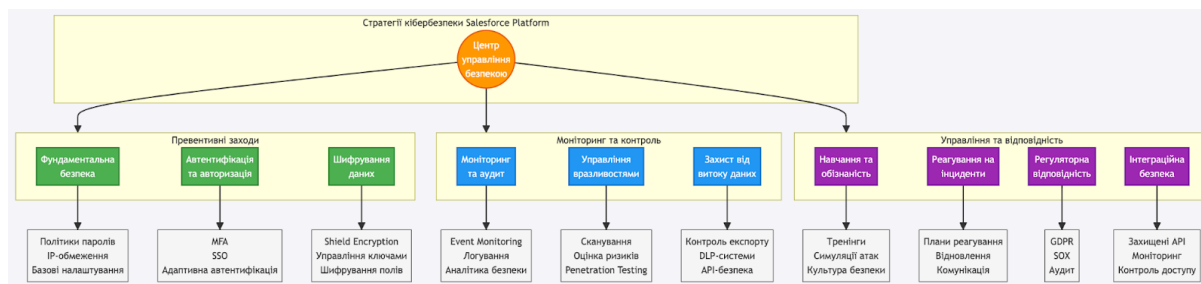


Рис.1. Стратегія кібербезпеки Salesforce Platform.

Постійний моніторинг активності користувачів та системних подій є критично важливим для виявлення потенційних загроз. Salesforce Event Monitoring забезпечує детальне логування всіх дій у системі, включаючи доступ до даних, зміни конфігурації та виконання запитів. На цьому рівні налаштовуються системи раннього попередження про підозрілу активність та автоматизовані сповіщення про порушення безпеки. Регулярний аудит логів та аналіз патернів використання допомагають виявляти потенційні вразливості.

Проактивне управління вразливістю включає регулярне сканування системи на наявність потенційних загроз, оцінку ризиків та впровадження необхідних оновлень безпеки. На цьому рівні здійснюється моніторинг повідомлень про нові вразливості від Salesforce Security та впровадження рекомендованих заходів захисту. Важливим аспектом є також регулярне тестування на проникнення та оцінка ефективності існуючих механізмів захисту.

Запобігання витоку даних вимагає впровадження спеціалізованих інструментів та політик. Salesforce надає можливості для налаштування правил експорту даних, обмеження масового завантаження та контролю за діями користувачів з критичною інформацією. На цьому рівні впроваджуються системи класифікації даних, політики використання зовнішніх пристроїв та контроль за передачею інформації через API.

Підвищення обізнаності користувачів у питаннях кібербезпеки є критично важливим елементом загальної стратегії захисту. Регулярні тренінги, симуляції фішингових атак та навчання з питань безпечного використання системи допомагають формувати культуру кібербезпеки в організації. На цьому рівні розробляються навчальні матеріали, проводяться практичні заняття та оцінюється ефективність навчальних програм [4].

Ефективне реагування на інциденти безпеки вимагає наявності чітко визначених процедур та інструментів. На цьому рівні розробляються плани реагування на різні типи інцидентів, визначаються ролі та відповідальності команди реагування, створюються процедури комунікації та відновлення після інцидентів. Важливим елементом є також регулярне тестування планів реагування та їх оновлення на основі отриманого досвіду [5].

Забезпечення відповідності різним регуляторним вимогам та стандартам безпеки є важливим аспектом загальної стратегії захисту. На цьому рівні впроваджуються механізми контролю відповідності, проводиться регулярний аудит

та підготовка необхідної документації. Особлива увага приділяється відповідності вимогам GDPR, SOX та інших релевантних стандартів.

Багаторівневий підхід до кібербезпеки в Salesforce Platform забезпечує комплексний захист корпоративних даних та мінімізує ризики кіберзагроз. Успішна реалізація всіх рівнів захисту вимагає постійної уваги, регулярного оновлення та адаптації до нових викликів у сфері кібербезпеки. Важливо розуміти, що кожен рівень безпеки є критично важливим елементом загальної системи захисту, і тільки їх комплексне впровадження може забезпечити надійний захист корпоративних даних.

Список використаної літератури

1. Johnson, R., & Martinez, K. (2024). «Multi-layered Security Strategies in Enterprise Cloud Solutions: Salesforce Platform Case Study.» *Journal of Cloud Computing Security*, 15(1), 23-45.
2. Коваленко, М.С., Іванчук, О.В. (2023). «Інтегровані системи захисту корпоративних даних у хмарних середовищах.» *Інформаційна безпека підприємства*, 8(4), 156-172.
3. Zhang, L., & Peterson, S. (2024). «Advanced Security Configurations in Salesforce: From Implementation to Monitoring.» *International Journal of Information Security*, 18(2), 89-107.
4. Wilson, M., Thompson, R., & Baker, C. (2023). «Corporate Data Protection Strategies: A Comprehensive Guide to Salesforce Security Features.» *Enterprise Information Systems Security*, 11(3), 234-256.
5. Harris, E., & Lee, S. (2024). «Zero Trust Architecture in Salesforce Implementations: A Strategic Approach.» *Cybersecurity Management Review*, 9(1), 45-67.

НАПРЯМКИ ОПТИМІЗАЦІЇ МІКРОСЕРВІСНОЇ АРХІТЕКТУРИ

*Кулагін В.П.,
аспірант,
ПВНЗ «Європейський університет»*

Мікросервісна архітектура (Microservice Architecture) являє собою підхід до розробки програмного забезпечення, за якого система формується як сукупність малих, автономних сервісів, що взаємодіють між собою за допомогою легких інтерфейсів. Кожен із таких сервісів виконує окрему бізнес-функцію, розробляється і розгортається незалежно від інших, що сприяє більшій гнучкості у підтримці та масштабуванні системи. Такий підхід сприяє прискоренню процесів розробки та оновлення продукту, дозволяючи невеликим командам паралельно працювати над окремими сервісами, відкриває можливості для кращої масштабованості та стійкості до збоїв завдяки ізольованості кожного компоненту. Архітектура такого типу підтримує високий рівень гнучкості, автономності та масштабованості компонентів, надаючи можливість розгортання кожного сервісу у власному середовищі - контейнері, віртуальній машині чи навіть на фізичному сервері. Це дозволяє ефективно локалізувати збої та зменшувати їхній вплив на функціонування всієї системи, оскільки відмова одного сервісу не паралізує роботу решти компонентів [1]. Але із збільшенням кількості таких незалежних елементів виникають нові

труднощі, зокрема щодо управління ресурсами, забезпечення безперебійної комунікації між сервісами, гарантування безпеки та ефективного оновлення системи. Оптимізація мікросервісної архітектури є надзвичайно важливою та водночас складною тематикою, яка охоплює широкий спектр технічних, організаційних та інфраструктурних аспектів, адже вона безпосередньо впливає на ефективність, масштабованість, продуктивність, надійність і витрати на обслуговування програмних систем.

Одним із ключових напрямків оптимізації є балансування навантаження між мікросервісами. Сучасні алгоритми, такі як Round Robin, Least Connections, а також адаптивні методи із застосуванням машинного навчання, дозволяють розподілити запити таким чином, що забезпечується рівномірне використання обчислювальних ресурсів і зменшуються затримки [2]. Це важливо для систем із високою пропускнуою здатністю, де навіть невеликі оптимізації можуть мати суттєвий вплив на загальну продуктивність.

Управління конфігурацією також виступає важливим напрямком оптимізації. Сучасні рішення, що базуються на централізованих сервісах (Spring Cloud Config, Consul, Etcd), забезпечують автоматизоване оновлення налаштувань, що є критичним для динамічно масштабованих середовищ [3]. Використання таких систем дозволяє зменшити людський фактор, який може призводити до помилок при ручному керуванні конфігурацією, і забезпечує узгодженість роботи всіх компонентів системи.

Ще одним перспективним напрямком є автоматизований моніторинг та діагностика з використанням алгоритмів машинного навчання. Сучасні підходи дозволяють аналізувати лог-файли та часові ряди для виявлення аномалій і прогнозування потенційних збоїв у роботі мікросервісів [4]. Таке рішення забезпечує своєчасне реагування на проблеми та мінімізує час відновлення після збоїв.

Важливим аспектом оптимізації є оптимізація моделей комунікації між сервісами. Різні протоколи взаємодії (REST, gRPC, GraphQL) мають свої переваги та недоліки з точки зору латентності, пропускнуої здатності та ефективності використання ресурсів [5]. Проведення експериментальних досліджень для визначення оптимального протоколу взаємодії дозволить адаптувати систему до специфічних умов експлуатації.

Також необхідно звернути увагу на забезпечення безпеки та відмовостійкості. Впровадження патернів, таких як Circuit Breaker, Retry та Bulkhead, дозволяє ізолювати збої окремих сервісів та зменшити ризик каскадних помилок. Розробка власних моделей безпеки, що адаптуються до характеристик кожного окремого сервісу, дозволяє підвищити загальний рівень захисту системи, зменшуючи потенційну атаковану поверхню [1], [3].

Кожен з розглянутих напрямків оптимізації вимагає глибокого аналізу, моделювання та практичної перевірки. Сучасні дослідження демонструють, що інтеграція алгоритмів машинного навчання для прогнозування збоїв, адаптивне балансування навантаження та автоматизація процесів управління конфігурацією є надзвичайно перспективними. Водночас необхідно враховувати додаткове ускладнення, пов'язане з підвищенням кількості точок входу та збільшенням потоку

інформації між сервісами, що вимагає посиленних заходів безпеки та систем моніторингу.

Таким чином, оптимізація мікросервісної архітектури є складною задачею, яка охоплює не лише технічні аспекти, але й включає питання управління командами, процесами розробки та підтримки. Удосконалення підходів до оптимізації має вирішальне значення для забезпечення високої ефективності, стійкості до збоїв, кібербезпеки та керованості сучасних розподілених програмних систем. Подальші наукові дослідження та практичні впровадження у цій сфері відкривають нові можливості для розвитку складних ІТ-інфраструктур, підвищення якості сервісів та зменшення витрат на їх обслуговування.

Список використаної літератури

1. G. D. Maayan, «5 Ways to Reduce the Attack Surface for Microservices,» [Online]. Available: <https://securityboulevard.com/2023/04/5-ways-to-reduce-the-attack-surface-for-microservices>.
2. J. Johnson and L. Shiff, «What Is Microservice Architecture? Microservices Explained,» [Online]. Available: <https://www.bmc.com/blogs/microservices-architecture>.
3. «8 Ways to Secure Your Microservices Architecture,» [Online]. Available: <https://www.okta.com/resources/whitepaper/8-ways-to-secure-your-microservices-architecture>.
4. C. Walia, Mastering Cloud-Native Microservices: Designing and Implementing Cloud-Native Microservices for Next-Gen Apps, BPB Publications, 2023, p. 298.
5. S. Newman, Building Microservices: Designing Fine-Grained Systems, O'Reilly, 2021, p.

ДОСЛІДЖЕННЯ АЛГОРИТМІВ ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ ВИЯВЛЕННЯ АНОМАЛЬНОЇ ПОВЕДІНКИ ГРАВЦІВ В ОНЛАЙН-ІГРАХ

*Курючкін А.Е.,
магістрант факультету інформаційних систем та технологій,
ПВНЗ «Європейський університет»*

В умовах сучасної цифрової економіки онлайн-ігри займають важливе місце серед розваг та комунікаційних платформ. Зі зростанням популярності масових багатокористувацьких ігор (MMORPG) та кіберспорту з'являються нові виклики, пов'язані з виявленням шахрайської та аномальної поведінки гравців. Створення ефективних систем виявлення аномалій є надзвичайно важливим для підтримки чесності гри та захисту від шахрайських атак, таких як використання чітів або ботів. У зв'язку з цим, алгоритми штучного інтелекту (ШІ) відіграють ключову роль у розробці таких систем.

Одним з підходів до виявлення аномальної поведінки є використання методів глибокого навчання, які дозволяють автоматично виявляти відхилення від нормального поведінкового патерну гравців. Наприклад, згідно з дослідженням Benedict Wilkins, Chris Watkins та Kostas Stathis (2020), один із найбільш ефективних методів включає використання метрик навчання для виявлення аномалій у відеоіграх, зокрема для аналізу взаємодії гравців з грою та їх поведінки в межах гри [1].

Штучний інтелект також застосовується для виявлення шахрайських дій у багатокористувацьких шутерах. Стаття «Cheat Detection in a Multiplayer First-Person

Shooter Using Artificial Intelligence» (2021) описує методи, за допомогою яких ШІ здатен виявляти використання читів у реальному часі, аналізуючи відхилення від звичних патернів руху та взаємодії гравців з ігровими об'єктами [2].

У свою чергу, застосування пояснювального ШІ дозволяє підвищити прозорість алгоритмів детекції, що особливо важливо для розробників ігор та гравців. Стаття «Explainable AI for Cheating Detection and Churn Prediction in Online Games» (2024) демонструє використання таких алгоритмів для прогнозування відтоку гравців і виявлення шахрайських дій, зокрема аналізуючи аномалії в поведінці користувачів [3].

Не менш важливою є роль аналізу внутрішньоігрових даних, які допомагають виявляти аномалії у поведінці гравців, особливо в еспортівних іграх. Так, у статті «Anomaly Detection in eSport Games Through Periodical In-Game Data» (2024) представлений метод, що дозволяє використовувати періодичний моніторинг даних для виявлення відхилень у поведінці гравців, що може свідчити про шахрайство або використання ботів [4].

Крім того, розробники ігор застосовують спеціальні фреймворки для виявлення ботів у рольових іграх, використовуючи методи оцінки поведінки та відстеження аномалій. Одним з таких методів є фреймворк, розроблений в рамках дослідження «Behavior Evaluation and Anomaly Tracking: Game Bot Detection Framework in RPG Games» (2023), який дозволяє ефективно ідентифікувати боти на основі аналізу їх поведінки в грі [5].

І, зрештою, систематичний огляд методів виявлення шахрайства в MMORPG та інших онлайн-іграх представлений у дослідженні «Cheating and Detection Method in Massively Multiplayer Online Role-Playing Game: Systematic Literature Review» (2024). Це дослідження надає вичерпну інформацію про різноманітні підходи до виявлення шахрайства в онлайн-іграх та застосування штучного інтелекту в цій сфері [6].

Застосування ШІ для виявлення аномальної поведінки в онлайн-іграх є потужним інструментом для боротьби з шахрайством. Алгоритми, що використовують методи глибокого навчання, пояснювального ШІ та моніторингу внутрішньоігрових даних, дозволяють досягти високої точності виявлення шахрайства та забезпечити безпеку та справедливість гри [7].

Список використаної літератури

1. Benedict Wilkins, Chris Watkins, Kostas Stathis «A Metric Learning Approach to Anomaly Detection in Video Games» // IEEE Conference on Games. – 2020. – С. 216-222.
2. Cheat Detection in a Multiplayer First-Person Shooter Using Artificial Intelligence // *ACM Digital Library*. – 2021. – С. 105-113.
3. Explainable AI for Cheating Detection and Churn Prediction in Online Games// *ResearchGate*. – 2024. – С. 1-9.
4. Anomaly Detection in eSport Games Through Periodical In-Game Data// *ScitePress*. – 2024. – С. 52-58.
5. Behavior Evaluation and Anomaly Tracking: Game Bot Detection Framework in RPG Games// *ACM Digital Library*. – 2023. – С. 45-50.
6. Cheating and Detection Method in Massively Multiplayer Online Role-Playing Game: Systematic Literature Review // *ResearchGate*. – 2024. – С. 1-10.

7. Скляренко, О., Покидько, Д. (2023). Методологія розробки навчальних цифрових ресурсів у вигляді онлайн ігор. *Електронне фахове наукове видання «Кибербезпека: освіта, наука, техніка»*, 2(22), 249–256. <https://doi.org/10.28925/2663-4023.2023.22.249256>

ВПЛИВ РОЗВИТКУ ІНФОРМАЦІЙНИХ ТЕХНОЛОГІЙ НА МЕТОДИ ТА ЗАСОБИ НАВЧАННЯ СТУДЕНТІВ

*Левченко С.В.,
старший викладач кафедри комп'ютерних
наук та програмної інженерії,
Савченко Я.О.,
викладач циклової комісії комп'ютерних
наук та програмної інженерії
ПВНЗ «Європейський університет»*

На сучасному етапі розвитку ІТ відбувається суттєва трансформація освітніх технологій, що змушує заклади вищої освіти переглядати традиційні педагогічні підходи. Використання віртуальних середовищ, хмарних сервісів, систем управління навчанням (LMS), а також штучного інтелекту сприяє підвищенню доступності, гнучкості та інтерактивності освітнього процесу. Все це зумовлено потребою у формуванні нових освітніх стратегій, здатних забезпечити якісну підготовку фахівців у цифрову епоху.

Інформаційні технології стали потужним ресурсом для розвитку гнучких форм навчання. Онлайн-курси, мобільні додатки, гейміфікація, віртуальна та доповнена реальність сприяють підвищенню мотивації студентів, залученню до активного пізнання та розвитку практичних навичок. Зокрема, інтерактивні платформи дозволяють студентам навчатися у власному темпі, в будь-який час і з будь-якого місця, що особливо актуально в умовах змішаного та дистанційного навчання.

Університети, зокрема, Європейський університет, адаптують свої освітні стратегії шляхом впровадження нових цифрових засобів навчання, таких як штучний інтелект для персоналізованих рекомендацій, аналітичні системи для відстеження прогресу студентів, симулятори та віртуальні лабораторії. Це дозволяє оптимізувати освітній процес, посилити індивідуальний підхід до кожного студента, а також підвищити ефективність викладання [1].

Дуже важливою є роль держави у цифровій трансформації освіти. Підтримка впровадження ІТ у навчальні заклади, фінансування цифрової інфраструктури, забезпечення кібербезпеки, а також підготовка педагогічних кадрів — є критичними чинниками успішної модернізації системи вищої освіти. Крім того, міжнародне співробітництво та обмін цифровими ресурсами сприяють інтеграції української освіти у глобальний освітній простір.

Попри значні переваги, впровадження інформаційних технологій супроводжується низкою викликів в умовах воєнного стану. Серед основних — цифрова нерівність серед студентів, відсутність стабільного інтернет-з'єднання, низький рівень цифрової грамотності частини викладачів та студентів. Також

викликає занепокоєння проблема академічної доброчесності в онлайн-середовищі, що вимагає нових підходів до оцінювання знань.

Ще одним викликом є збереження балансу між технологічною ефективністю та людським фактором у навчанні. Надмірна автоматизація може призвести до зменшення живої взаємодії між викладачем і студентом, що є важливою складовою освітнього процесу. Саме тому ключовим завданням є розробка змішаних моделей навчання, які поєднують переваги цифрових інструментів із традиційними педагогічними методами.

Загалом, розвиток інформаційних технологій кардинально змінює підходи до організації та реалізації освітнього процесу. Вони створюють нові можливості для інноваційного навчання, стимулюють академічну мобільність, готують студентів до викликів цифрової економіки. Цифрова трансформація освіти стає не лише відповіддю на виклики сьогодення, а й потужним рушієм формування майбутньої освітньої системи.

Список використаної літератури

1. Скляренко О.В., Ягодзінський С.М., Ніколаєвський О.Ю., Невзоров А.В. Цифрові інтерактивні технології навчання як невід'ємна складова сучасного освітнього процесу // Науковий журнал «Інноваційна педагогіка», Випуск 68, Т.2. – 2024. - с. 250-257.
2. Bond, M., Zawacki-Richter, O., & Nichols, M. (2021). Revisiting five decades of educational technology research: A content and authorship analysis of the British Journal of Educational Technology. *British Journal of Educational Technology*, 52(1), 55–73.
<https://doi.org/10.1111/bjet.13030>
3. UNESCO (2022). *Technology in education: A tool on whose terms? Global Education Monitoring Report*. <https://unesdoc.unesco.org/ark:/48223/pf0000382062>

ПОРІВНЯННЯ АРХІТЕКТУР НЕЙРОННИХ МЕРЕЖ ДЛЯ ВИЯВЛЕННЯ АНОМАЛІЙ У ЧАСОВИХ РЯДАХ У КОНТЕКСТІ КІБЕРБЕЗПЕКИ

*Литвин А.А.,
аспірант,
ПВНЗ «Європейський університет»*

Використання глибокого навчання у сфері кібербезпеки стає все більш поширеним підходом до виявлення шахрайських дій та кібератак. Однією з ключових проблем у цій галузі є обробка даних часових рядів, які мають структуровану природу, але відрізняються від класичних табличних даних. Це створює певні складнощі, оскільки традиційні методи машинного навчання, розроблені для табличних даних, не можуть бути застосовані безпосередньо. Наукова спільнота визнає специфіку цієї проблеми, про що свідчить зростаюча кількість досліджень у цьому напрямку. Так, у роботі Chen та ін. (2025) проведено систематичний огляд розвитку методів глибокого навчання для виявлення фінансових шахрайств, який демонструє значний прогрес у створенні сучасних

архітектур нейронних мереж, спеціально адаптованих для обробки часових рядів у контексті безпеки [2].

Часові ряди у кібербезпеці, такі як послідовності мережевих пакетів, журнали подій системи чи дані транзакцій, мають специфічну структуру, яка відрізняється від класичних табличних даних своєю послідовною природою та часовою залежністю, на відміну від простого переліку характеристик у стандартних табличних даних. Це вимагає особливого підходу до їх обробки та аналізу. Один із дієвих підходів полягає у використанні нейронних мереж для отримання векторних представлень (ембеддінгів) часових рядів з подальшим застосуванням методів виявлення аномалій для ідентифікації підозрілої активності.

Виникає питання, чи буде підхід, що використовує архітектури нейронних мереж, спеціально розроблених для табличних даних, для побудови ембеддінгів часових рядів, демонструвати значно відмінну поведінку порівняно із архітектурами, що не розроблені саме для табличних задач. У межах дослідження не ставилося за мету визначити загальну перевагу того чи іншого підходу, натомість акцент було зроблено на виявленні потенційних відмінностей у їхній поведінці під час взаємодії з різними алгоритмами виявлення аномалій.

У цьому контексті було проведено порівняльний аналіз двох архітектур нейронних мереж: TabNet, спеціально розробленої для обробки табличних даних [1], та згорткової нейронну мережу (CNN), традиційно використовуваної для обробки послідовних даних. Вибір саме цих архітектур був зумовлений їхньою принциповою відмінністю в підході до вхідних даних: TabNet - оптимізований для роботи зі структурованими табличними даними, тоді як CNN - це дуже проста і розповсюджена архітектура, що використовуються у широкому спектрі задач від машинного зору до аналізу часових рядів.

У рамках аналізу були згенеровані синтетичні часові ряди з різними пропорціями аномалій та різною силою аномальних відхилень. Для виявлення аномалій були використані два алгоритми: DBSCAN, алгоритм кластеризації, що базується на щільності розподілу даних [3], та Isolation Forest, ансамблевий метод для виявлення викидів, що використовує випадкове розбиття простору ознак [4].

Результати дослідження вказали на суттєві відмінності у поведінці алгоритмів виявлення аномалій при роботі з ембеддінгами, отриманими різними архітектурами. При використанні алгоритму DBSCAN, TabNet загалом демонструє кращі результати для виявлення аномалій. Натомість при застосуванні Isolation Forest, CNN стабільно перевершує TabNet при всіх пропорціях та силах аномалій. Під час симуляцій використовувались метрики F1 для оцінки якості виявлення аномалій, довжина часових рядів – 100 (відрізків), кількість записів для кожної симуляції – 2000, по 30 симуляцій на кожну з комбінацій «сили аномалії» та «частки аномалії у вибірці», частка аномалій коливалась від 1% до 15%.

Таблиця 1

Результати симуляцій

Detector	Outlier Proportion	Outlier Strength	CNN_F1 - TabNet_F1	t-test p-value
DBSCAN	0,01	0,50	-0,03880	0,00000

DBSCAN	0,01	1,00	-0,04880	0,00000
DBSCAN	0,01	2,00	-0,03370	0,00020
DBSCAN	0,01	3,00	-0,04820	0,00000
DBSCAN	0,05	0,50	-0,13740	0,00000
DBSCAN	0,05	1,00	-0,16440	0,00000
DBSCAN	0,05	2,00	-0,11240	0,00000
DBSCAN	0,10	0,50	-0,08890	0,00010
DBSCAN	0,10	1,00	-0,06350	0,00030
DBSCAN	0,10	2,00	-0,06100	0,00030
DBSCAN	0,10	3,00	0,06890	0,01710
DBSCAN	0,15	3,00	0,05760	0,00000
IFOREST	0,05	0,50	0,08230	0,00000
IFOREST	0,05	1,00	0,05640	0,00040
IFOREST	0,05	2,00	0,09880	0,00000
IFOREST	0,05	3,00	0,28440	0,00000
IFOREST	0,10	0,50	0,08000	0,00000
IFOREST	0,10	1,00	0,07960	0,00050
IFOREST	0,10	2,00	0,24140	0,00000
IFOREST	0,10	3,00	0,48010	0,00000
IFOREST	0,15	0,50	0,13220	0,00000
IFOREST	0,15	1,00	0,05810	0,00750
IFOREST	0,15	2,00	0,23080	0,00000
IFOREST	0,15	3,00	0,33160	0,00000

Отримані результати демонструють, що вибір архітектури нейронної мережі для генерації ембедінгів суттєво впливає на ефективність виявлення аномалій у часових рядах, причому різні алгоритми виявлення аномалій по-різному взаємодіють з векторними представленнями, отриманими різними архітектурами. Подальші дослідження в цьому напрямку можуть бути спрямовані на тестування цих підходів на реальних даних кібератак та розробку гібридних архітектур, що поєднують переваги обох підходів.

Список використаної літератури

1. Arik, S. O., & Pfister, T. (2021). *TabNet: Attentive interpretable tabular learning*. *Proceedings of the AAAI Conference on Artificial Intelligence*, 35(8):6679-87. <https://doi.org/10.48550/arXiv.1908.07442>.
2. Chen, Y., Zhao, C., Xu, Y., & Nie, C. (2025). *Year-over-Year Developments in Financial Fraud Detection via Deep Learning: A Systematic Literature Review*. <https://doi.org/10.48550/arXiv.2502.00201>
3. Ester, M., Kriegel, H.-P., Sander, J., & Xu, X. (1996). *A Density-Based Algorithm for Discovering Clusters in Large Spatial Databases with Noise*. *Proceedings of the 2nd International Conference on Knowledge Discovery and Data Mining (KDD-96)*.
4. Liu, F. T., Ting, K. M., & Zhou, Z.-H. (2008). *Isolation Forest*. *IEEE International Conference on Data Mining (ICDM)*. DOI: 10.1109/ICDM.2008.17.

МЕТОДИ МОНІТОРИНГУ ОБ'ЄКТІВ ЗАХИСТУ ІНФОРМАЦІЇ НА ОСНОВІ ТЕХНОЛОГІЇ DATA MINING

Литвиненко Б.В.

аспірант,

*Національний технічний університет
«Харківський політехнічний інститут»*

Завдяки стрімкому розвитку інформаційних технологій наразі відбувається всеохоплююча інтеграція технологій у різні сфери людського життя. Такі сфери як банківська справа, електронна комерція, уряд та медіа все більше ґрунтуються на інформаційних технологіях. В той же час кібератаки стають більш розповсюдженими і все більш складними, що створює все труднощі для точного виявлення вторгнень.

Забезпечення безпеки інформаційних та комунікаційних систем і мереж здебільшого може бути досягнуто за допомогою систем виявлення вторгнень (Intrusion Detection System, IDS) та систем запобігання вторгненням (Intrusion Prevention System). За практичним спрямуванням IDS/IPS систем можна виділити 2 напрямки [7]:

1. Виявлення зловживань (Misuse detection) — підхід орієнтований на виявлення відомих вторгнень на основі підходів синтаксичного порівняння відповідності структурних (патернів/сигнатур), інваріантних та кореляційних ознак процесу (системи) з існуючою базою відомих шаблонів. Недоліки підходу — неможливість автоматичного додавання нових шаблонів та неспроможність виявлення нових модифікацій чи атак нульового дня (0-day).

2. Виявлення аномалій (Anomaly detection) — підхід орієнтований на виявлення невідомих вторгнень на основі виявлення набору ознак, що не відповідає очікуваній поведінці користувача чи системи – шаблони характеристик, які не задовольняють визначеному поняттю нормальної поведінки фіксуються як аномалії.

Однак, згідно досліджень [10,12] традиційні методи виявлення вторгнень показують високу ефективність у виявленні відомих загроз (зловживань), відсоток виявлення варіюється від 51% до 93%, в той же час для невідомих загроз (аномалій) вони є неефективними — відсоток виявлення може не досягати навіть 20% залежно від типу атак. Таким чином, традиційні методи моніторингу не завжди здатні ефективно виявляти атаки, що може призводити до втрат інформації чи значних збитків. Вирішенням проблеми може бути інтеграція сучасних технологій, таких як Data Mining у процеси моніторингу об'єктів захисту інформації задля підвищення ефективності виявлення зловживань та аномалій.

Data Mining (інтелектуальний аналіз даних) — процес виявлення шаблонів і вилучення корисної інформації з великих наборів даних. Дану технологію можна використовувати для оцінки як структурованих, так і неструктурованих даних для виявлення нової інформації. Методи аналізу даних можна використовувати для спостереження та прогнозування поведінки користувачів, включаючи виявлення шахрайства у електронній комерції, банківській сфері тощо [14].

Виходячи з того, що технологія базується на обробці даних, основні проблеми теж пов'язані з даними. Найбільшою проблемою для технології є доступність даних, які можуть бути використані для аналізу — через особливості реалізації систем, необхідні дані можуть не зберігатися або зберігатися в недостатньому об'ємі; також дані потребують додаткового опрацювання перед використанням [8]. Також в процесі розробки IDS/IPS систем і при проведенні досліджень надзвичайно важливим завданням є вибір наборів даних (масив даних або Dataset) для проведення тестувань. Згідно досліджень [2,4,5,11] найбільш поширеними масивами даних для виявлення вторгнень є DARPA, KDD CUP, NSL-KDD, CAIDA та KYOTO. З перелічених вище масивів DARPA (приблизно 20% досліджень) та KDD CUP (приблизно 40% досліджень) є найбільш вживаними у наукових дослідженнях. Масив даних KDD CUP вважається загальноприйнятим еталоном для тестової перевірки ефективності систем IDS/IPS.

Проаналізувавши наукові праці [2,11,13,15] можна виділити 4 найефективніші алгоритми інтелектуального аналізу даних для виявлення вторгнень:

1. Random Forest — метод поєднує пакетування та випадковий вибір ознак для побудови групи дерев рішень із контрольованою варіацією. Продуктивність таблиці рішень і класифікаторів випадкового лісу використовуються для прогнозування точності класифікації. Виходячи з цього, Random Forest перевершує більшість інших технік [9].

2. J48 — створює двійковий кадр і класифікує будь-який новий елемент на основі дерева рішень, побудованого з атрибутів навчальних даних [9]. Алгоритм спочатку був розроблений як розширена версія алгоритму C4.5, але зі зменшеними вимогам до обчислювальних ресурсів [11].

3. SVM (Support Vector Machines) — Алгоритм можна описати у вигляді лінії розділення (плану прийняття рішень), яка ділить об'єкти на дві підмножини так, щоб у кожній підмножині всі елементи були подібними. Основна перевага алгоритму це можливість використовувати як для лінійної, так і для нелінійної класифікації [6].

4. Naïve Bayes — алгоритм, класифікатор якого заснований на висновках імовірнісних графічних моделей, які визначають імовірнісні залежності, що лежать в основі конкретної моделі, використовуючи структуру графа (вузли представляють випадкові величини, а дуги — припущення умовної залежності). Таким чином, він забезпечує компактне представлення спільних розподілів ймовірностей [9].

Порівнявши відсоток спрацювань, хибних спрацювань та нерозпізнаних атак, для різних реалізацій алгоритмів що були протестовані на масивах даних DARPA та KDD CUP, можна вивести наступні результати: алгоритм Random Forest має показник точності що варіюється від 99% до 99.92%, що робить його найточнішим з наведених. На другому місці за точністю спрацювань знаходиться J48 з точністю 96.8-99.91%, далі алгоритм SVM з показником точності 93-99.21% та найменша точність у алгоритму Naïve Bayes — лише 89.59-97%.

З наведених вище результатів можна стверджувати, методи моніторингу об'єктів захисту інформації на основі технології Data Mining з використанням алгоритму Random Forest може забезпечити успішне виявлення вторгнень з високою точністю.

Список використаної літератури

1. І. Ю. Субач, В. В. Фесьоха, *Модель виявлення аномалій в інформаційно – телекомунікаційних мережах органів військового управління на основі нечітких множин та нечіткого логічного виводу. Збірник наукових праць ВІТІ № 3, сс. 158-164, 2017.*
2. С. В. Довбешко, С. В. Толюпа, Я. В. Шестак, *Застосування методів інтелектуального аналізу даних для побудови систем виявлення атак. Сучасний захист інформації. № 1, сс. 6-15, 2019.*
3. A. Chiche, M. Meshesha, *An intelligent network intrusion detection system using data mining and knowledge based system. Journal of Theoretical and Applied Information Technology, Vol. 95, No. 17, pp. 4273-4282, 2017.*
4. B. Dhafian, I. Ahmad, A. AL-Ghamid, *An Overview of the current Classification Techniques in Intrusion Detection, in International Conference Security and Management, pp. 82-88, 2015.*
5. C. Azad, V. K. Jha. *Data mining in intrusion detection: a comparative study of methods, types and data sets. International Journal of Information technology and Computer Science (IJITCS). Jul1; Vol. 5, No. 8, pp. 75-90, 2013.*
6. C. Nilă, V-V. Patriciu, I. Bica, *Machine Learning Datasets for Cyber Security Applications, CONFSEC - 3rd International Scientific Conference (CONFSEC), Vol. 1, No. 5, pp. 117-120, 2019.*
7. D. Y. Yeung, C. Chow, «Prazen-window Network Intrusion Detectors», *In: 16th International Conference on Pattern Recognition, pp:11–15, 2002.*
8. G. M. Weiss, *Data Mining in the Real World: Experiences, Challenges, and Recommendations. International Conference on Data Mining (DMIN), pp. 124–130, 2009.*
9. H. Nguyen, C. Deokjai, *Application of data mining to network intrusion detection: Classifier selection model. Proceedings of Asia-Pacific Network Operations and Management Symposium, pp. 399–408, 2008.*
10. J. W. Haines, D. J. Fried and J. Korba, *Analysis and results of the 1999 DARPA off-line intrusion detection evaluation, in Proc. Intl. Symposium on Recent Advances in Intrusion Detection, Springer-Verlag (2000), pp. 162-182.*
11. M. Ahsan, K. E. Nygard, R. Gomes, M. M. Chowdhury, N. Rifat, J. F. Connolly, *Cybersecurity threats and their mitigation approaches using machine learning—a review, Journal of Cybersecurity and Privacy, Vol. 2, No. 3, pp. 527–555, 2022.*
12. R. Lippmann, D. Fried and I. Graf, *Evaluating intrusion detection systems: the 1998 DARPA off-line intrusion detection evaluation, in Proc. 2000 DARPA Information Survivability Conference and Exposition (2000), pp. 12-26.*
13. S. O. Al-Mamory, F. S. Jassim, *Evaluation of Different Data Mining Algorithms with KDD CUP 99 Data Set», Journal of Babylon University/Pure and Applied Sciences, Vol. 21, No. 8, pp. 2663-2681, 2013.*
14. *What is data mining? [Електронний ресурс]: <https://www.ibm.com/think/topics/data-mining>*
15. Z. Dewa and L. A. Maglaras, «Data Mining and Intrusion Detection Systems,» *International Journal of Advanced Computer Science and Applications, vol. 7, no. 1, pp. 33-42, 2016.*

ВПРОВАДЖЕННЯ ШТУЧНОГО ІНТЕЛЕКТУ У СИСТЕМИ КІБЕРЗАХИСТУ: МОЖЛИВОСТІ ТА ВИКЛИКИ

*Литвиненко Л.О.,
кандидат технічних наук, доцент кафедри
комп'ютерних наук та програмної інженерії,
Бойко М.М.
викладач кафедри комп'ютерних наук та програмної інженерії,
ПВНЗ «Європейський університет»*

В умовах зростаючої складності та інтенсивності кіберзагроз, впровадження штучного інтелекту (ШІ) у системи кіберзахисту постає як стратегічно важливий напрям забезпечення інформаційної безпеки в державному, корпоративному та інфраструктурному секторах. Традиційні засоби кіберзахисту, що базуються на сигнатурному аналізі, дедалі частіше виявляються недостатньо ефективними проти новітніх атак, які мають складну архітектуру та здатні до адаптації. У цьому контексті інструменти ШІ, зокрема машинне навчання, глибинне навчання та нейронні мережі, демонструють високу результативність у виявленні аномалій, аналізі поведінки користувачів і систем, а також у прогнозуванні потенційних загроз. Впровадження ШІ дозволяє не лише автоматизувати виявлення загроз, а й забезпечити адаптивне реагування на нові типи атак.

ШІ дедалі активніше впроваджується у системи виявлення та реагування на інциденти безпеки. Зокрема, алгоритми машинного навчання дозволяють:

- прогнозувати потенційні атаки на основі аналізу історичних даних [3];
- виявляти аномалії у поведінці користувачів і пристроїв, створюючи профілі «нормальної» активності, які дають змогу ідентифікувати підозрілу поведінку;
- реагувати на інциденти в реальному часі без участі людини, що особливо важливо при атаках нульового дня (zero-day);
- навчатися на нових загрозах, адаптуючись до змін у поведінці зловмисників та оновлюючи моделі захисту [2].

Наприклад, системи на основі нейронних мереж активно використовуються для розпізнавання шкідливих програм, фішингових листів, ботнет-активності тощо. Деякі провайдери безпеки також інтегрують рішення з обробки природної мови (NLP) для аналізу технічної документації атак або розслідування внутрішніх загроз. Однією з ключових переваг використання ШІ в системах кібербезпеки є здатність до обробки великих обсягів неоднорідних даних у режимі реального часу. Це дозволяє формувати динамічні профілі загроз і автоматизовано виявляти відхилення від нормальної поведінки в інформаційній системі. Так, алгоритми на основі *unsupervised learning* здатні ідентифікувати нові типи атак, які ще не були формалізовані у вигляді сигнатур. Крім того, використання *explainable AI* (XAI) забезпечує прозорість і підзвітність автоматизованих рішень, що є критично важливим для побудови довіри між системами захисту та їхніми користувачами.

Розглянемо ключові виклики та обмеження використання штучного інтелекту. Попри значний потенціал, впровадження ШІ в системи кіберзахисту не є позбавленим ризиків, зокрема:

- високий рівень помилкових спрацьовувань (false positives) може призводити до ігнорування справжніх загроз або перевантаження операторів безпеки;
- adversarial attacks — тобто атаки, що навмисно вводять ШІ в оману шляхом подачі модифікованих даних;
- відсутність прозорості у процесі прийняття рішень ШІ-системами (black-box models), що ускладнює аудит, перевірку та юридичну відповідальність;
- необхідність великих і якісних масивів навчальних даних, які не завжди доступні в достатньому обсязі, особливо для нових або рідкісних типів атак [1].

Крім того, важливим аспектом залишається етичність застосування інтелектуальних технологій: автоматизовані системи мають дотримуватися прав людини, забезпечувати недопущення дискримінації та відповідати нормам законодавства про захист персональних даних.

Інший важливий аспект — це кадрове забезпечення: ефективне впровадження ШІ у системи кіберзахисту потребує фахівців, які одночасно володіють знаннями в галузях безпеки, аналітики даних і розробки алгоритмів машинного навчання. З цієї точки зору важливу роль відіграє міждисциплінарна підготовка кадрів, а також співпраця між академічними установами, ІТ-компаніями та урядовими структурами.

Отже, штучний інтелект відкриває нові горизонти для підвищення ефективності систем кіберзахисту, проте його впровадження має супроводжуватися комплексним аналізом ризиків та розробкою нормативно-правових механізмів регулювання. Найефективнішим підходом є поєднання інтелектуальних алгоритмів з людським контролем, де ШІ виконує рутинну аналітику та виявлення загроз, а остаточні рішення ухвалюються фахівцями з безпеки. У майбутньому вирішальним фактором стане не лише технічна досконалість алгоритмів, а й здатність створити довірливе, відповідальне середовище їх застосування. Баланс між інноваційністю, безпекою та етичними стандартами стане визначальним у формуванні ефективної стратегії кіберзахисту на основі ШІ.

Список використаної літератури

1. Коломоєць, М.І. (2023). Штучний інтелект у системах кіберзахисту. *Інформаційні технології та безпека*, №2.
2. Russell, S., & Norvig, P. (2021). *Artificial Intelligence: A Modern Approach*. Pearson.
3. Buczak, A. L., & Guven, E. (2016). *A Survey of Data Mining and Machine Learning Methods for Cyber Security Intrusion Detection*. *IEEE Communications Surveys & Tutorials*, 18(2), 1153–1176.

СУЧАСНІ ІНСТРУМЕНТИ АВТОМАТИЗОВАНОГО ТЕСТУВАННЯ БЕЗПЕКИ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ

Максимович М.В.

*аспірант кафедри безпеки інформаційних технологій
Національний університет «Львівська політехніка»*

На тлі стрімкої цифровізації та зростання кількості кібератак, розвиток ефективних механізмів захисту інформаційних систем та програмного забезпечення (ПЗ) є дуже важливим на сьогоднішній день. Програмне забезпечення повинно не тільки виконувати якусь функцію бізнесу, але й повинно бути безпечним, вільним від вразливостей та стійким до атак, тому можна сказати що забезпечення безпеки є насамперед відповідальністю розробників [1]. Одним з ефективних засобів для мінімізації ризику виникнення вразливостей є впровадження автоматизованого тестування безпеки в життєвий цикл розробки, яке дозволяє своєчасно виявляти та усувати вразливості на ранніх етапах розробки [2].

Автоматизоване тестування безпеки представлене кількома типами аналізу: статичний (SAST), динамічний (DAST), інтерактивний (IAST) та композитний аналіз (SCA). SAST забезпечує аналіз вихідного коду без виконання програми, дозволяючи швидко ідентифікувати вразливості на ранніх етапах. DAST потребує запуску програми, виявляючи вразливості через симуляцію реальних атак на систему. IAST комбінує переваги SAST і DAST, проводячи тестування в реальному часі. SCA аналізує бібліотеки та залежності на наявність відомих вразливостей [3].

Сьогодні на ринку широко представлені як комерційні інструменти, так і рішення на основі відкритого коду. Комерційні продукти, такі як Jit Security, Veracode, Fortify, забезпечують високу функціональність і зручність інтеграції, підтримують різні типи тестування. Наприклад, Jit Security пропонує просту інтеграцію з популярними платформами розробки (GitHub, AWS, GCP), використовуючи SAST, DAST та SCA [3]. Veracode та Fortify також надають комплексні рішення з інтегрованим інтерфейсом для тестування безпеки.

Однак використання комерційних рішень має свої недоліки. Основним є висока вартість використання, що виходить за рамки бюджетів стартапів та малих компаній. Також, комерційні інструменти часто обмежують користувачів вибором тарифних планів, на основі яких пропонують певний функціонал, що знижує гнучкість налаштувань і кастомізації процесу тестування.

Інструменти відкритого коду, такі як OWASP ZAP, Semgrep та OWASP Dependency-check доступні безкоштовно. Вони мають високу гнучкість і можливість налаштування під конкретні потреби проекту. OWASP ZAP ефективно використовується для DAST-тестування, виявляючи SQL-ін'єкції та XSS-атаки. Semgrep підтримує статичний аналіз різних мов програмування, дозволяючи знаходити вразливості ще до розгортання програми. Dependency-check забезпечує композитний аналіз, перевіряючи сторонні бібліотеки на наявність відомих вразливостей. Недоліком використання інструментів на основі відкритого коду є складність впровадження та підтримки, тому можна сказати що дослідження в

області покращення методів впровадження та підтримки інструментів на основі відкритого коду є актуальними.

Важливим напрямом розвитку автоматизованого тестування є вдосконалення його методів та інтеграція нових технологій, зокрема штучного інтелекту (ШІ). Використання ШІ в автоматизованому тестуванні дозволяє підвищити ефективність аналізу, проактивно виявляти потенційні вразливості на ранніх етапах розробки програмного забезпечення, та надавати оцінку використовуваних рішень в контексті безпеки [5]. Серед методів застосування ШІ можна виділити використання машинного навчання для класифікації та пріоритизації вразливостей, прогнозування нових типів атак на основі історичних даних, автоматичне генерування тестових сценаріїв, а також розпізнавання аномалій у поведінці додатків у реальному часі.

Розглянемо приклад впровадження методів штучного інтелекту. Платформа Jit Security нещодавно представила інноваційних агентів на основі штучного інтелекту для автоматизації завдань в контексті безпеки додатків [6]. Як зазначається в статті, агенти можуть брати на себе рутинну роботу, суттєво прискорюючи аналіз та реагування на вразливості. Серед основних агентів штучного інтелекту згадуються такі як AppSec Agent, Compliance Agent та Security Ops Agent. Розглянемо коротко кожен з них:

- AppSec Agent будує і постійно актуалізує модель ризиків для додатку та надає розробникам контекстні підказки з безпеки під час так званого процесу перегляду коду (code review).

- Compliance Agent автоматично зіставляє стан середовища з вимогами стандартів і генерує звіти відповідності.

- Security Ops Agent створює завдання на виправлення, розставляє пріоритети та супроводжує процес усунення вразливостей. Таким чином, платформа з елементами ШІ може працювати як «член команди», що постійно стежить за безпекою: від сканування коду до контролю конфігурацій хмарної інфраструктури – і все це з урахуванням бізнес-контексту (критичності сервісів, потенційного впливу на користувачів тощо).

Інтеграція ШІ дозволяє фільтрувати тисячі знайдених проблем та виділяти серед них найбільш ризиковані для компанії, зменшуючи навантаження на фахівців. В результаті ефективність роботи безпекових команд зростає, адже вони можуть зосередитися на дійсно пріоритетних загрозах, які на даний момент не можна делегувати рішенням на основі штучного інтелекту [7].

Висновки. Зростання кількості кібератак і ускладнення програмних систем робить автоматизоване тестування безпеки невід’ємним елементом в життєвому циклі розробки програмного забезпечення.

Розглянуті сучасні інструменти (SAST, DAST, IAST, SCA) забезпечують багаторівневий захист, доповнюючи один одного. Практика показує, що для досягнення належного рівня безпеки варто використовувати комбінацію підходів: статичний аналіз для перевірки коду, динамічний – для сканування працюючих застосунків, інтерактивний – для глибокої перевірки під час тестування, а SCA – для контролю сторонніх компонентів. Комерційні платформи значно спрощують таку інтеграцію, але потребують значних фінансових ресурсів, тоді як безкоштовні

рішення доступніші й більш гнучкі, проте вимагають більше зусиль на впровадження і підтримку.

Водночас, кіберпростір постійно еволюціонує - з'являються нові типи загроз (наприклад, потенційні атаки з використанням методів штучного інтелекту), нові платформи і фреймворки, які також потребують перевірки. Це ставить перед науковцями та інженерами завдання постійного вдосконалення інструментів автоматизованого тестування безпеки. Перспективними напрямками подальших досліджень є розробка методів, що спрощують інтеграцію різного роду сканерів у єдину екосистему, зменшують кількість хибних спрацювань та автоматично адаптуються до змін у додатку, а також підходів що базуються на використанні методів штучного інтелекту.

Безперервне дослідження і вдосконалення в цій галузі дозволить забезпечувати надійний захист програмного забезпечення при мінімально необхідних затратах ресурсів.

Список використаної літератури

1. Pathirathna, P. P. W., Ayesha, V. A. I., Imihira, W. A. T., Wasala, W. M. J. C., Kodagoda, N., & Edirisinghe, E. A. T. D. (2017). *Security testing as a service with docker containerization*. 2017-December. Scopus. <https://doi.org/10.1109/SKIMA.2017.8294109>
2. Putra, A. M., & Kabetta, H. (2022). *Implementation of DevSecOps by Integrating Static and Dynamic Security Testing in CI/CD Pipelines*. ICOSNIKOM 2022 - 2022 IEEE International Conference of Computer Science and Information Technology: Boundary Free: Preparing Indonesia for Metaverse Society. Scopus. <https://doi.org/10.1109/ICOSNIKOM56551.2022.10034883>
3. Tudela, F. M., Higuera, J.-R. B., Higuera, J. B., Montalvo, J.-A. S., & Argyros, M. I. (2020). *On combining static, dynamic and interactive analysis security testing tools to improve owasp top ten security vulnerability detection in web applications*. Applied Sciences (Switzerland), 10(24), 1–26. Scopus. <https://doi.org/10.3390/app10249119>
4. Top 11 DevOps Security Tools. (2024, August 8). Jit. <https://www.jit.io/resources/appsec-tools/top-11-devops-security-tools>
5. *AI-Driven Security Testing: The Ultimate Guide 2025* - Thinksys Inc. (n.d.). Retrieved April 19, 2025, from <https://thinksys.com/qa-testing/ai-driven-security-testing/>
6. Williams, C. (2025, April 8). *Jit Launches an AI Workforce That's Helping Security Teams Keep Up*. Tech Times. URL: <https://www.techtimes.com/articles/309856/20250408/jit-launches-ai-workforce-thats-helping-security-teams-keep.htm>
7. *Will AI automate penetration testing?* (2024, May 10). Claranet Cyber Security. URL: <https://www4.claranet.com/us/blog/2024-05-10-will-ai-automate-penetration-testing>

АНАЛІЗ СТАНУ КІБЕРБЕЗПЕКИ ФІНАНСОВИХ УСТАНОВ ТА МЕТОДИ ПРОТИДІЇ КІБЕРАТАКАМ

Мойсеєнко Є.В.,

аспірант,

Яровий Р.О.,

к.т.н., доцент, доцент кафедри кібербезпеки

та захисту інформації

ПВНЗ “Європейський університет”

Авторами розглянуто сучасний стан кібербезпеки фінансового сектору, зокрема, банківських установ, проаналізовано основні вектори атак і вразливості, які найбільшою мірою експлуатують зловмисники. Представлено огляд ключових технологій захисту — управління вразливостями, веб-застосункові міжмережеві екрани (WAF), поведінковий аналіз у рішеннях EDR/NGAV, а також роль Центру операцій з безпеки (SOC), сегментації мережі і захисту від DDoS-атак. На основі сучасних статистичних даних та стандартів ISO/IEC 27001:2022 сформульовано рекомендації для побудови системного підходу до кіберзахисту. Наведемо далі основні з них.

1. Критичність забезпечення кібербезпеки в банківському секторі. Банківські установи залишаються пріоритетною цілью кібератак через високу цінність збереженої інформації і потенційні фінансові втрати у разі успіху зловмисників. Щорічно кількість виявлених IT-вразливостей зростає: у 2023 році було опубліковано 26 447 уразливостей, що на понад 1 500 більше, ніж 2022 року, що підкреслює динаміку ескалації загроз у глобальній IT-екосистемі (blog.qualys.com). Більше того, за даними OWASP, 94% протестованих додатків містили принаймні одну форму ін'єкції, а середній рівень інцидентності склав 3,37%, що демонструє високу поширеність базових векторів атак, таких як SQL-ін'єкції та XSS (owasp.org).

2. Сучасні вектори атак на фінансові вебсервіси. Аналіз щорічного звіту OWASP Top 10 показує, що ін'єкційні атаки (SQL Injection, XSS) займають топові позиції серед вразливостей веб-додатків, хоча в 2021 році вони опустилися на третє місце за частотою уражень, але залишаються критичними з огляду на кількість екземплярів (274 000) та середню інцидентність 3% (owasp.org). Додатково зловмисники використовують Illegal Resource Access, Bad Bots і Remote File Inclusion (RFI) для обходу традиційних захисних механізмів та стійкого проникнення в інфраструктуру (owasp.org).

3. Управління вразливостями (Vulnerability Management). Впровадження систем автоматизованого сканування, таких як Qualys Vulnerability Management, дозволяє відстежувати динаміку виявлення й усунення слабких місць у режимі реального часу. Звіт Qualys TruRisk Research 2023 показує, що сканери виявили понад 25 млн вразливостей, з яких 33% відносяться до категорії Misconfiguration (OWASP A05) (qualys.com). Центральною практикою є регулярне тестування за методиками NIST SP 800-115, що гарантує відповідність процесу рекомендаціям урядових стандартів (sprinto.com).

4. Веб-застосункові міжмережеві екрани (WAF). WAF-рішення фільтрують HTTP-запити в реальному часі, блокуючи атаки OWASP Top 10. Хмарна платформа Imperva Incapsula забезпечує автоматичне оновлення правил захисту та масштабується відповідно до навантаження, відбиваючи SQL-ін'єкції, XSS та інші загрози (imperva.com). On-premise рішення SecureSphere властиві розгорнуті механізми аналізу трафіку з глибокою перевіркою контексту сеансу, що дозволяє застосовувати кастомні політики безпеки під кожний бізнес-сервіс (imperva.com).

5. Поведінковий аналіз у EDR/NGAV системах. Системи нового покоління (EDR/NGAV) здатні виявляти аномальні патерни діяльності, що дозволяє ідентифікувати Zero-Day експлойти та складні малварні кампанії до їх реальної активації. Згідно з Forrester Wave: Endpoint Security Q4 2023, лідери ринку демонструють високі показники з оперативного реагування та кореляції подій на кінцевих точках (forrester.com).

6. Центр операцій з безпеки (SOC). SOC централізує збір, агрегацію та аналіз подій від різних джерел — WAF, EDR, SIEM. Стандарт NIST SP 800-137 визначає безперервний моніторинг, що дозволяє своєчасно виявляти і реагувати на загрози (nvlpubs.nist.gov). Грамотна архітектура SOC підвищує швидкість реагування та знижує середній час до виявлення (MTTD) і середній час до реагування (MTTR) (sprinto.com).

7. Сегментація мережі та контроль доступу. Логічна сегментація мінімізує межі ураження під час інциденту. Архітектурні практики CISA рекомендують створювати окремі зони для OT і IT мереж, обмежувати транзитні маршрути та застосовувати мікросегментацію для критичних сервісів (cisa.gov). До найкращих підходів належать: баланс між кількістю сегментів і рівнем адміністрування, автоматизовані політики доступу та аудит з'єднань UpGuard (upguard.com).

8. Захист від DDoS-атак. Cloudflare автоматично виявляє та пом'якшує DDoS-атаки за допомогою динамічних правил, що запобігають як мережевим, так і HTTP-флуд-атакам. У 2024 Q2 Cloudflare ліквідувала 10,2 трлн HTTP-запитів та 57 PB мережевого трафіку з атак (developers.cloudflare.com, blog.cloudflare.com). Рішення Arbor DDoS Protection забезпечують аналітику трафіку та інтегровані механізми реагування, що мінімізують простой сервісів і втрати SLA (netscout.com).

9. Системний підхід та безперервне вдосконалення. Стандарт ISO/IEC 27001:2022 визначає вимоги до впровадження ISMS, включаючи оцінку ризиків, керування активами та контроль доступу (iso.org, protiviti.com). Інтеграція ризик-орієнтованих процесів, регулярний перегляд контролів та адаптація до нових загроз забезпечують довгострокову кіберстійкість фінансових установ.

Список використаної літератури

1. R. Anderson, *Security Engineering: A Guide to Building Dependable Distributed Systems*, 3rd ed., Wiley, 2020.
2. B. Schneier, *Data and Goliath: The Hidden Battles to Collect Your Data and Control Your World*, W. W. Norton & Company, 2015.
3. OWASP Foundation, *OWASP Top 10 – 2021: The Ten Most Critical Web Application Security Risks*, 2021.
4. M. Scarfone, K. J. Houser, *Guide to Vulnerability Assessment (NIST SP 800-115)*, NIST, 2020.

5. Qualys, 2023 Threat Landscape Year in Review, 2023.
6. A. Sharma, P. Kumar, "Evaluating On-Premise WAF Solutions for Enterprise Security," *ACM Computing Surveys*, vol. 55, no. 2, pp. 1–28, 2023.
7. Imperva, Imperva Cloud WAF Datasheet, 2024.
8. Forrester Research, "The Forrester Wave™: Endpoint Security, Q4 2023," 2023.
9. NIST, *Information Security Continuous Monitoring (ISCM) for Federal Information Systems and Organizations (SP 800-137A)*, NIST, 2020.
10. S. Kent, T. Sommestad, "Designing Effective SOC Architectures in Financial Institutions," *IEEE Security & Privacy*, vol. 20, no. 1, pp. 34–42, 2022.
11. P. Raghavan, L. Chen, "Network Segmentation Techniques for Enhanced Cybersecurity," *Journal of Network and Computer Applications*, vol. 150, pp. 102–118, 2023.
12. CISA, *Layering Network Security Through Segmentation Infographic*, 2024.
13. Cloudflare, *DDoS Protection Service Overview*, 2024.
14. Arbor Networks, *The Total Economic Impact™ of NETSCOUT Arbor DDoS Protection*, 2023.
15. ISO/IEC 27001:2022, *Information security management systems — Requirements*.

СТРАТЕГІЧНА СЛІПОТА У КОРПОРАТИВНОМУ УПРАВЛІННІ ЯК ФАКТОР ЗРОСТАННЯ КІБЕРВРАЗЛИВОСТІ ЕКОНОМІКИ УКРАЇНИ В УМОВАХ ЗБРОЙНОГО КОНФЛІКТУ

*Ніколаєвський О.Ю.,
кандидат технічних наук, доцент, доцент кафедри
кібербезпеки та захисту інформації,
Ващенко М.В.,
викладач кафедри кібербезпеки та захисту інформації
ПВНЗ «Європейський університет»*

Гібридна війна проти України суттєво змінила уявлення про роль кібербезпеки в забезпеченні економічної стійкості держави. У сучасних умовах стратегічна сліпота корпоративного управління, що проявляється у недооцінці цифрових ризиків, стає одним із головних чинників зростання вразливості підприємств та економіки загалом [1].

Попри зростання кількості кіберінцидентів та їхніх економічних наслідків, багато організацій продовжують сприймати кібербезпеку як допоміжну функцію, що залишається на рівні IT-підрозділів і не інтегрується у загальні стратегії розвитку [2]. За даними World Economic Forum, понад 43% керівників бізнесу у 2024 році оцінюють цифрові ризики як «погано зрозумілі» в контексті їх впливу на довгострокову стійкість організацій [3].

Ігнорування системного управління кіберризиками сприяє не лише зростанню ймовірності успішних атак, але й формуванню ситуацій, коли атаки через ланцюги постачання або партнерські взаємодії призводять до каскадного впливу на економічні сектори [4]. Ініціатива CISA Shields Up підкреслює, що більшість успішних атак проти критичної інфраструктури реалізуються через непрямі вектори, використовуючи слабкі місця в корпоративному управлінні ризиками [5].

Формування цифрової стійкості має розглядатися як обов'язковий елемент корпоративної стратегії, що передбачає оцінку кіберризиків на рівні ради директорів,

розробку політик безпеки, адаптацію моделей бізнес-процесів до сучасних викликів. У цьому контексті Microsoft у своєму Digital Defense Report акцентує на необхідності включення аналізу кібервразливостей до щорічних звітів компаній як показника відповідальної бізнес-практики [4].

Підвищення уваги до управління цифровими ризиками має також базуватися на загальновизнаних міжнародних рекомендаціях щодо безпеки, таких як настанови ОЕСД для забезпечення економічної та соціальної стабільності через ефективне управління цифровими ризиками [6]. Оцінка організаційної резильєнтності, що враховує здатність підприємства реагувати на цифрові загрози, повинна стати ключовим елементом як внутрішнього контролю, так і зовнішньої оцінки стійкості бізнесу [7].

Таким чином, подолання стратегічної сліпоти у корпоративному управлінні є необхідною передумовою для зниження системної вразливості економіки в умовах ескалації кіберзагроз. Формування цифрової стійкості має стати інтегрованою складовою стратегії розвитку підприємств і частиною національної політики безпеки.

Список використаної літератури

1. ENISA. ENISA Threat Landscape 2023. European Union Agency for Cybersecurity. URL: <https://www.enisa.europa.eu/publications/enisa-threat-landscape-2023>
2. Accenture. Cyber Threat Intelligence Report 2023. Accenture. URL: <https://www.accenture.com/us-en/insights/security/cyber-threat-intelligence>
3. World Economic Forum. Global Risks Report 2024. URL: <https://www.weforum.org/publications/global-risks-report-2024/>
4. Microsoft. Digital Defense Report 2023. URL: <https://www.microsoft.com/en-us/security/business/security-intelligence-report>
5. CISA. Shields Up Campaign, 2023. URL: <https://www.cisa.gov/shields-up>
6. OECD. Recommendation on Digital Security Risk Management, 2021. URL: <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0456>
7. BSI. Organizational Resilience Index 2023. British Standards Institution. URL: <https://www.bsigroup.com/en-GB/our-services/organizational-resilience/organizational-resilience-index/>

РОЗПОДІЛЕНІ ЦИФРОВІ ДВІЙНИКИ В ПУБЛІЧНИХ МЕРЕЖАХ: ВИКЛИКИ КОНФІДЕНЦІЙНОСТІ ТА ПІДХОДИ НА ОСНОВІ БЛОКЧЕЙНУ

*Овсянко Д.О.,
аспірант*

Національний університет «Львівська Політехніка»

Поняття цифрового двійника (ЦД) стосується цифрової копії фізичної системи, пристрою чи об'єкта в режимі реального часу. У міру того, як галузі рухаються до взаємопов'язаних екосистем, з'явилося поняття розподілених цифрових двійників (DDT) - колекцій цифрових копій, що працюють у різних місцях і системах, часто взаємодіючи через загальнодоступні мережі [3].

Оскільки дані системи можуть бути сильно розподіленими в фізичному просторі, це унеможливорює побудову ізольованих, захищених каналів зв'язку.

Єдиним виходом є використання публічних мереж типу Інтернет. Конфіденційність DDT стає критично важливою проблемою в публічних мережах, де процеси передачі даних і управління ними більш схильні до перехоплення, маніпуляцій і несанкціонованого доступу.

Робота в публічних інфраструктурах, таких як Інтернет, створює численні вразливості. DDT часто обмінюються конфіденційною інформацією, в тому числі телеметрією, оперативними даними і сигналами управління, які можуть бути перехоплені або використані зловмисниками, якщо вони не зашифровані належним чином. Ефективна перевірка ідентичності має важливе значення для забезпечення доступу до DDT або взаємодії з нею лише уповноважених суб'єктів. У розподілених системах звичайні централізовані системи управління ідентифікацією борються з затримками, єдиними точками відмови та обмеженою масштабованістю довіри. Якщо навіть один DDT у розподіленому середовищі скомпрометований, зловмисники можуть змінювати набори даних, надсилати зловмисні команди або поширювати неправдиву інформацію. Тому забезпечення цілісності даних і контроль доступу є основою для забезпечення конфіденційності в екосистемах DDT.

Блокчейн забезпечує розподілені, незмінні реєстри ідентичності, де кожен цифровий двійник може бути криптографічно автентифікований. Завдяки децентралізованим ідентифікаторам (DID) та обліковим даним, які можна перевірити, можна встановити довіру між агентами, не покладаючись на центральну владу [1]. Смарт-контракти - самодостатні скрипти, розгорнуті на блокчейні - пропонують потужний механізм для автоматизованого та захищеного від несанкціонованого втручання застосування політик конфіденційності. Наприклад, вони можуть:

- перевіряти права доступу до обміну даними;
- визначати динамічні правила конфіденційності в залежності від контексту або ролі;
- записувати взаємодію у зручний для аудиту спосіб, що піддається перевірці.

Вбудовуючи логіку для збереження конфіденційності безпосередньо в мережевий рівень, смарт-контракти зменшують залежність від людського нагляду і централізованого управління. Хоча блокчейн забезпечує незмінність і можливість аудиту, для збереження конфіденційності даних на нього можуть бути накладені розширені криптографічні протоколи, такі як докази з нульовим знанням і безпечні багатосторонні обчислення. У таких умовах DDT можуть доводити дотримання правил, не розкриваючи при цьому приватних даних, що лежать в їх основі.

Розглянемо наявні проблеми застосування технології блокчейн в децентралізованих цифрових двійниках. Масштабованість та затримка в системах блокчейн. Загальнодоступні блокчейни часто страждають від обмеженої пропускну здатності та високої затримки. Застосування блокчейну в середовищах DDT, де продуктивність у реальному часі є критично важливою, вимагає рішень другого рівня (Layer-2) або спеціальних приватних блокчейнів.[2]

Наступною проблемою є інтероперабельність між протоколами DDT і блокчейн. Сучасні промислові платформи DDT не призначені для взаємодії з інфраструктурою блокчейну. Подолання цього розриву потребує нових архітектур

проміжного програмного забезпечення, стандартизованих API та безпечних оракулів, які можуть передавати реальні дані в логіку ланцюжка.

Останніми, проте не менш важливими є правові та етичні проблеми. Автоматизоване забезпечення конфіденційності за допомогою смарт-контрактів викликає питання щодо дотримання вимог, підзвітності та механізмів відшкодування. Крім того, незмінні журнали даних повинні відповідати правилам захисту даних, таким як GDPR, які передбачають право на забуття.

Швидкий розвиток розподілених цифрових двійників у публічних мережевих середовищах створює серйозні виклики для конфіденційності та довіри. Традиційні моделі безпеки не в змозі впоратися з динамічною, децентралізованою природою цих систем. У цій статті підкреслюється потенціал технологій блокчейн і смарт-контрактів для заповнення цієї прогалини шляхом створення автономних засобів контролю конфіденційності, які можна перевірити і які є стійкими до несанкціонованого втручання. Однак цей підхід має свої технічні та регуляторні складнощі. Необхідні подальші дослідження для розробки масштабованих блокчейн-архітектур, покращення інтеоперабельності DDT-ланцюгів та приведення управління смарт-контрактами у відповідність до глобальних стандартів конфіденційності.

Список використаної літератури

1. Casino, F., Dasaklis, T. K., & Patsakis, C. (2019). *A systematic literature review of blockchain-based applications: Current status, classification and open issues. Telematics and Informatics*, 36, 55–81.
2. Christidis, K., & Devetsikiotis, M. (2016). *Blockchains and smart contracts for the Internet of Things. IEEE Access*, 4, 2292–2303.
3. Tao, F., Zhang, H., Liu, A., & Nee, A. Y. C. (2019). *Digital Twin in Industry: State-of-the-Art. IEEE Transactions on Industrial Informatics*, 15(4), 2405–2415.

ВПЛИВ ШТУЧНОГО ІНТЕЛЕКТУ НА БЕЗПЕКУ І ЯКІСТЬ ПРОГРАМНИХ СИСТЕМ: ПОГЛЯД З ПОЗИЦІЇ QA-СПЕЦІАЛІСТА

*Одольський М.Ю.,
магістрант факультету інформаційних систем та технологій,
ПВНЗ «Європейський університет»*

У сучасному цифровому середовищі питання якості програмного забезпечення посідає одне з провідних місць у забезпеченні стабільної роботи різних бізнес-процесів, державних структур та персональних даних. Із поширенням систем на базі штучного інтелекту зростає також складність забезпечення їх коректної та передбачуваної поведінки з обов'язковим урахуванням такого важливого фактору, як безпека. Це зумовлює перегляд ролі фахівця з контролю якості програмного забезпечення.

Штучний інтелект активно застосовується для оптимізації процесів контролю якості. Зокрема, автоматизоване створення тест-кейсів, аналіз логів та пошук

помилки стали поширеними практиками. Водночас виникає проблема непередбачуваності поведінки моделей, що навчаються саме на змінних даних. У ситуаціях, коли системи, що працюють з використанням штучного інтелекту, приймають рішення, спираючись на змінні дані, передбачити їхню поведінку стає доволі складно. А ще складніше виходить пояснити результати звичайними методами. Тому, в такій ситуації, фахівець з контролю якості повинен не лише володіти технічними навичками, а й також має розуміти, як саме модель приймає рішення, і чому система поводить певним чином у конкретних умовах [3].

Саме в подібних умовах, коли машинне навчання відбувається на змінних даних, питання безпеки набуває особливого значення. Інструменти, які створені для виявлення вразливостей, можуть бути використані не лише для захисту. Бувають такі ситуації, що інструменти, які спочатку створювалися для підвищення безпеки, потрапляють до рук зловмисників, які намагаються використати їх у своїх інтересах, знаходячи слабкі місця в системі [1]. У таких випадках ризики значно зростають, адже передбачити всі можливі варіанти розвитку подій вже не виходить, а стандартні методи перевірки та забезпечення якості теж далеко не завжди дають потрібний результат, і це ускладнює процес контролю ще більше.

З часом усе це змінює уявлення про те, ким є спеціаліст із QA. Його завдання вже не зводиться лише до пошуку помилок та багів, а тепер він намагається зрозуміти, як поводить система та виявити потенційно небезпечні місця ще до того, як вони стануть проблемою, і допомогти уникнути потенційних загроз. В такій ситуації, дуже важливою стає щоденна командна робота, детальна комунікація QA-спеціаліста зі спеціалістами з кібербезпеки, спільний аналіз і бажання підтримати стабільність системи.

Отже, в якості висновку, буде логічним зазначити, що штучний інтелект може бути як інструментом підвищення якості, зручності та автоматизації так і фактором потенційного ризику саме для безпеки програмного забезпечення [2]. Тому для ефективного використання ШІ спеціалістам з контролю якості потрібен новий підхід з акцентом на прозорість, контроль і кіберетичну відповідальність. QA-інженер майбутнього повинен бути не просто тестувальником чи «вартовим» дотримання вимог, а й аналітиком ризиків, який працює на перетині якості, безпеки та технологій ШІ, враховуючи та передбачаючи поведінку різних моделей, які навчаються на змінних даних. Для цього, з умов безпеки, є сенс ретельно тестувати та вивчати поведінку ШІ саме на локальних тестових задачах та симуляціях перед тим, як вводити цей інструмент в роботу у реальні програмні системи.

Список використаної літератури

1. *Europol*. Criminal Use of ChatGPT: Cautionary tale about large language models. 2023. URL: <https://www.europol.europa.eu/media-press/newsroom/news/criminal-use-of-chatgpt-cautionary-tale-about-large-language-models>
2. *LinuxSecurity*. How AI and Machine Learning Are Transforming Cybersecurity and Quality Assurance. 2025. URL: <https://linuxsecurity.com/features/how-ai-and-machine-learning-are-transforming-cybersecurity-quality-assurance>
3. *OWASP Foundation*. OWASP Top 10 for Machine Learning Security. 2023. URL: <https://owasp.org/www-project-machine-learning-security-top-10/>

ОГЛЯД КІБЕРАТАКИ НА ЕНЕРГЕТИЧНУ ТА ТЕЛЕКОМУНІКАЦІЙНУ ІНФРАСТРУКТУРУ УКРАЇНИ: АНАЛІЗ ІНЦИДЕНТІВ З «УКРЕНЕРГО» ТА «КІЇВСТАР»

*Опольський М.В.,
старший викладач кафедри математичних дисциплін
та інноваційного проектування,
Позняк А.А.
викладач Фахового бізнес коледжу,
ПВНЗ «Європейський університет»*

У контексті повномасштабної агресії РФ проти України кіберпростір став ще одним фронтом війни. Зокрема, особливої уваги набули цілеспрямовані атаки на об'єкти критичної інфраструктури — енергетичні компанії та телекомунікаційних операторів. Два яскравих приклади таких атак — інциденти, пов'язані з НЕК «Укренерго» та мобільним оператором «Київстар». Ці кейси демонструють зміну підходів до кібервійни: від саботажу та психологічного тиску до комплексної деструкції стратегічно важливих цифрових систем.

Перша масштабна кібератака на українську енергетичну систему була зафіксована ще у грудні 2015 року, коли атака на енергокомпанію «Прикарпаттяобленерго» призвела до відключення електроенергії для понад 200 тисяч споживачів. Але саме у 2016 році атака на НЕК «Укренерго» стала поворотним моментом, адже зловмисники вперше використали спеціалізоване шкідливе ПЗ — Industroyer (CrashOverride), яке було здатне безпосередньо взаємодіяти з протоколами енергосистем (IEC 101/104, OPC тощо).

Хакери, імовірно з угруповання Sandworm, проникли в системи управління компанії через фішинг та вразливості у віддаленому доступі. Після збору розвідданих вони ініціювали виведення з ладу частини автоматизованих систем SCADA, що керують передачею електроенергії. Попри обмеженість впливу за обсягом, атака засвідчила наявність у зловмисників глибоких технічних знань щодо енергетичної інфраструктури.

Після початку повномасштабної війни у 2022 році «Укренерго» знову зазнало атак — як кібернетичних, так і фізичних (удари по підстанціях та лініях передачі). У 2022–2023 роках фіксувалися спроби використання модифікованих варіантів Industroyer та інших інструментів шкідливого впливу, проте завдяки покращеним заходам кіберзахисту, частину атак вдалося виявити й нейтралізувати.

Найбільш резонансною кіберподією 2023 року стала атака на мобільного оператора «Київстар» у грудні. Унаслідок зламу було повністю знищено велику частину ІТ-інфраструктури компанії, що призвело до масштабного відключення мобільного зв'язку, мобільного інтернету та сервісів, що базувалися на цих каналах (банкінг, державні послуги, зв'язок ЗСУ). Порушення тривало кілька днів, а відновлення окремих сервісів — до кількох тижнів.

За даними СБУ, атаку було здійснено угрупованням Sandworm, пов'язаним із ГРУ РФ. Зловмисники проникли у внутрішню мережу «Київстару» за кілька місяців до атаки, приховано переміщувалися в інфраструктурі та поступово готували

середовище до повного знищення. Кінцевий етап — використання скриптів і деструктивного ПЗ для стирання вмісту серверів (атака типу wiper), включно з резервними копіями.

Особливість цієї атаки — її деструктивний характер. Якщо попередні кібератаки мали переважно розвідувальний або тимчасовий вплив, то тут була мета — знищення, а не зупинка. Імовірно, одночасно з технічним руйнуванням планувався психологічний вплив на користувачів і державу, зокрема, демонстрація здатності РФ паралізувати цифрову інфраструктуру України.

Кібератаки на об'єкти критичної інфраструктури, зокрема національну енергетичну компанію «Укренерго» та телекомунікаційного оператора «Київстар», засвідчують, що кіберзагрози для таких систем є не гіпотетичними, а цілком реальними та масштабними за своїми наслідками. Зазначені інциденти свідчать про трансформацію характеру кіберзагроз — від ізольованих точкових атак до комплексних, скоординованих операцій, спрямованих на повне або часткове виведення з ладу критичної інфраструктури. У цьому контексті Україна перебуває в ситуації, коли недостатньо лише реагувати на вже здійснені вторгнення; необхідним є перехід до проактивних моделей кібероборони. Зокрема, йдеться про впровадження глибокої сегментації мереж, створення хмарних резервних копій критичних систем, розбудову багаторівневого моніторингу кіберзагроз, а також активне міжнародне співробітництво у сфері кібербезпеки як один з інструментів стримування та протидії кібератакам.

Ці інциденти підкреслюють необхідність глибшої інтеграції кібербезпеки в національну систему оборони, а досвід України може стати взірцем для інших країн, які прагнуть посилити захист критичної інфраструктури в умовах зростання гібридних загроз.

Список використаної літератури

1. NIST. Zero Trust Architecture (Special Publication 800-207) / Rose S., Borchert O., Mitchell S., Connelly S. — National Institute of Standards and Technology, 2020. — 59 p.
2. Закон України «Про основні засади забезпечення кібербезпеки України» № 2163-VIII від 05.10.2017 // Відомості Верховної Ради України. — 2017. — №45.
3. Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data (General Data Protection Regulation, GDPR). — Official Journal of the European Union, L119, 04.05.2016.
4. Directive (EU) 2022/2555 of the European Parliament and of the Council of 14 December 2022 on measures for a high common level of cybersecurity across the Union (NIS2 Directive). — Official Journal of the European Union, L333, 27.12.2022.
5. ENISA. Zero Trust Maturity Model and Implementation Guidelines / European Union Agency for Cybersecurity. — Heraklion: ENISA, 2023. — 38 p.
6. ISO/IEC 27001:2022. Information security, cybersecurity and privacy protection — Information security management systems — Requirements / International Organization for Standardization. — Geneva: ISO, 2022. — 35 p.

МЕТОДИ ЗАХИСТУ В МЕРЕЖАХ WPA2 ТА WPA3: ВІД KRACK ДО DRAGONBLOOD

*Павлів Т. Ю.
студентка кафедри безпеки інформаційних технологій,
Коробейнікова Т. І.
к.т.н., доцент
доцент кафедри безпеки інформаційних технологій,
Національний університет «Львівська політехніка»*

Протоколи WPA (Wi-Fi Protected Access), включаючи вдосконалені версії WPA2 і WPA3, захищають бездротові мережі, але не є повністю безпечними. Дві відомі атаки – KRACK та Dragonblood – продемонстрували серйозні вразливості навіть у найновіших стандартах.

KRACK (Key Reinstallation Attack) націлена на WPA2, використовуючи слабкість чотиристороннього рукостискання (рис. 1). Зловмисник змушує пристрій жертви повторно використовувати ключ шифрування, що дозволяє йому перехоплювати трафік без знання пароля Wi-Fi (рис. 2). Це уможливорює крадіжку персональних даних, таких як облікові дані для входу, паролі, банківська інформація тощо.

WPA/WPA2 4 Ways Handshake

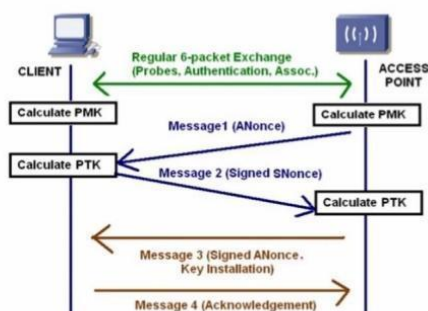


Рис. 1. Чотиристороннє рукостискання WPA2.

Належить джерелу [4].

Хоча WPA3 забезпечує більш надійне шифрування, атака Dragonblood виявила вразливості в його рукостисканні SAE. Зловмисники можуть здійснювати атаки через побічні канали або знизити рівень захисту до WPA2. Це створює можливість зламати паролі та отримати доступ до мережевого трафіку, особливо якщо не використовується додаткове шифрування на рівні додатків, наприклад, HTTPS.

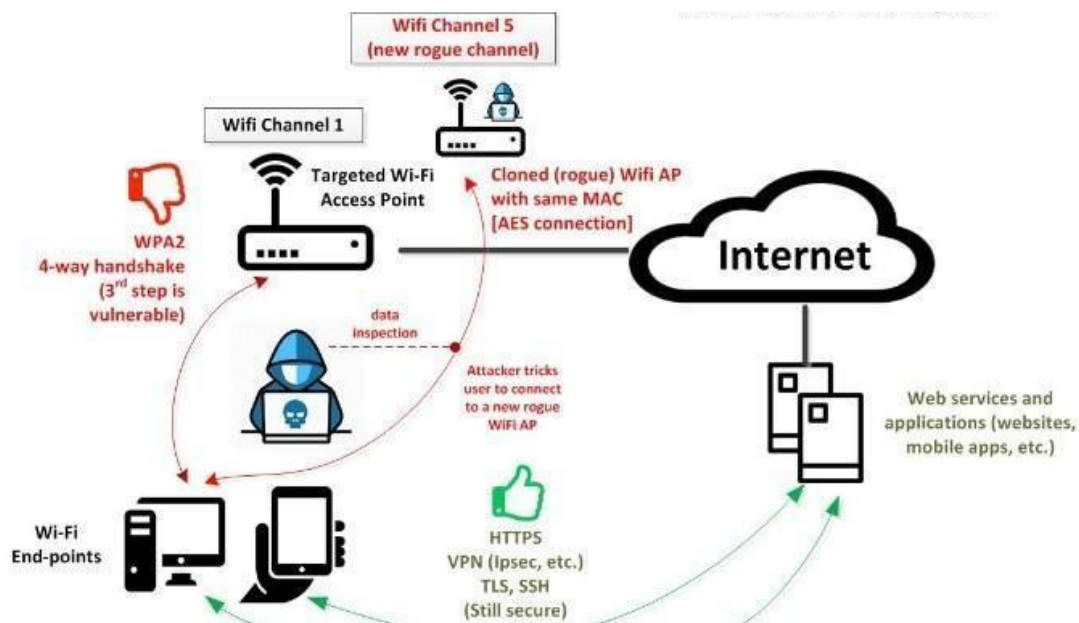


Рис. 2. KRACK атака.

Належить джерелу [4].

Щоб зменшити ризик, рекомендується регулярно оновлювати прошивку пристрою, використовувати WPA3 без режиму сумісності, створювати надійні паролі та застосовувати VPN – особливо в публічних мережах. Постійні оновлення та правильна конфігурація значно підвищують безпеку.

Сучасні Wi-Fi атаки діляться на дві основні категорії: пасивні та активні. Пасивні атаки не втручаються в трафік, а аналізують його. Наприклад, прослуховування пакетів дозволяє зловмисникам перехоплювати незашифровані дані за допомогою таких інструментів, як Wireshark.

Активні атаки безпосередньо впливають на мережу:

- DoS-атаки (відмова в обслуговуванні) перевантажують мережу запитами, спричиняючи збій в роботі.
- MITM (Man-in-the-Middle) дозволяє зловмиснику перехоплювати або змінювати дані, що передаються між користувачем і сервером.
- Evil Twin – підроблена точка доступу, яка імітує легальну мережу Wi-Fi. Користувачі несвідомо підключаються до неї і передають конфіденційну інформацію зловмиснику. Важливо відрізнити атаку Evil Twin від неавторизованої точки доступу: перша імітує справжню мережу, тоді як друга підключається до легітимної без дозволу.

Інший ризик – фізичний доступ: зловмисник може отримати контроль над маршрутизатором, змінити його конфігурацію або повністю вимкнути його. Часто зловмисники використовують атаку грубої сили, витік облікових даних і методи соціальної інженерії, наприклад, фішингові веб-сайти, які збирають паролі до Wi-Fi.

Наступною атакою є підміна (спуфінг) MAC-адреси, у цьому типі злому, Wi-Fi зловмисник змінює MAC-адресу (Media Access Control) свого пристрою на адресу авторизованого комп'ютера в мережі. В результаті зловмисник може увійти в мережу, не запитуючи дані для входу користувача.

Крім усіх цих атак, щоб перешкодити роботі мережі Wi-Fi, зломисники надсилають сигнал (наприклад, білий шум) на тій самій довжині хвилі, що й мережа. Ці атаки можуть призвести до того, що мережа працюватиме повільніше або може зникнути повністю.

Можемо підсумувати, що компанії, які працюють з конфіденційною інформацією, не повинні шукати неідентифіковані або несанкціоновані Wi-Fi, оскільки вони можуть бути налаштовані на атаку або захоплення їхніх Wi-Fi мереж. Це серйозна проблема, оскільки атака може дозволити хакерам отримати облікові дані для входу в систему і використовувати їх для доступу до мереж компанії або навіть повністю відключити Wi-Fi. Захист від мережевих вразливостей вимагає постійної оцінки, моніторингу та адаптації загроз. Безпека Wi-Fi залежить як від технічних заходів, так і від обізнаності користувачів.

Список використаної літератури

1. Гальчинський, Л. Ю., Корольова, В. Р. (2023). Оцінка захищеності бездротових мереж у міському середовищі з урахуванням використання протоколів безпеки. *Інфокомунікаційні та комп'ютерні технології*, 2(06), 9-21.
2. R. Banakh and A. Piskozub, «Attackers' Wi-Fi Devices Metadata Interception for their Location Identification,» 2018 IEEE 4th International Symposium on Wireless Systems within the International Conferences on Intelligent Data Acquisition and Advanced Computing Systems (IDAACS-SWS), Lviv, Ukraine, 2018, pp. 112-116.
3. Скидан, Д., Гальчинський, Л. (2024). Оцінка рівня захищеності протоколів безпеки для бездротових мереж. *Collection of scientific papers «SCIENTIA»*, (June 7, 2024; Antwerp, Belgium), 113-117.
4. An infographic about KRACK (aka WIFI WPA2 vulnerability). URL: https://www.linkedin.com/redir/redirect?url=https%3A%2F%2Fwww%2Elinkedin%2Ecom%2Fpulse%2F infographic-krack-aka-wifi-wpa2-vulnerability-jiang-ciso-cissp-1&urlhash=MK4k&trk=public_post_feed-article-content

МІЖ СТРАТЕГІЧНОЮ ІНДИФЕРЕНТНІСТЮ ТА ЦИФРОВОЮ СТІЙКІСТЮ: ПРОБЛЕМА УСВІДОМЛЕННЯ ВАЖЛИВОСТІ КІБЕРБЕЗПЕКИ КЕРІВНИЦТВОМ КОМПАНІЙ В УМОВАХ ВІЙНИ

*Пашорін В.І.,
к.т.н., професор,
завідувач кафедри кібербезпеки та захисту інформації
Ващенко М. В.
викладач кафедри кібербезпеки та захисту інформації
Саморай О. К.
викладач кафедри кібербезпеки та захисту інформації
ПВНЗ «Європейський університет»*

У сучасних умовах повномасштабної війни кібератаки стали системною загрозою, що впливає як на державний сектор, так і на комерційні структури. Попри зростання інтенсивності таких атак, спостерігається низький рівень усвідомлення

важливості кібербезпеки з боку керівників підприємств і організацій. Це формує вразливість цифрового середовища бізнесу, що, у свою чергу, негативно впливає на загальну кіберстійкість держави.

У 2021 році в Україні було зафіксовано 147 кіберінцидентів [1]. У 2022 році, після початку повномасштабного вторгнення, їхня кількість стрімко зросла до 5 970 [2]. У 2023 році CERT-UA опрацювала 2 543 інциденти [3], а в 2024 — вже 4 315 [4]. Така нерівномірна динаміка демонструє, що кіберактивність супротивника адаптується до фаз конфлікту, і підтверджує необхідність системного підходу до кіберзахисту незалежно від «піків» атак.

Попри загрозливу статистику, кібербезпека часто не входить до кола стратегічних пріоритетів бізнесу [5]. Вона сприймається як другорядна або суто технічна функція, яку можна делегувати без інтеграції в загальний процес управління ризиками.

Додатковим свідченням знецінення цієї теми є висловлювання міністра цифрової трансформації України Михайла Федорова, який у 2019 році заявив, що «роль кібербезпеки трохи перебільшена» [6].

Комунікаційна розбіжність між технічними спеціалістами й керівниками знижує ефективність рішень щодо інвестування у захист [7]. У багатьох компаніях технічні звіти не потрапляють до уваги управлінців або подаються у формі, яка не дозволяє зважено оцінити масштаби ризику.

Міжнародний досвід показує: компанії, які впровадили систему ISO/IEC 27001, знижують втрати від атак до 40–60% [8]. Це підтверджує, що кібербезпека — не витрата, а стратегічна інвестиція у стабільність, конкурентоспроможність і довіру клієнтів.

Для посилення участі керівництва у питаннях кібербезпеки необхідно впроваджувати сучасні підходи, які одночасно формують розуміння природи ризиків і залучають до процесів прийняття рішень. Насамперед це моделювання загроз, де для кожного сценарію оцінюються потенційні фінансові, операційні та репутаційні наслідки. Також ефективними є симуляційні навчання (tabletop), які відтворюють ситуації кризового реагування з урахуванням особливостей конкретного бізнесу. Регулярний аналіз реальних кейсів атак на конкурентів дозволяє управлінському складу усвідомити наслідки недооцінки безпеки. Участь CISO у стратегічних засіданнях і формування зрозумілих KPI (наприклад, середній час реагування, кількість навчань) сприяють культурі цифрової відповідальності.

В умовах війни формування культури цифрової стійкості на рівні керівництва — необхідна передумова для збереження бізнесу та підтримки національної безпеки. Системний підхід передбачає постійний моніторинг загроз, внутрішню звітність з безпеки та взаємодію з державними структурами. Жодна стратегія кіберзахисту не дасть бажаного ефекту без участі керівництва.

Список використаної літератури

1. CERT-UA. Звіт 2021 року. https://cert.gov.ua/files/pdf/SOC_Annual_Report_2022.pdf
2. CERT-UA. Звіт 2022 року. URL: <https://cip.gov.ua/ua/news/u-2022-roci-kilkist-zareyestrovanih-kiberincidentiv-viroslo-maizhe-vtrichi-zvit>

3. CERT-UA. Звіт 2023 року. URL: <https://cip.gov.ua/ua/news/uryadova-komanda-cert-ua-v-2023-roci-opracuyovala-2543-kiberincidenti>
4. CERT-UA. Оперативна інформація, 2024. URL: <https://cip.gov.ua>
5. Ponemon Institute. The Cost of Cybercrime. Accenture, 2023. URL: <https://iapp.org/resources/article/the-cost-of-cybercrime-annual-study-by-accenture/>
6. Федоров М. Роль кібербезпеки трохи перебільшена : інтерв'ю // LB.ua. – 29.11.2019. URL: https://lb.ua/tech/2019/11/29/443542_rol_kiberbezopasnosti_nemnogo.html
7. ENISA. Threat Landscape 2023. URL: <https://www.enisa.europa.eu/publications/enisa-threat-landscape-2023>
8. ISO. ISO/IEC 27001:2022 – Information security, cybersecurity and privacy protection. URL: <https://www.iso.org/standard/27001>
9. CISA. Shields Up Initiative. URL: <https://www.cisa.gov/shields-up>

БЕЗПЕКА ХМАРНИХ ОБЧИСЛЕНЬ: ЗАХИСТ ДАНИХ ТА ІНФРАСТРУКТУРИ ПРИ ВИКОРИСТАННІ IAAS, PAAS, SAAS РІШЕНЬ

*Переверзєв А.М.,
вчитель інформаційних технологій
Фаховий бізнес-коледж Європейського університету
Подвиженко А.В.,
вчитель інформаційних технологій
Фаховий бізнес-коледж Європейського університету
Левченко С.В.,
старший викладач кафедри комп'ютерних наук
та програмної інженерії
ПВНЗ «Європейський університет»*

Хмарні обчислення значно спростили спосіб зберігання даних, розробки додатків та ведення бізнесу. Такі моделі як: Інфраструктура як послуга (IaaS), Платформа як послуга (PaaS) та Програмне забезпечення як послуга (SaaS) пропонують гнучкість, масштабованість та економічну ефективність. Однак, разом із перевагами виникають і нові виклики, пов'язані з безпекою даних та інфраструктури.

Зберігання корпоративних даних та робочих навантажень у хмару означає передачу частини контролю над інфраструктурою та даними третій стороні, яка являється хмарним провайдером. Це створює унікальні ризики безпеки, які відрізняються від традиційних локальних середовищ. До основних проблем належать наступні.

1. Розмита відповідальність. Наприклад, при використанні моделі спільної відповідальності, не завжди чітко зрозуміло, де закінчується відповідальність провайдера і починається відповідальність клієнта за безпеку.
2. Ризики несанкціонованого доступу, витоку, втрати або пошкодження даних через помилки конфігурації, атаки злоумисників або збої у хмарного провайдера.
3. Вразливості інфраструктури (неправильна конфігурація віртуальних мереж, серверів, сховищ або платформ) можуть створити точки входу для атак.

4. Складнощі відповідності нормативним вимогам, дотримання галузевих стандартів та законодавчих норм (наприклад, GDPR, PCI DSS) у хмарному середовищі через розподілений тип інфраструктури.

5. Недостатня видимість та контроль. Наприклад, клієнтам може бракувати повного уявлення про те, що відбувається в інфраструктурі провайдера, що ускладнює моніторинг та реагування на інциденти.

Ці проблеми вимагають ретельного аналізу та впровадження комплексних заходів безпеки, адаптованих до конкретної моделі хмарних послуг (IaaS, PaaS, SaaS), що використовується.

Забезпечення безпеки в хмарі ґрунтується на моделі спільної відповідальності. Суть її полягає в тому, що хмарний провайдер відповідає за безпеку самої хмари (фізична інфраструктура, мережа, гіпервізор), тоді як клієнт відповідає за безпеку в хмарі (дані, додатки, конфігурації, управління доступом). Рівень відповідальності клієнта суттєво відрізняється залежно від моделі послуг.

1. IaaS (Інфраструктура як послуга).

-Відповідальність провайдера фізична безпека дата-центрів, базова мережева інфраструктура, гіпервізори.

- Відповідальність клієнта: *найвища*. Включає безпеку операційних систем (встановлення оновлень, патчів), проміжного ПЗ, додатків, конфігурацію мережевих правил (фаєрволи, групи безпеки), управління ідентифікацією та доступом (IAM), шифрування даних.

- Ключові заходи безпеки для клієнта. Регулярне сканування на вразливості та оновлення, налаштування мережевої сегментації та правил брандмауера, впровадження надійних політик IAM (принцип найменших привілеїв, багатофакторна автентифікація - MFA), шифрування даних при зберіганні та передачі, резервне копіювання.

2. PaaS (Платформа як послуга):

- відповідальність провайдера розширюється на ОС, проміжне ПЗ (наприклад, бази даних, середовища виконання);

- відповідальність клієнта зосереджена на безпеці розроблених або розгорнутих додатків та даних, управлінні доступом користувачів до платформи та додатків, конфігурації сервісів платформи (наприклад, налаштування безпеки бази даних, що надається як сервіс);

- ключові заходи безпеки для клієнта: безпечна розробка коду (Secure SDLC), захист API, управління залежностями та вразливостями в коді, конфігурація політик доступу на рівні платформи та додатків, шифрування чутливих даних у додатках, моніторинг активності додатків.

3. SaaS (Програмне забезпечення як послуга):

- відповідальність провайдера найвища, він управляє майже всім стеком: інфраструктурою, ОС, платформою, самим додатком.

- відповідальність клієнта найменша, але критично важлива; включає управління доступом користувачів до сервісу, конфігурацію параметрів безпеки, доступних у додатку (наприклад, налаштування політик паролів, MFA, ролей доступу), класифікацію та захист даних, що завантажуються або створюються в сервісі, розуміння та контроль за тим, як дані обробляються провайдером;

- ключові заходи безпеки для клієнта: ретельний вибір провайдера (перевірка його сертифікацій безпеки та політик), налаштування надійних політик IAM та MFA, використання інструментів CASB (Cloud Access Security Broker) для видимості та контролю, навчання користувачів правилам безпечного використання сервісу, розуміння політик збереження та видалення даних.

Загальні практики безпеки для всіх моделей наступні.

1. Управління ідентифікацією та доступом (IAM). Використання сильних паролів, MFA, принципу найменших привілеїв, регулярний перегляд прав доступу.
2. Захист даних: Шифрування даних при зберіганні та передачі, класифікація даних, впровадження політик запобігання втраті даних.
3. Моніторинг та реагування. Збір та аналіз логів (аудит дій користувачів, події безпеки), використання систем SIEM (Security Information and Event Management), розробка плану реагування на інциденти.
4. Управління конфігураціями. Застосування безпечних базових конфігурацій, регулярний аудит налаштувань для виявлення відхилень.
5. Оцінка ризиків та відповідність. Регулярна оцінка ризиків, пов'язаних із хмарним середовищем, та забезпечення відповідності регуляторним вимогам.

Висновки. Хмарні обчислення надають значні переваги, але безпека залишається ключовим пріоритетом. Ефективний захист даних та інфраструктури в хмарі вимагає чіткого розуміння моделі спільної відповідальності для обраного типу послуг (IaaS, PaaS або SaaS) та впровадження відповідних заходів безпеки на рівні клієнта.

Незалежно від моделі, критично важливими є надійне управління ідентифікацією та доступом, шифрування даних, постійний моніторинг, а також правильна конфігурація сервісів. Для IaaS клієнти несуть найбільшу відповідальність за безпеку інфраструктурних компонентів. У PaaS фокус зміщується на безпеку додатків та даних. У SaaS основна увага приділяється управлінню доступом користувачів, конфігурації сервісу та контролю над даними.

Побудова комплексної стратегії хмарної безпеки, що враховує специфіку використовуваних сервісів та покладається на передові практики, дозволяє організаціям мінімізувати ризики та впевнено використовувати переваги хмарних технологій. Безпека в хмарі – це не одноразове завдання, а безперервний процес, що вимагає постійної уваги, адаптації та вдосконалення.

Список використаної літератури

1. Application Container Security Guide
<https://nvlpubs.nist.gov/nistpubs/SpecialPublications/NIST.SP.800-190.pdf>
2. AWS Shared Responsibility Model
https://aws.amazon.com/compliance/shared-responsibility-model/?nc1=h_ls
3. Cloud Controls Matrix (CCM)
<https://cloudsecurityalliance.org/research/cloud-controls-matrix>
4. Cloud security best practices center <https://cloud.google.com/security/best-practices>

ПУБЛІЧНІ WI-FI МЕРЕЖІ ТА ЗАХИСТ ОСОБИСТОЇ ІНФОРМАЦІЇ

Подвиженко А.В.

вчитель інформаційних технологій Фахового бізнес-коледжу

Переверзєв А.М.

вчитель інформаційних технологій Фахового бізнес-коледжу

ПВНЗ «Європейський університет»

Зручність публічних мереж Wi-Fi у різних громадських місцях, таких як кафе, аеропорти, готелі та бібліотеки, незаперечна в нашу сучасну цифрову епоху. Проте вкрай важливо визнати, що в цих мережах часто відсутні достатні заходи безпеки, що створює потенційні ризики для конфіденційності особистої інформації користувачів. У цьому дослідженні розглядаються основні вразливості та досліджуються ефективні стратегії захисту особистих даних під час використання загальнодоступних мереж Wi-Fi.

Основні ризики, пов'язані з персональними даними при використанні публічних мереж Wi-Fi наступні.

1. Перехоплення даних. Зловмисники мають просте завдання перехопити дані, що передаються через загальнодоступні мережі через відсутність шифрування. Відсутність безпеки призводить до розкриття конфіденційної інформації, такої як паролі, особисті дані та інші конфіденційні відомості.

2. Атаки типу Man-in-the-Middle також відомі як атаки перехоплення, викликають серйозне занепокоєння у сфері кібербезпеки. Ці атаки відбуваються, коли зловмисник перехоплює та змінює зв'язок між двома сторонами без їх відома чи згоди. Це форма прослуховування, яка дозволяє зловмиснику отримати несанкціонований доступ до конфіденційної інформації, такої як облікові дані для входу, фінансові дані або особисті дані. Процес атаки «Людина посередині» передбачає, що зловмисник розташовується між відправником і одержувачем повідомлення. Вони можуть досягти цього, використовуючи вразливі місця в мережевій інфраструктурі або використовуючи оманливі методи, такі як створення шахрайських мереж Wi-Fi або розгортання фішингових електронних листів. Після того, як зловмисник успішно підключився до каналу зв'язку, він може контролювати та маніпулювати даними, що передаються. Це може включати захоплення та аналіз трафіку для отримання цінної інформації або впровадження шкідливого коду для порушення цілісності зв'язку. Наслідки атаки «Людина посередині» можуть бути серйозними. Це може призвести до крадіжки особистих даних, фінансових втрат, несанкціонованого доступу до конфіденційних систем і навіть компрометації національної безпеки. Тому для окремих осіб і організацій вкрай важливо застосовувати надійні заходи безпеки, такі як протоколи шифрування та безпечні канали зв'язку, щоб захистити від цих атак. Підсумовуючи, можна сказати, що атаки Man-in-the-Middle становлять значну загрозу для конфіденційності, цілісності та доступності конфіденційної інформації. Розуміючи тактику зловмисників і застосовуючи відповідні заходи безпеки, ми можемо зменшити ризики, пов'язані з цими атаками, і захистити наші цифрові активи. Перехоплюючи або змінюючи дані, зловмисники мають можливість опинитися між користувачем і точкою доступу. В

результаті вони отримують несанкціонований доступ до особистих повідомлень, логінів і паролів.

3. Підробні зони Wi-Fi. Зловмисники мають можливість створювати підроблені точки доступу, які неймовірно схожі на справжні мережі. Користувачі, які нічого не підозрюють, можуть несвідомо підключатися до цих оманливих мереж, залишаючи свої дані вразливими для перехоплення або використання для подальших атак.

4. Шкідливе програмне забезпечення. Загальнодоступні мережі можна використовувати для розповсюдження вірусів і шкідливих програм, які збирають конфіденційну інформацію або залишають зловмисникам віддалено отримувати доступ до пристроїв.

Способи захисту персональних даних наступні.

1. Використання віртуальної приватної мережі (VPN). VPN шифрує весь трафік між користувачами та сервером VPN, що робить перехоплення даних практично неможливим. Це один з найефективніших методів захисту в публічних мережах Wi-Fi.

2. Шифрування даних. Використання веб-сайту з використанням протоколу HTTPS забезпечує шифрування даних між браузером і веб-сервером. Це гарантує, що дані, що передаються між користувачами і сайтом, не будуть перехоплені.

3. Автентифікація та складні паролі. Використовуйте двофакторну автентифікацію та складні паролі, щоб зменшити ризик несанкціонованого доступу до вашого запису. Двофакторна автентифікація додає додатковий рівень безпеки, вимагаючи додаткового коду підтвердження.

4. Оновлення програмного забезпечення. Регулярне оновлення операційної системи та програмного забезпечення може допомогти захистити від відомих вразливостей. Нові версії програмного забезпечення часто відповідають патчі безпеки, які усувають уразливість, якими можуть бути зловмисники.

5. Використовуйте мобільну мережу, якщо можливо, найкраще використовувати мобільні мережі (3G, 4G, 5G) для доступу до Інтернету, тому вони забезпечують вищий рівень безпеки, ніж загальнодоступні мережі Wi-Fi.

Поради користувачам публічної мережі Wi-Fi наступні.

1. Уникайте введення конфіденційної інформації. Не вводьте паролі та особисті дані під час підключення до публічної мережі Wi-Fi. Використовуйте безпечні мережі для здійснення фінансових операцій або доступу до важливої інформації.

2. Вимкнути автоматичне підключення. Вимкнути автоматичне підключення до відкритої мережі, щоб запобігти випадковим підключенням до незахищених точок доступу. Це дозволить уникнути підключення до фейкових мереж.

3. Перевірте наявність шифрування. Використовуйте лише мережу Wi-Fi, яка пропонує шифрування (WPA3 або WPA2). Це забезпечити додатковий рівень захисту вашої інформації.

Отже, захист особистих даних при використанні публічних мереж Wi-Fi є важливою складовою безпечного використання Інтернету. Впровадження заходів захисту, таких як використання VPN, шифрування даних, двофакторна аутентифікація та регулярне оновлення програмного забезпечення, допоможе значно зменшити ризики, пов'язані з використанням відкритих мереж. Користувачам також

варто дотримуватися основних рекомендацій щодо безпеки, щоб забезпечити захист своїх даних під час роботи в публічних WiFi мережах. Повністю забезпечити свій захист в загальнодоступних Wi-Fi мережах - неможливо, хіба що не користуватись Wi-Fi.

Список використаної літератури

1. *Комп'ютерні мережі: [Книга 1. Технології комп'ютерних мереж]: Навчальний посібник / Євсєєв С.П., Дженюк Н.В., Толкачов М.Ю та ін. – Харків, – Львів: Видавництво ПП «Новий Світ – 2000», 2024р.*
2. *Gupta, R., Madan, S., Sahu, I., & Goyal, R. (2023). The Dangers of Public Wi-Fi: How to Protect Your Device Through Cyber Security. International Journal on Recent and Innovation Trends in Computing and Communication, 11(9), 3016–3022.*

ПІДВИЩЕННЯ СТІЙКОСТІ КІБЕРБЕЗПЕКИ УКРАЇНИ В УМОВАХ ВІЙНИ ШЛЯХОМ ВПРОВАДЖЕННЯ ЕФЕКТИВНИХ ПРАКТИК АВТОМАТИЗОВАНОГО ТЕСТУВАННЯ БЕЗПЕКИ

*Позняк А.А.
викладач Фахового бізнес-коледжу,
ПВНЗ «Європейський університет»*

У сучасних умовах війни кібербезпека стала критичним компонентом національної безпеки України. Однією з ключових проблем є відсутність системного, масштабованого та вимірюваного підходу до забезпечення безпеки програмного забезпечення. Традиційні методи ручного тестування не здатні впоратися з високим темпом розробки, великою кількістю компонентів і постійною зміною загроз.

Ефективним рішенням цієї проблеми є впровадження автоматизованого тестування безпеки (security autotesting) на всіх етапах життєвого циклу програмного забезпечення. Цей підхід дозволяє не лише своєчасно виявляти вразливості, але й формувати відтворювані, вимірювані та контрольовані процеси забезпечення безпеки.

На прикладі практик, запропонованих командою Cossack Labs, автоматизоване тестування включає в себе:

- інтеграцію security-тестів у CI/CD пайплайни;
- використання тестів для контролю стабільності безпеки між релізами;
- створення метрик, які дозволяють оцінити рівень захищеності систем;
- використання спеціалізованих тестів для перевірки шифрування, контролю доступу, захисту даних тощо.

Особливої актуальності такі практики набувають у сфері публічних цифрових сервісів України (зокрема, у проєктах «Дія», цифрових реєстрах), де постійні оновлення потребують контрольованого процесу безпеки. Застосування автоматизованих security-тестів дозволяє гарантувати, що кожна нова функція чи

зміна не знижує рівень безпеки, а навпаки — підтверджується тестами, орієнтованими на сучасні кіберзагрози.

У військовому та критичному інфраструктурному ІТ-секторі це також дозволяє пришвидшити цикл оновлень без втрати контролю над безпековими аспектами, що є вирішальним для збереження функціональності в умовах кібер конфліктів.

Таким чином, масштабування практик автоматизованого тестування безпеки є одним із ключових елементів побудови кіберстійкої, передбачуваної та захищеної цифрової екосистеми України, що забезпечує швидке та безперервне сканування на вразливості, що критично важливо під час постійних атак. Автоматизація також зменшує кількість людських помилок та гарантує стандартизовані процедури тестування, що допомагає виявляти широкий спектр загроз, включаючи шкідливе програмне забезпечення та DDoS-атаки.

Отже, інтеграція автоматизованого тестування в процес розробки допомагає виявляти та усувати вразливості до того, як ними зможуть скористатися зловмисники. Постійна перевірка засобів контролю безпеки та виявлення слабких місць підвищує загальну кіберстійкість. Крім того, автоматизація звільняє фахівців з безпеки для виконання більш стратегічних завдань.

Список використаної літератури

1. Cossack Labs. *Security autotests for measurable and stable application security processes* [Електронний ресурс]. — 2023. — Режим доступу: <https://www.cossacklabs.com/blog/security-autotests-for-measurable-and-stable-application-security-processes>
 2. Apple Developer Documentation. *Security Overview* [Електронний ресурс]. — Apple Inc., 2023. — Режим доступу: <https://developer.apple.com/documentation/security>
 3. NIST. *Security and Privacy Controls for Information Systems and Organizations (Special Publication 800-53 Rev. 5) / Joint Task Force*. — National Institute of Standards and Technology, 2020. — 487 p.
- Федорів І.І., Савчук М.М. *Кібербезпека в інформаційних системах: навч. посібник*. — Київ: КНЕУ, 2021. — 224 с.
4. Мельник В.С. *Захист інформації в комп'ютерних системах: навч. посібник*. — Харків: ХНУРЕ, 2022. — 198 с.

СОЦІАЛЬНА ІНЖЕНЕРІЯ В УМОВАХ ВІЙНИ: КОНЦЕПЦІЯ СТВОРЕННЯ БОТА ДЛЯ ВИЯВЛЕННЯ ЦИФРОВИХ ЗАГРОЗ У МЕСЕНДЖЕРАХ

*Д.Ю. Покидько,
аспірант,
ПВНЗ «Європейський Університет»*

Соціальна інженерія залишається ключовим інструментом кібератак, особливо в умовах війни в Україні, де фішингові кампанії та дезінформація в Telegram підсилюють цифрові загрози [1-8]. Оскільки наявні рішення, як антивірусні боти чи інструменти перевірки посилань, не повністю враховують специфіку соціальної інженерії, розробка нового бота для безпечного аналізу вмісту стає актуальним

рішенням. Цей бот має аналізувати фото, текст повідомлень і поведінку відправника, оцінюючи ймовірність шахрайства, щоб захистити користувачів від необережних дій.

Запропонований бот задуманий як зручний помічник, що перевіряє підозрілі повідомлення в Telegram. Наприклад, якщо хтось отримує фото чи текст від невідомого контакту, він просто пересилає вміст боту, а той швидко оцінює, чи є ризик. Бот аналізує метадані файлів, шукає в тексті ознаки фішингу, як-от підозрілі слова чи посилання, і звертає увагу на те, хто відправив повідомлення. Результат видається у вигляді простої оцінки, наприклад: «Ймовірність шахрайства — 90%». Це дозволяє користувачу одразу зрозуміти, чи варто відкривати файл або клікати на посилання.

Для створення бота було використано Python через його простоту і доступність бібліотек. Основою став набір HTTP-методів Telegram API, який дозволяє легко налаштувати бота через @BotFather. Наприклад, бібліотека python-telegram-bot допомагає обробляти вхідні повідомлення та файли. Фото, отримані від користувача, можна перевіряти за допомогою бібліотеки Pillow, яка аналізує метадані на наявність прихованих URL чи скриптів. Якщо в повідомленні є текст, бот шукає ключові слова, характерні для фішингу, як-от «терміново» чи «введи код», і перевіряє посилання через сервіси типу VirusTotal. Поведінка відправника теж важлива: якщо це новий акаунт або невідомий контакт, ризик зростає.

Щоб оцінити, наскільки повідомлення може бути шахрайським, бот комбінує кілька факторів. Наприклад, підозрілі метадані у фото додають певний рівень ризику, так само як фішингові слова в тексті чи невідомий відправник. Усі ці фактори об'єднуються в загальну оцінку, яка показує ймовірність шахрайства у відсотках. Такий підхід допомагає користувачу швидко зрозуміти ситуацію без складних пояснень. У майбутньому оцінку планується зробити точнішою, навчивши бота за допомогою машинного навчання. Для цього знадобиться зібрати приклади реальних фішингових і безпечних повідомлень, щоб модель могла розпізнавати закономірності.

Щоб бот залишався актуальним, він може підключатися до баз даних про кіберзагрози. Наприклад, PhishTank пропонує API для перевірки фішингових посилань, а звіти CERT-UA можна використати для створення локальної бази загроз, актуальних для України. VirusTotal також стане в пригоді для сканування файлів і посилань, хоча безкоштовна версія має обмеження. Підключення до баз даних дозволяє боту швидше виявляти відомі загрози і попереджати користувачів.

Щоб переконатися, що бот працює адекватно, його протестували на реальних прикладах. В одному з варіантів успішного тестування використовувались підозрілі повідомлення, які отримує користувач. Було зібрано кілька фішингових і кілька безпечних повідомлень для класифікації ботом. Далі оцінювалось, наскільки точними є висновки бота. Для перевірки була залучена фокус-група, яка надсилала боту свої повідомлення, що дало змогу оцінити зручність створеного бота та отримати зворотній зв'язок від користувачів. В результаті перевірки бот правильно визначив більшість загроз і був достатньо швидким, що не дратувало користувачів і було зручним у використанні.

Перевірка підтвердила ефективність бота, його можна запропонувати ширшій аудиторії. Для цього було перекладено всі повідомлення українською мовою і зроблені максимально простими інструкції. Розміщення бота на хмарному сервері, наприклад, AWS, забезпечить його стабільну роботу. Щоб потенційні користувачі дізналися про рішення, в планах співпраця з Мінцифри чи CERT-UA, які просувають кібербезпеку в Україні. Публікація коду на GitHub також дозволить іншим розробникам долучитися до вдосконалення. Частину коду представлено на рис. 1.

```
from PIL import Image
from PIL.ExifTags import TAGS

def analyze_file(file_path):
    img = Image.open(file_path)
    exif_data = img._getexif()
    if exif_data:
        for tag_id, value in exif_data.items():
            tag = TAGS.get(tag_id, tag_id)
            if "http" in str(value): # Якщо є посилання в метаданих
                return "Увага: Фото містить підозрілі дані (90% шахрайство)"
    return "Фото виглядає безпечним"
```

Рис. 1 Перевірка метаданих файлу

Таким чином, створення нового бота для безпечного аналізу вмісту не просто захищає окремих користувачів, а й зменшує вплив соціальної інженерії в ширшому контексті. У час війни фішингові атаки часто стають частиною інформаційних операцій, спрямованих на поширення паніки чи збір даних. Бот, який допомагає розпізнавати загрози в Telegram, може стати частиною національної стратегії кіберзахисту, особливо якщо врахувати популярність цього месенджера в Україні.

Список використаної літератури

1. Державна служба спеціального зв'язку та захисту інформації України, «Аналіз фішингових кампаній, спрямованих на українські інформаційні ресурси у 2022–2024 роках», Київ, Україна, 2024. [Online]. Доступно: <https://cip.gov.ua/>.
2. Жмурко О., «Соціальна інженерія як загроза кібербезпеці: методи запобігання та захисту», Педагог. безпеки, т. 9, № 1, с. 37–42, 2024. doi: 10.31649/2524-1079-2024-9-1-037-042.
3. Кіберполіція України, «Комп'ютерні віруси та як від них захиститися: поради кіберполіції», Київ, Україна, 2023. [Online]. Доступно: <https://cyberpolice.gov.ua/>.
4. Міністерство цифрової трансформації України, «Кібергігієна для громадян: як захиститися від фішингу та шахрайства в месенджерах», Київ, Україна, 2024. [Online]. Доступно: <https://thedigital.gov.ua/>.
5. Рада національної безпеки і оборони України, «Річний аналітичний огляд: ключові події, тенденції та виклики у сфері кібербезпеки (жовтень 2023 – вересень 2024)», Київ, Україна, 2025. [Online]. Доступно: <https://www.rnbo.gov.ua/>.
6. CERT-UA, «Звіт про кіберінциденти в Україні за 2023–2024 роки», Державна служба спеціального зв'язку та захисту інформації України, Київ, Україна, 2024. [Online]. Доступно: <https://cert.gov.ua/>.
7. Шевчук О., «Кібербезпека у 2025 році: виклики, технології та ключові аспекти», ІТС.ua, 2024. [Online]. Доступно: <https://itc.ua/ua/statti/kiberbezpeka-u-2025-rotsi-vyklyky-tehnologiyi-ta-klyuchovi-aspekty/>.

8. Якименко Ю. М., Рабчун Д. І., та Запорожченко М. М., «Місце соціальної інженерії в проблемі витоку даних та організаційні аспекти захисту корпоративного середовища від фішингових атак з використанням електронної пошти,» *Кібербезпека: освіта, наука, техніка*, т. 1, № 13, с. 6–15, 2021. doi: 10.28925/2663-4023.2021.13.615.

ШЛЯХИ ПІДВИЩЕННЯ КІБЕРСТІЙКОСТІ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ: СУЧАСНІ МОВИ ПРОГРАМУВАННЯ ТА СУПУТНІ МЕТОДОЛОГІЇ

*Рихальський О.Ю.,
викладач кафедри комп'ютерних наук та програмної інженерії
ПВНЗ «Європейський університет»,
ст. викладач «DAN.IT Education».*

У сучасних умовах широкого впровадження цифрових технологій програмне забезпечення відіграє ключову роль у функціонуванні критично важливих систем — від фінансової сфери й медицини до безпеки, енергетики та промислової автоматизації. З огляду на це, питання кіберстійкості — здатності програмних рішень протистояти зловмисним впливам, залишаючись функціональними та безпечними — набуває першочергового значення.

Одним із базових чинників кіберстійкості є коректність програмного забезпечення, тобто відповідність його поведінки формалізованим вимогам і специфікаціям у всіх допустимих ситуаціях [1]. Недоліки в архітектурі, логіці чи реалізації програм призводять до появи вразливостей, які потенційно використовуються для здійснення кібератак. З огляду на це, актуальним є пошук шляхів зменшення ймовірності виникнення таких вразливостей ще на ранніх етапах життєвого циклу програмного продукту.

Сучасні мови програмування, зокрема такі як Scala [2], Rust [3], пропонують концептуальні й інструментальні засоби, які сприяють підвищенню безпеки програм на рівні компіляції. Йдеться про контроль доступу до пам'яті, контроль типів та обов'язкову обробку всіх варіантів значень, запобігання нульовим посиланням, які знижують імовірність появи типових помилок.

Крім мов програмування, у забезпеченні кіберстійкості дедалі ширше застосовуються супутні методології: формальна верифікація, статичний аналіз коду, використання шаблонів проєктування, орієнтованих на безпеку, а також інструменти автоматизованої перевірки. У сукупності ці підходи утворюють підґрунтя для створення програм, здатних протистояти кіберзагрозам.

Попри зростання технічної зрілості індустрії програмного забезпечення, практики реальної розробки залишаються вразливими до системних помилок [1], що впливають на коректність, надійність і безпеку створюваних продуктів. Існує низка викликів, зумовлених не лише технологічними обмеженнями, але й організаційними, методологічними та кадровими факторами.

Однією з ключових проблем є надмірна орієнтація на швидкість розробки, особливо в умовах комерційного тиску або стартап-культури. Програмне

забезпечення часто створюється без належного технічного резерву на якісне проєктування, валідацію вимог, написання тестів і детальну перевірку логіки. Як наслідок, у коді залишаються суттєві прогалини.

Ігнорування обробку виняткових ситуацій, відкладаючи її «на потім» або покладаючись на загальний механізм винятків без аналізу конкретних ризиків. Наприклад, операція ділення без перевірки на нуль, доступ до елемента масиву без перевірки меж, чи обробка неініціалізованого або звільненого об'єкта — усе це призводить до `panic`, `null reference`, `segmentation fault` або логічних помилок із непередбачуваними наслідками.

Відсутність попереднього аналізу предметної області та бізнес-процесів, що призводить до неправильного моделювання сутностей, невірному визначення граничних значень, відсутності перевірки припустимих станів. Часто системи не мають чіткої специфікації, і логіка реалізується у вигляді набору припущень розробника, які не витримують перевірки реальними даними або змінюваними вимогами.

Відсутність перевірки крайових ситуацій, які не охоплюються тестуванням або взагалі не розглядаються в логіці коду. Наприклад, масив нульової довжини, значення, що дорівнює мінімальному або максимальному значенню, порожні файли, їх відсутність, неправильний формат тощо. Система стає крихкою, вразливою до несподіваних помилок і аварійної зупинки.

Відсутність перевірки вхідних даних — як за типом, так і за обсягом, структурою, форматом, розміром. «Неправильні» дані можуть спричинити некоректну роботу бізнес-логіки або компрометацію системи безпеки. Зокрема, часто зустрічаються переповнення буфера, SQL-ін'єкції, обробка надто довгих рядків, відсутність обмежень на кількість переданих параметрів, або перевантаження API надмірними запитами. Особливо критичною є ситуація, коли вхідні дані зчитуються повністю в пам'ять, без попереднього аналізу розміру. Це призводить до вичерпання оперативної пам'яті та відмови сервісу.

Неповні умовні конструкції, як-от відсутність `else` у гілці `if` або пропущені варіанти в операторах `switch`, `match` [2][3]. Такі помилки призводять до неочікуваної поведінки або залишають систему у неузгодженому стані.

Вихід за межі виділеної пам'яті, що виникає внаслідок неправильного доступу до масивів або буферів, або помилок у циклічній обробці колекцій.

Значна частина зазначених проблем пов'язана з відсутністю або недостатністю механізмів статичного аналізу та перевірки повноти логіки на етапі компіляції. Багато мов і компіляторів не здатні автоматично виявити необроблені варіанти, неперевірені параметри, часткові функції чи відсутність обробки крайових умов. Це створює об'єктивну потребу в використанні та інтеграції більш «суворих» інструментів — як-от статичні аналізатори коду, системи контрактного програмування, типи з явною обробкою помилок (`Option` [2], `Result` [3]), тощо. Усвідомлення цього факту є першим кроком до впровадження ефективніших інструментів, практик і методів, які забезпечують надійність цифрових систем ще на етапі розробки.

Більшість з окреслених проблем можуть бути вирішені «автоматично», за умови використання сучасних компіляторів та мов програмування але для вирішення деяких з них потрібно створювати та змінювати культуру розробки.

Контроль повноти обробки варіантів (exhaustiveness checking). Мови програмування Scala, Rust мають цей функціонал вбудований. Якщо в операторі match не перелічені всі можливі варіанти enum [3] / sealed trait [2] — цей код просто не буде скомпільованим. Також ці компілятори можливо налаштувати, щоб конструкція if-then не компілювалася, коли відсутній else.

Формування повних функцій та уникнення часткових обчислень. Прикладом таких функцій є конвертація $String \Rightarrow Int$, коли рядок не може бути конвертований в число, або $\frac{a}{b}$ без перевірки, що $b \neq 0$. Класичним вирішенням цієї проблеми є зміна сигнатури функції на $String \Rightarrow Option[Int]$. Вбудований тип $Option[A]$ підкреслює факт того що результат може і не бути отриманий.

Обов'язкова перевірка параметрів, особливо отриманих через API, коли ми не можемо їм «довіряти». Для вирішення цієї проблеми використовуються бібліотека Refined [4]. Функція яка задекларована $(Int \text{ Refined Positive}) \Rightarrow Int$ та буде викликана зі значенням -5 , просто не буде скомпільована.

Обробка великих обсягів даних. Обробка великих файлів або потоків даних без контролю за споживанням пам'яті призводить до переповнення пам'яті та аварійного завершення. Для запобігання цьому в мові програмування Scala існує бібліотека fs2 [2], а в мові Rust існує концепція iterators [3], які можуть зчитувати наступний об'єкт в пам'ять, коли обробка попереднього буде завершена.

Вихід за межі виділеної пам'яті неможливий в мовах програмування Rust, Scala, оскільки компілятор гарантує коректність, тому що операції описуються декларативно і програміст не звертається до пам'яті безпосередньо [2][3].

Деякі інструменти й підходи вимагають певної зміни мислення серед розробників. Наприклад, примус до обробки всіх варіантів у match може сприйматись як надлишковий. Водночас це ціна, яку варто платити за надійність і передбачуваність поведінки системи.

Отже, проблема забезпечення кіберстійкості залишається однією з найактуальніших у сучасній інженерії [1]. Аналіз помилок, що призводять до вразливостей, свідчить про те, що їх часто спричиняють не складні алгоритмічні помилки, а порушення базових принципів безпечної розробки. Сучасні мови програмування, орієнтовані на безпеку та надійність, такі як Rust, Scala тощо, пропонують інструменти, що дозволяють виявляти широкий спектр помилок ще на етапі компіляції. Вони доповнюються супутніми методологіями — формальними методами [1], статичним аналізом коду, практиками «secure by design» — які забезпечують системний підхід до створення стійких до атак програм. Таким чином, підвищення кіберстійкості можливе не лише за рахунок зовнішніх засобів захисту, а й через усвідомлений вибір мов програмування та інженерних практик. Такий підхід сприяє формуванню культури розробки та забезпечує надійну основу для створення ефективних та безпечних програм.

Список використаної літератури

1. *Correctness 2024: Eighth International Workshop on Software Correctness for HPC Applications // SC24-W: Workshops of the International Conference for High Performance Computing, Networking, Storage and Analysis, 17–22.11.2024 p., Atlanta, GA, USA. – IEEE, 2025. – <https://doi.org/10.1109/SCW63240.2024.00024>*
2. *Odersky M. Programming in Scala. – 5th ed. / M. Odersky, L. Spoon, B. Venners, F. Sommers. – Walnut Creek, CA: Artima Press, 2021. – 668 p.*
3. *Sharma R. Mastering Rust: Learn about memory safety, type system, concurrency, and the new features of Rust 2018 edition. – 2nd ed. / R. Sharma, V. Kaihlavirta. – Birmingham: Packt Publishing, 2019. – 554 p.*
4. *Herbert K., Plante D. Refining Existing Types in Scala With Refined. <https://www.baeldung.com/scala/refined-types>*

III-ОПТИМІЗОВАНІ SIEM-СИСТЕМИ: ПІДВИЩЕННЯ ШВИДКОСТІ ТА ТОЧНОСТІ ВИЯВЛЕННЯ ІНЦИДЕНТІВ

*Саморай О. К.
викладач кафедри кібербезпеки та захисту інформації,
Ващенко М. В.
викладач кафедри кібербезпеки та захисту інформації
ПВНЗ «Європейський університет»*

У сучасних умовах зростаючого обсягу подій безпеки й ускладнення загроз традиційні SIEM-системи дедалі частіше демонструють обмежену здатність своєчасно виявляти атаки без значної кількості хибних спрацьовувань. За оцінками, до 90 % усіх сповіщень у класичних SIEM є хибно позитивними [1], і це призводить до додаткової безкорисної роботи для аналітиків та затримок у реагуванні на справжні інциденти. У багатьох рішеннях кореляція подій реалізується через жорстко прописані правила, а аналітика поведінки користувачів обмежена простими правилами, тому складні, багатоетапні атаки часто пролітають повз увагу системи. Їх виявляють тоді, коли вже запізно.

Використання машинного навчання й штучного інтелекту в контексті SIEM систем дозволяє подолати ці обмеження. Замість фіксованих правил - III-модулі які навчані на реальних даних, автоматично корелюють різні події за допомогою моделей класифікації, виявляють аномалії і пріоритезують інциденти за ризиком. Такі підходи показують високу ефективність у виявленні нетипової поведінки та зменшенні навантаження на SOC спеціалістів в організації.

Загальна архітектура III-оптимізованої SIEM включає п'ять шарів: шар збору даних (логи, мережевий трафік, DNS, EDR), функціональне сховище для попереднього обчислення ознак у реальному часі, офлайн-середовище для поетапного тренування та валідації моделей, движок логічного виводу для детекції та кореляції в режимі реального часу, а також шар візуалізації й інтеграції зі SOAR. [2] Така схема дозволяє розділити pipeline на оффлан (тренування) та онлайн (детекція) етапи, забезпечуючи швидкість і надійність роботи системи.

Перший та найпоширеніший патерн — детекція аномалій за допомогою автоенкодерів і кластеризації. Автоенкодери навчаються відтворювати «нормальний» мережевий трафік, а події з похибкою відновлення вище встановленого порога вважаються аномальними. Застосування поетапного навчання дозволяє оновлювати моделі без зупинки сервісу та виявляти нові відхилення в реальному часі.

Другий патерн — графова кореляція подій. Вузли графа можуть представляти користувачів, хости й процеси, а ребра — події доступу чи мережеві сесії. За допомогою методів прогнозування зв'язку (link prediction) і баєсових мереж (Bayesian networks) формується єдина картина інциденту: наприклад, серія дрібних анонімних подій об'єднується в ланцюжок Advanced Persistent Threat. Це дозволяє виявляти багатоетапні атаки, які важко відстежити традиційними механізмами.

Сортування та пріоритизація інцидентів здійснюється за допомогою навчених ML-моделей. На вхід моделі надходять ознаки: час доби, геолокація IP, історичний ризиковий профіль користувача чи хоста. Модель навчається на розмічених SOC-аналітиками даних і присвоює новим подіям ризиковий бал. Зворотний зв'язок від експертів (feedback loop) дозволяє регулярно перетреновувати модель і підтримувати високу точність класифікації.

Автоматизоване реагування (SOAR+AI) інтегрується через playbook-сценарії та ChatOps-ботів у Slack чи Telegram. Модуль AI пропонує варіанти дій — ізолювати хост, заблокувати IP, запустити forensic-скрипт — і надсилає їх аналітику з кнопками «Підтвердити» чи «Відхилити». Це скорочує час від виявлення до реагування з десятків хвилин до кількох секунд.

У ряді досліджень докладно описано методи забезпечення прозорості AI-рішень. Explainable AI-техніки, такі як SHAP і LIME, допомагають зрозуміти, які ознаки вплинули на рішення моделі, що критично для довіри з боку керівництва й відповідності нормативам (GDPR, NIST SP 800-53) [4].

Впровадження AI-SIEM відбувається поетапно: спочатку ґрунтовне очищення та нормалізація даних із формуванням функціонального сховища, далі вибір і валідація моделей, розгортання движка логічного виводу з Kafka/Flink, а потім організація feedback loop і моніторинг відхилень. Дослідження показують, що системи з MLOps-підходом підтримують стабільно високі показники навіть при динамічній зміні ландшафту кіберзагроз.

Насамкінець, AI-оптимізовані SIEM-системи довели свою здатність суттєво знижувати рівень хибно позитивних спрацьовувань, прискорювати реакцію SOC спеціалістів і адаптуватися до нових загроз. Подальший розвиток наряду лежить у глибшому впровадженні Explainable AI для підвищення довіри, автоматизації сценаріїв реагування з ШІ-асистентами та розширенні MLOps-інфраструктури для підтримки безперервного навчання і моніторингу моделей. Це відкриває нові можливості для побудови безпечних, масштабованих і автономних систем кіберзахисту на базі штучного інтелекту.

Список використаної літератури

1. Терейковський, І. А., Корченко, А. О., Парацук, Т. І., & Педченко, Є. М. (2018). Аналіз відкритих систем виявлення вторгнень. *Ukrainian Scientific Journal of Information Security*, 24(3), 201-216.
2. Субач, І. Ю., & Власенко, О. В. (2023). Архітектура інтелектуальної SIEM-системи для виявлення кіберінцидентів у базах даних інформаційно-комунікаційних систем військового призначення. *Системи і технології зв'язку, інформатизації та кібербезпеки*.
3. Lanz, Joel. «The Updated NIST Cybersecurity Framework.» *The CPA Journal* 94, no. 5/6 (2024): 70-72.
4. NIST SP 800-53 Rev. 5, «Security and Privacy Controls for Information Systems and Organizations» 2020.

АНАЛІЗ СУЧАСНОГО СТАНУ КІБЕРЗАГРОЗ ДЛЯ ОБ'ЄКТІВ КРИТИЧНОЇ ІНФРАСТРУКТУРИ ТА ІСНУЮЧИХ ПІДХОДІВ ДО ЇХ КЛАСИФІКАЦІЇ

*Світличний Д.О.,
аспірант кафедри кібербезпеки,
Національний технічний університет
«Харківський політехнічний інститут», Харків, Україна*

Об'єкти критичної інфраструктури (ОКІ) існують ще з давніх часів та їх захист є першочерговим для існування суспільства, країни або держави. У теперішній час серед існуючих загроз до найбільш актуальних можна віднести кіберзагрози. Їх дослідження, аналіз та класифікація стали потужними викликами для вчених в останні роки.

В сучасному світі ОКІ пов'язані між собою прямо або побічно, що може зумовити каскадне ураження ОКІ. У роботах [1, 2] виділяється кілька типів залежностей:

- фізичні (Physical interdependency),
- кібернетичні (Cyber (inter)dependency),
- географічні (Geographical (inter)dependency),
- логічні (Logical (inter)dependency),
- соціальні (Social dependency).

Такі залежності можуть існувати як всередині однієї інфраструктурної галузі, так і між різними секторами, що ускладнює оцінку ризиків та управління ними.

У роботі [2] розглянуто різні підходи до моделювання:

- модель вхідно-вихідної непрацездатності (Input-Output Inoperability Model (ІМ)), яка оцінює ефект каскадних збоїв;
- топологічні моделі, що використовують теорію мереж для аналізу взаємозв'язків між елементами інфраструктури;
- комплексні підходи, що поєднують якісні та кількісні методи для відображення складних міжсистемних залежностей.

Наведені підходи дозволяють провести пріоритетний та всебічний аналіз кіберзагроз для ОКІ.

Підходи до класифікації кіберзагроз мають розглядати не тільки показники пристроїв безпеки, які можна чітко виміряти (кількість виявлених вірусів і т.п.), а й інші якісні показники і критерії (інформаційна безпека персоналу, безпека користувачів та інше). У роботі [3] запропоновано класифікацію метрик кібербезпеки, що базується на шести доменах:

- критичність активів,
- типи послуг безпеки,
- наслідки атак,
- функціональна форма метрик,
- орієнтація на управління ризиками,
- фінансові витрати на атаки або захист.

Така структура дозволяє здійснювати багатовимірний аналіз кіберзагроз і заходів протидії для ОКІ.

Підхід до класифікації критеріїв, які не можливо чітко виміряти, розглянуто у роботі [4]. Головна ідея цієї класифікації базується на застосуванні таких основних критеріїв:

- конфіденційність (Confidentiality),
- цілісність (Integrity),
- доступність (Availability)

до важливих частин ОКІ (субдоменів):

- безпека програм та системи (Application and System Security),
- інформаційна безпека (Information Security), мережева безпека (Network Security),
- безпека вузлів (Endpoint Security),
- безпека критичної інфраструктури (Critical Infrastructure Security),
- управління, ризики та відповідність (Governance, Risk and Compliance),
- безпека мобільних пристроїв (Mobile Security),
- безпека користувача (End User Security),
- безпека зберігання даних (Storage Security).

Ці критерії використовуються для досягнення наступних цілей:

- визначити (Identify),
- захистити (Protect),
- знайти (Detect),
- відповісти (Respond),
- відновити (Recover).

Наведена архітектура дозволяє провести загальний аналіз захисту від кіберзагроз для ОКІ.

Таким чином, проведені дослідження показали сучасний стан кіберзагроз для ОКІ. Було розглянуто та проаналізовано існуючі підходи до класифікації кіберзагроз для ОКІ. Результати цього дослідження можуть бути використанні в подальшому удосконаленні існуючих підходів протидії кіберзагрозам для ОКІ.

Список використаної літератури

1. Roberto Setola, Eric Luijff, Marianthi Theocharidou (2017) *Cyber Threats Impacting Critical Infrastructures, Chapter 1 Critical Infrastructures, Protection and Resilience*
https://link.springer.com/chapter/10.1007/978-3-319-51043-9_1
2. Roberto Setola, Marianthi Theocharidou (2017) *Cyber Threats Impacting Critical Infrastructures, Chapter 1 Critical Infrastructures, Protection and Resilience*
https://link.springer.com/chapter/10.1007/978-3-319-51043-9_2
3. Yevseiev, S., Milov, O., Oprisky, I., Dunaievska, O., Huk, O., Pogorelov, V., Bondarenko, K., Zviertseva, N., Melenti, Y., Tomashevsky, B. (2022). Development of a concept for cybersecurity metrics classification. *Eastern-European Journal of Enterprise Technologies*, 4 (4 (118)), 6–18. doi: <https://doi.org/10.15587/1729-4061.2022.263416>
4. Noluntu Mpekoe (2024) *An Analysis of Cybersecurity Architectures*
https://www.researchgate.net/publication/379221397_An_Analysis_of_Cybersecurity_Architectures

АНАЛІЗ АНОМАЛЬНОЇ АКТИВНОСТІ В МЕРЕЖАХ ЗА ДОПОМОГОЮ МАШИННОГО НАВЧАННЯ

*Семенко А.І.,
доктор технічних наук, професор,
професор кафедри комп'ютерних наук та програмної інженерії
Бойко М.М.,
викладач кафедри комп'ютерних наук та програмної інженерії,
ПВНЗ «Європейський університет»*

У часи стрімкої цифровізації та зростання кіберзагроз особливої актуальності набуває виявлення нетипової поведінки у мережових середовищах. Багато атак обходять класичні засоби захисту, зокрема, фаєрволи чи антивірусне програмне забезпечення, завдяки використанню шкідливого коду, який не має відомих сигнатур [2]. У такому середовищі машинне навчання (ML) стає критично важливим інструментом для виявлення аномалій у мережевому трафіку, оскільки здатне розпізнавати навіть ті загрози, що не були відомі раніше [1].

Аномальна активність у мережі може вказувати на:

- несанкціонований доступ до системи;
- спроби експлуатації вразливостей;
- дії зловмисників із середовища самої організації;
- активність zero-day загроз або нових шкідливих програм.

Для виявлення таких відхилень застосовуються різні моделі ML. Найпоширенішими підходами є:

- метод опорних векторів (SVM);
- дерева рішень та ансамблеві методи (наприклад, Random Forest);
- глибокі нейронні мережі (DNN);
- методи кластеризації та виявлення відхилень без учителя (unsupervised learning).

Supervised-моделі використовуються у випадках, коли доступні попередньо марковані дані з прикладами нормальної та шкідливої активності. Натомість unsupervised-підходи дозволяють самостійно знаходити шаблони без попередніх

знань про тип загрози, що особливо корисно для zero-day атак або внутрішніх порушень.

Впровадження ML для аналізу аномалій у реальному часі вимагає:

- збирання, фільтрації та нормалізації мережевого трафіку;
- побудови векторів ознак, які якісно описують суть трафіку (наприклад, кількість з'єднань, частота запитів, обсяг переданих даних тощо);
- оптимізації моделей для забезпечення точності, масштабованості та швидкості роботи [1].

Проте такі системи мають ризик надлишкової чутливості, що може викликати помилкові спрацьовування (false positives), або, навпаки, прогавити загрозу при недостатньо точному налаштуванні. Важливою складовою також є періодичне перенавчання моделей для відображення змін у мережевій поведінці.

Отже, машинне навчання є потужним інструментом для виявлення аномальної активності у мережевому трафіку, здатним значно посилити ефективність кіберзахисту. Завдяки здатності обробляти великі масиви даних та адаптуватися до нових загроз ML-системи можуть стати основою нової генерації систем виявлення вторгнень. Однак для досягнення високої точності необхідно забезпечити якість вхідних даних, постійне оновлення моделей та поєднання автоматичних рішень із людською експертизою. Використання ML у кіберзахисті — це не лише технічне, а й стратегічне рішення, що потребує постійного вдосконалення.

Список використаної літератури

1. Ткаченко, С.В. (2022). Машинне навчання у виявленні кіберзагроз. *Кібербезпека: наукові підходи*, №1.
2. Яровий, Р., Улічев, О., Скляренко, О., Пашорін, В. (2024). Моделирование мультиагентных систем защиты информационных ресурсов. *Herald of Khmelnytskyi National University. Technical Sciences*, 337(3(2)), 278-284. <https://doi.org/10.31891/2307-5732-2024-337-3-42>

ОКРЕМІ АСПЕКТИ ВИКОРИСТАННЯ МЕТОДІВ OSINT

*Семенюк О.О.,
кандидат юридичних наук, доцент,
доцент кафедри права
ПВНЗ «Європейський університет»*

OSINT (Open Source Intelligence) — це метод збору та аналізу інформації з публічно доступних джерел для аналітичного використання у різних сферах, таких як журналістика, безпека та дослідження [1].

З розвитком цифрових технологій OSINT набуває все більшого значення в правоохоронній діяльності, кібербезпеці, бізнес-розвідці та журналістиці. Однак процес збору та обробки інформації вимагає дотримання правових норм, щоб уникнути порушень законодавства та прав людини.

Збір інформації у рамках OSINT здійснюється на основі таких принципів: Легальність – отримання даних лише з відкритих та доступних законним способом

джерел; Етичність – повага до приватного життя та прав людини; Достовірність – використання перевірених та надійних джерел. Анонімність – захист особистих даних дослідника та мінімізація ризику ідентифікації.

Безумовно з дотриманням вказаних принципів повинні застосовуватися і методи OSINT. В цілому методи OSINT можна поділити на активні та пасивні. Пасивними методами OSINT є способи збору інформації без безпосереднього втручання в діяльність об'єкта дослідження. Такі методи насамперед мінімізують ризик виявлення суб'єкта збору даних. Серед найбільш ефективних пасивних методів OSINT наступні.

- Аналіз пошукових систем. Використання пошукових систем (Google, Bing, Yandex, DuckDuckGo та ін.) для отримання даних з веб-сайтів, архівів і баз даних. Це може бути: використання операторів розширеного пошуку (Google Dorking) для точнішого знаходження інформації; пошук у кешованих версіях сторінок; використання інструментів зокрема, Wayback Machine для перегляду архівних версій сайтів; аналіз індексації веб-ресурсів для отримання прихованих або недостатньо захищених даних.

- Моніторинг соціальних мереж. Це дослідження активності користувачів у соціальних мережах (Facebook, Instagram, Twitter, LinkedIn, TikTok та ін.) для збору інформації про їхню діяльність, зв'язки та геолокацію. Таке дослідження включає зокрема: аналіз відкритих профілів і їхнього контенту (фото, відео, коментарі, вподобання, поширення); використання спеціалізованих сервісів для збору даних із соцмереж (наприклад, Maltego, Social Searcher, Twitonomy); відстеження зв'язків між користувачами через аналіз друзів, підписок і взаємодій; геолокаційний аналіз публікацій і метаданих зображень.

- Вивчення новинних порталів і блогів. Пошук згадок про події, особи, організації та тенденції у засобах масової інформації, блогах та форумах. Для цього застосовується: автоматизовані агрегатори новин (Google News, Yahoo News, RSS-канали); аналіз платформ (Medium, LiveJournal, Reddit, Telegram та ін.); використання сервісів для відстеження трендів і популярних тем у новинах; моніторинг фактчекінгових платформ для виявлення дезінформації.

- Робота з базами даних. Аналіз відкритих реєстрів і баз даних, що містять інформацію про компанії, фінансові операції, нерухомість, судові справи та ін. Такими найпоширенішими ресурсами є: публічні реєстри компаній (OpenCorporates, EDGAR, D&B Hoovers та ін.); земельні кадастри та реєстри нерухомості; судові реєстри, бази даних про банкрутства, фінансові звіти компаній; витоки баз даних, що стали загальнодоступними через злом або витік інформації.

- Збір інформації з даркнету (Dark Web). Даркнет є джерелом досить специфічної інформації, зокрема витоки даних, дані про нелегальну діяльність, а також анонімні форуми. Найпоширенішими методами збору даних з даркнету є: використання спеціалізованих браузерів (Tor, I2P) для доступу до закритих ресурсів; аналіз торгових майданчиків, форумів та обмінників даних; моніторинг витоків баз даних та паролів, що можуть бути опубліковані в даркнеті; виявлення злочинних груп або окремих осіб, що діють в тіньовому сегменті інтернету.

- Аналіз метаданих файлів. Метадані електронних ресурсів містять приховану інформацію, яка може розкрити важливі деталі про об'єкт дослідження.

Наприклад: отримання EXIF-даних із зображень для визначення місця, часу та пристрою зйомки; аналіз властивостей документів (Word, PDF, Excel) для отримання інформації про автора, дату створення та редагування; дослідження метаданих з електронної пошти та листування (наприклад, заголовки листів, IP-адреси відправників); виявлення та аналіз водяних знаків, що можуть свідчити про джерело походження файлу.

На відміну від пасивних, активними є методи OSINT які передбачають взаємодію з об'єктами дослідження, що може призвести до виявлення суб'єкта збору даних. До них відносять:

- Запити через відкриті API: використання програмних інтерфейсів (API) для отримання даних із публічних сервісів (наприклад, Twitter API, Google Maps API та ін.); збір інформації про акаунти, розташування, події та зв'язки;
- Взаємодія із соціальними групами та форумами. Реєстрація та участь у тематичних спільнотах для збору необхідної інформації; встановлення контактів з потенційними джерелами даних;
- Перевірка електронних адрес та телефонних номерів. Аналіз витоків баз даних для виявлення зв'язків та реєстрацій користувачів на різних сайтах; використання спеціалізованих сервісів (наприклад, Have I Been Pwned, PhoneInfoga, Getcontact та ін).
- Аналіз цифрових слідів. Перевірка, ресурсів пов'язаних з IP-адресою, доменом або іншим цифровим ідентифікатором; використання WHOIS-запитів, аналіз DNS-запитів та ін.
- Соціальна інженерія. До такого методу OSINT відносять: взаємодію з об'єктом дослідження для отримання необхідної інформації (метод вимагає дотримання етичних та правових норм); використання методів фішингу, соціальної маніпуляції та психологічного аналізу; Імітація довіри або створення обставин, що спонукають цільову особу до розголошення інформації; використання підроблених профілів або легендованих персонажів для встановлення контакту; аналіз поведінки людини та використання її когнітивних упереджень для отримання конфіденційної інформації.

Таким чином, методи OSINT відіграють ключову роль у зборі та аналізі інформації з відкритих джерел, що є важливим для кібербезпеки, правоохоронної діяльності, бізнес-аналітики та журналістики. Вони диференціюються на пасивні та активні, кожен із яких має свої переваги та ризики. Пасивні методи дозволяють непомітно збирати дані, використовуючи пошукові системи, соціальні мережі, новинні портали та відкриті бази даних. Активні методи, такі як API-запити, соціальна інженерія та аналіз цифрових слідів, можуть забезпечити глибший рівень дослідження, але водночас підвищують ризик викриття дослідника.

Важливим аспектом використання OSINT є дотримання правових норм та етичних принципів, зокрема легальності, достовірності, анонімності та поваги до приватного життя. Розвиток цифрових технологій постійно розширює можливості OSINT, роблячи його ефективним інструментом у багатьох сферах, проте водночас вимагає відповідального підходу до збору та використання інформації.

Список використаної літератури

1. OSINT (розвідка на основі відкритих джерел) [Електронний ресурс] // Termin.in.ua. – Режим доступу: <https://termin.in.ua/osint-rozvidka-na-osnovi-vidkrytykh-dzherel/> – Дата звернення: [28.03.2025].

ПОТЕНЦІЙНІ РИЗИКИ ЗАСТОСУВАННЯ ГЕНЕРАТИВНОГО ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ АВТОМАТИЗОВАНОГО СТВОРЕННЯ ДОКУМЕНТАЦІЇ ТА ТЕХНІЧНИХ СПЕЦИФІКАЦІЙ.

Сенишин Р.В.¹

Коцун В.І.²

¹аспірант,

*²к. т. н., доц., завідувач кафедри математики та комп'ютерних дисциплін,
Львівська філія ПВНЗ «Європейський університет»*

Генеративні моделі штучного інтелекту (ШІ), зокрема великі мовні моделі (LLM) на кшталт GPT-4, дедалі частіше використовують для швидкого формування інструкцій, користувацьких посібників і внутрішніх технічних специфікацій. Організації очікують зниження витрат на технічне письмо та скорочення часу виходу продукту на ринок. Однак стрімке впровадження технології випереджає усвідомлення її ризиків.

Висока переконливість машинного тексту поєднується з відсутністю гарантованої точності. Спільнота розробників уже відчула проблему: Stack Overflow тимчасово заборонив відповіді ChatGPT через «високий рівень неправильних, але правдоподібних відповідей» [1]. Коли подібні матеріали потрапляють у технічну документацію, помилки можуть трансформуватися у матеріальні збитки або загрозу безпеці, що робить дослідження ризиків особливо актуальним.

Найбільш очевидний технічний ризик — «галюцинації» ШІ, коли модель упевнено створює неіснуючі факти чи параметри. Показовий приклад — юридичний меморандум, у якому адвокати подали шість вигаданих судових справ, згенерованих ChatGPT; суд оштрафував їх на 5000 USD [2]. У складних інженерних проектах подібна вигадка може закласти в специфікацію критично помилкові величини.

Інша небезпека — неповнота. У випадку з літаком Boeing 737 MAX система MCAS взагалі не була згадана в мануалах для пілотів, що стало одним із чинників двох катастроф і 346 загиблих [3]. Якщо ШІ опустить опис важливого режиму роботи, наслідки можуть бути аналогічними.

Автоматизовані інструменти збільшують і ризик одиничних конвертаційних помилок. Історичний приклад — втрата апарата Mars Climate Orbiter через незбіг метричних та англійських одиниць у документації [4]. LLM здатна зробити таку ж «тиху» помилку, замаскувавши її бездоганним стилем.

Додатково існує ризик неусвідомленого плагіату: модель може відтворити фрагменти навчальних даних, порушивши авторські права. Це вимагає

формалізованих процедур експертної перевірки, журналу змін і тестів на унікальність.

Генеративний ШІ є цінним інструментом для підготовки чернеток, але не може слугувати автономним джерелом достовірної технічної інформації. Поєднання машинного письма зі строгою людською валідацією, перевіркою на відповідність стандартам і аудитом авторських прав — мінімально необхідна умова безпечного впровадження. Організаціям варто:

- розробити внутрішні політики роботи з LLM;
- визначити критичні розділи, де обов'язковий огляд фахівця;
- зберігати прозорий журнал редагувань;
- навчати персонал принципом «довіряй, але перевіряй».

Лише так автоматизація створення документації принесе ефективність, не створюючи неприйнятних технічних і юридичних ризиків.

Список використаної літератури

1. *Stack Overflow bans ChatGPT for 'substantially harmful' answers. The Register, 05.12.2022.* [Enterprise Technology News and Analysis](#)
2. *U.S. lawyers sanctioned for using fake ChatGPT cases. Reuters, 22.06.2023.* [Reuters](#)
3. *Duckworth, T. Boeing failed to disclose MCAS information to pilots. Сенат США, 2023.* [Tammy Duckworth](#)
3. *Mars Probe Lost Due to Simple Math Error. Los Angeles Times, 01.10.1999.* [Los Angeles Times](#)

ЗАХИСТ ПЕРСОНАЛЬНИХ ДАНИХ КОРИСТУВАЧІВ З ІНВАЛІДНІСТЮ У ПРОСТОРОВИХ ОБЧИСЛЮВАЛЬНИХ СЕРЕДОВИЩАХ

Скицький Т.Р.¹

Коцун В.І.²

¹ аспірант,

*² к. т. н., доц., завідувач кафедри математики та комп'ютерних дисциплін,
Львівська філія ПВНЗ «Європейський університет»*

Інноваційні технології просторових обчислень, зокрема доповнена/віртуальна реальність (AR/VR) та Інтернет речей (IoT), відкривають нові можливості для людей з інвалідністю, підвищуючи їхню автономність і доступ до цифрових середовищ [1]. У той же час, ці просторові системи збирають безпрецедентну кількість персональних даних користувачів, включаючи чутливу біометричну та контекстну інформацію, яка становить серйозний ризик для приватності. Незважаючи на значні переваги, використання просторових обчислень супроводжується багатьма викликами захисту персональних даних осіб з інвалідністю.

Персональні дані, які люди з інвалідністю генерують у просторових обчислювальних середовищах, відзначаються підвищеною чутливістю та різноманітністю. По-перше, інформація про стан здоров'я чи інвалідність належить до

особливих даних, які потребують особливого захисту [2]. По-друге, платформи AR/VR відстежують багато біометричних і поведінкових параметрів користувача, таких як рух очей, міміку обличчя, частоту серцевих скорочень, жести та переміщення у просторі тощо [3]. Ці дані збираються в режимі реального часу для створення інтуїтивного та адаптивного досвіду, але в той же час можуть розкривати приховану інформацію про особу. Зокрема, унікальні шаблони погляду чи ходи користувача будуть виконувати функцію своєрідного біометричного ідентифікатора, навіть якщо користувач діє під псевдонімом. По-третє, асистивні пристрої IoT (розумні протези, голосові помічники тощо) також генерують чутливі дані, наприклад голосові записи команд чи медичні показники та геолокацію користувача [4]. Такі дані дозволяють системам реагувати на запити користувача (напр., моніторити стан здоров'я чи забезпечувати безпечне пересування), але їх збір може мимоволі вказувати на діагноз користувача, його фізичні можливості чи повсякденну активність. Таким чином, у контексті просторових обчислень персональні дані людей з інвалідністю охоплюють як звичайну приватну інформацію, так й нові типи даних (біометричні, поведінкові, контекстуальні), які потенційно розкривають дуже детальний «портрет» користувача та його потреб. Це вимагає особливо ретельного захисту таких даних.

Користувачі з інвалідністю, що взаємодіють із AR/VR і IoT, піддаються тим же кіберзагрозам, що й інші, але враховуючи специфіку своїх пристроїв та сценаріїв використання ці загрози можуть мати більш серйозні наслідки. Найважливіші технічні загрози включають:

- Вразливості пристрою. Недостатньо захищене апаратне або програмне забезпечення в просторовій системі може стати точкою проникнення для зловмисника. Невиправлені помилки у прошивці медичних протезів чи інвалідних візків з підключенням до мережі можуть призвести до відмови в їх роботі. Такі атаки особливо критичні для людей з інвалідністю, які покладаються на ці пристрої у своєму повсякденному житті [4].
- Небезпеки мережевого рівня. Дані від сенсорів та пристроїв (напр., біометричні показники або відеопотік з камери AR-окулярів) переважно передаються через бездротові мережі. Без належного шифрування та захисту каналів зв'язку персональна інформація може бути перехоплена сторонніми особами [4].
- Соціальна інженерія та фішинг. Просторові технології відкривають новий фронт для фішингових атак. Зловмисники можуть створити правдоподібні віртуальні сцени або об'єкти в AR/VR, щоб ввести користувача в оману і змусити розкрити конфіденційну інформацію. Користувачі з інвалідністю, які повністю довіряють допоміжним цифровим підказкам, можуть бути мішенню таких атак, якщо не будуть введені додаткові механізми перевірки достовірності контенту [5].

Таким чином, спектр кіберзагроз у просторових обчисленнях є широким – від суто технічних (експлуатація вразливостей, мережеві атаки) до психологічних (фішинг у AR/VR). Для користувачів з інвалідністю наслідки таких атак особливо серйозні, оскільки вони часто довіряють технологіям критичні аспекти свого життя (здоров'я, мобільність, комунікацію).

Для ефективного захисту персональних даних у просторових обчислювальних середовищах необхідний багаторівневий підхід, який поєднує кілька технічних засобів. Основні методи включають:

- Шифрування даних. Надійне шифрування – це основний інструмент захисту: усі чутливі дані користувачів повинні передаватися та зберігатися в зашифрованому вигляді. Використання сучасних протоколів шифрування (TLS, end-to-end encryption) забезпечує конфіденційність інформації та унеможливорює її прочитання третіми особами навіть у випадку перехоплення трафіку [4].
- Анонімізація та псевдонімізація. Мінімізація зв'язку між зібраними даними та особою користувача знижує ризик ідентифікації. Впровадження практики анонімізації означає зберігання даних у безособовому вигляді: персональні ідентифікатори (ім'я, адреса, ID пристрою) відокремлюються від основного масиву інформації або замінюються на псевдоніми [6].
- Принципи «конфіденційність за задумом» (privacy by design). Ці принципи полягають у тому, що механізми захисту даних інтегруються в архітектуру та логіку системи з самого початку її розробки, а не додавалися постфактум. Для інклюзивних технологій це означає врахування потреб і ризиків для людей з інвалідністю на ранніх етапах проектування. Зокрема, розробникам потрібно за замовчуванням обмежувати обсяг даних, що збираються в користувача, та встановити високі налаштування захисту даних, надаючи користувачам можливість явно погоджуватися на збір додаткової інформації [6].

Важливою перспективою є розробка галузевих стандартів і керівництв з безпеки для інклюзивних технологій. Завдяки співпраці між виробниками пристроїв, розробниками ПЗ, науковцями та правозахисниками можна виробити узгоджені підходи для захисту даних користувачів з інвалідністю.

Вбудована й продумана політика безпеки «за задумом» не лише захищає користувачів, але й спрощує подальше дотримання більш жорстких регуляторних вимог (таких як закони про захист даних, стандарти щодо штучного інтелекту тощо). У довгостроковій перспективі інвестиції у безпеку і приватність окупляться підвищенням довіри споживачів та ширшим впровадженням технологій серед людей з інвалідністю. Як зазначають дослідники, інтеграцією приватності у сам дизайн інклюзивних систем, можна повністю розкрити потенціал просторових обчислень без жертвування фундаментальними правами користувачів [7]. Відтак, майбутнє просторових обчислювальних середовищ має розвиватися за принципом: інклюзивність + безпека = цифрова автономність без ризиків.

Список використаної літератури

1. Cogent Infotech. (2025, January 1). *Spatial Computing: The Next Frontier in Digital Transformation*. URL: <https://www.cogentinfo.com/resources/spatial-computing-driving-the-next-wave-of-digital-transformation>
2. Martinez, C. (2019, January 31). *FPF Report: IoT Devices Should Deal with Privacy Impacts for People with Disabilities*. *Future of Privacy Forum*. URL: <https://fpf.org/blog/fpf-report-iot-devices-should-deal-with-privacy-impacts-for-people-with-disabilities/>
3. TrustArc. (2024). *Privacy in Augmented and Virtual Reality Platforms: Challenges and Solutions for Protecting User Data*. URL:

<https://trustarc.com/resource/privacy-augmented-virtual-reality-platforms/>

4. Ibrahim, A., & Gebali, F. (2025). Enhancing Security and Efficiency in IoT Assistive Technologies: A Novel Hybrid Systolic Array Multiplier for Cryptographic Algorithms. *Applied Sciences*, 15(5), 2660. <https://doi.org/10.3390/app15052660>
5. Tripwire. (2024, November 14). Exploring the Security Risks of VR and AR. *Tripwire State of Security Blog*. URL: <https://www.tripwire.com/state-of-security/exploring-security-risks-vr-and-ar>
6. MirrAR. (2023). 3 Strategies for AR Data Privacy to Safeguard Customer Information in the Beauty-Tech Realm. URL: <https://www.mirrAR.com/blogs/3-strategies-for-augmented-reality-data-privacy-to-safeguard-customer-information-in-the-beauty-tech-realm>
7. Forward Pathway. (2025, April 1). Privacy and Security Challenges in VR/AR: Gaze Data, Military Applications, and Mitigation Strategies. URL: <https://www.forwardpathway.us/privacy-and-security-challenges-in-vr-ar-gaze-data-military-applications-and-mitigation-strategies>

БЕЗПЕЧНЕ ЗБЕРЕЖЕННЯ ТА ПЕРЕДАЧА ДАНИХ У МОБІЛЬНИХ ЗАСТОСУНКАХ

Скляренко О.А.,
аспірант,
Позняк А.А.,
викладач Фахового бізнес-коледжу,
ПВНЗ «Європейський університет»

У сучасних умовах мобільні застосунки стали інтегрованим елементом повсякденної діяльності, охоплюючи широке коло функцій — від комунікації та фінансових операцій до доступу до державних, медичних і сервісних послуг, а також керування технічними пристроями. Особливої значущості ці цифрові інструменти набули в умовах воєнного стану та нестабільності енергетичної інфраструктури, виконуючи критичні функції в системах оповіщення, координації дій під час надзвичайних ситуацій, волонтерських ініціативах та організації процесів евакуації.

Однак, зі зростанням функціональності застосунків спостерігається зростання обсягу чутливої інформації, що ними обробляється, що, у свою чергу, зумовлює виникнення низки серйозних викликів у сфері інформаційної безпеки [1-4]. Зокрема, актуальними залишаються питання захищеного зберігання даних на кінцевих пристроях, а також їх безпечного передавання мережею [5,6]. Серед найбільш поширених загроз, пов'язаних із цими процесами, можна визначити наступні: зберігання чутливих даних у незашифрованому вигляді (наприклад, токенів, логінів, паролів); використання незахищеного з'єднання, наприклад, передача паролів через HTTP; витоки через зламані або root/jailbreak-пристрої, де злоумисник може легко дістатися до файлів апки; фішинг, перехоплення трафіку (MITM-атаки), особливо в публічних Wi-Fi мережах.

Для уникнення вище перелічених загроз можна розбити рішення на наступні категорії: збереження, керування доступом, комунікація мережею, захист самого пакету застосунку. Щоб забезпечити безпечне збереження даних на пристрої

розробник має використовувати системні інструменти з можливістю шифрування, які відповідають складності й структурі даних, зазначені далі.

1. Локальні сховища з можливістю збереження простих неоднорідних даних типу ключ-значення:

- Android: EncryptedSharedPreferences, Keystore;
- iOS: Keychain, Secure Enclave.

2. Локальні бази даних з шифруванням складних систем даних з відносинами, наприклад, SQLCipher або шифрування баз у Room / Realm.

Також важливо уникати зберігання зайвих даних, які не є критично необхідними для роботи застосунку або відновлення сесії. Наприклад, якщо токен не потрібен, то краще його не зберігати взагалі.

Біометрична автентифікація — додатковий шар безпеки, завдяки якому можна надати або обмежити доступ до чутливих функцій завдяки біометричним даним, таким як Face ID, Touch ID, відбиток пальця. Перевагами цього механізму є те, що біометрія не передається на сервер, а зберігається локально в безпечних місцях, гарантованих операційною системою, а саме:

- iOS: LAContext (LocalAuthentication);
- Android: BiometricPrompt + захищений Keystore.

Для захисту передачі даних, що передаються між апкою та сервером потрібно дотримуватися таких рекомендацій: використовувати HTTPS (TLS 1.2 або новіше), перевіряти сертифікат сервера через SSL Pinning, краще застосовувати OAuth2 або JWT замість Basic Auth, мінімізувати дані в запитах, ніколи не передавати паролі у відкритому вигляді.

Навіть за умов правильного шифрування даних, злоумисник може отримати .apk або .ipa файл, розпакувати його та проаналізувати логіку, знайти API-ключі, URL-и серверів чи іншу конфіденційну інформацію. Цей процес називається реверс-інжиніринг, тобто вивчення структури застосунку без доступу до його вихідного коду. Його часто застосовують для виявлення вразливостей або створення неофіційних модифікацій. Для протидії цьому використовують обфускацію або «заплутування» коду, яке значно ускладнює зворотну інженерію в залежності від платформи:

- Android: ProGuard, R8;
- iOS: Obfuscator-LLVM (рідше використовується).

Крім того, важливо видаляти логування у продакшн-збірках, а також реалізовувати перевірки на root/jailbreak, захист від дебагу та використання емуляторів.

Таким чином, забезпечення захищеного зберігання та передавання даних у мобільних застосунках є не лише технічним завданням, а й ключовим чинником формування довіри користувачів, збереження репутації розробників і гарантування інформаційної безпеки національного рівня. У цьому контексті особливої ваги набуває дотримання принципу мінімізації обробки даних, який передбачає збирання та обробку виключно необхідної інформації з одночасним застосуванням сучасних засобів кіберзахисту. В умовах воєнного стану та зростаючої інтенсивності кібератак такий підхід перестає бути рекомендацією й трансформується у критично важливу складову безпечної цифрової інфраструктури.

Список використаної літератури

1. Stallings W. «Network Security Essentials: Applications and Standards. — 6th ed». — Pearson, 2020. — 480 p.
2. OWASP Foundation. «Mobile Security Testing Guide (MSTG)». [Електронний ресурс]. — OWASP, 2023. — Режим доступу: <https://owasp.org/www-project-mobile-security-testing-guide/>
3. «Android Developers. Security best practices» [Електронний ресурс]. — Google, 2024. — Режим доступу: <https://developer.android.com/topic/security/best-practices>
4. «Apple Developer Documentation. Security Overview» [Електронний ресурс]. — Apple Inc., 2023. — Режим доступу: <https://developer.apple.com/documentation/security>
5. NIST. «Security and Privacy Controls for Information Systems and Organizations (Special Publication 800-53 Rev. 5) / Joint Task Force.» — National Institute of Standards and Technology, 2020. — 487 p.
6. Яровий, Р., Улічев, О., Скляренко, О., & Пашорін, В. (2024). Моделювання мультиагентних систем захисту інформаційних ресурсів. *Herald of Khmelnytskyi National University. Technical Sciences*, 337(3(2), 278-284. <https://doi.org/10.31891/2307-5732-2024-337-3-42>

АТАКИ ЧЕРЕЗ НЕПРЯМІ ВЕКТОРИ У ЦИФРОВОМУ СЕРЕДОВИЩІ ЯК ІНСТРУМЕНТ ПОСИЛЕННЯ КІБЕРЗАГРОЗ ПРОТИ ПІДПРИЄМСТВ В УМОВАХ ГІБРИДНОЇ ВІЙНИ

*Скляренко О.В.,
кандидат фізико-математичних наук, доцент,
завідувач кафедри математичних дисциплін
та інноваційного проектування
Ващенко М.В.,
викладач кафедри кібербезпеки та захисту інформації
ПВНЗ «Європейський університет»*

В умовах сучасної гібридної війни кібератаки, що здійснюються через непрямі вектори у цифровому середовищі, набули нового стратегічного значення. Вони слугують не лише інструментом спричинення локального збитку окремим суб'єктам господарювання, а й засобом системного впливу, спрямованого на дестабілізацію економічної стійкості держави та підлив суспільної довіри до критично важливих галузей національної інфраструктури [1].

Особливість атак через непрямі вектори полягає у компрометації слабких ланок бізнес-екосистеми, зокрема, постачальників, партнерських організацій і зовнішніх сервісів. Згідно з дослідженням ENISA, атаки на ланцюги постачання стали критичним вектором для складних операцій проти високозначущих цілей [1]. IBM X-Force у своєму останньому звіті підкреслює, що у 2023 році атаки через supply chain призвели до найбільших економічних втрат серед усіх типів атак [2].

У воєнних умовах атаки через ланцюги постачання (supply chain attacks) виступають ефективним інструментом деструктивного впливу, спрямованого на виведення з ладу об'єктів критичної інфраструктури та мають на меті порушення функціонування фінансових систем, логістичних ланцюгів і телекомунікаційних мереж, що в сукупності сприяє послабленню державного управління та зниженню стійкості ключових секторів економіки. За даними Verizon DBIR, понад 60% атак

через ланцюг постачання залишаються непоміченими протягом перших 180 днів, що суттєво ускладнює оперативне реагування [3]. Водночас, CISA у своїх рекомендаціях наголошує на тому, що успішна атака через одного вразливого партнера може призвести до каскадного ефекту в усій галузі [4].

Низький рівень контролю над процесами розробки, інтеграції і постачання програмних продуктів є одним із основних чинників вразливості цифрового середовища. Стандарти безпечної розробки, як-от Secure Software Development Framework (SSDF) від NIST, мають бути інтегровані у всі етапи життєвого циклу програмного забезпечення [5].

CrowdStrike у своєму Global Threat Report вказує, що супротивники дедалі частіше націлюються не на безпосередніх кінцевих споживачів послуг, а на менш захищених технологічних провайдерів, щоб урадити основну ціль через непрямі канали [6]. Gartner серед основних трендів кібербезпеки на 2024 рік визначає зміщення пріоритету із захисту периметра на управління довірою до партнерів та постачальників [7].

Таким чином, атаки через непрямі вектори у цифровому середовищі стали потужним інструментом системного впливу на економічну безпеку. Протидія таким атакам вимагає комплексного підходу, що поєднує технічні засоби контролю, розвиток культури цифрової гігієни серед партнерів і впровадження міжнародних стандартів управління ризиками ланцюгів постачання.

Список використаної літератури

1. ENISA. *Threat Landscape for Supply Chain Attacks 2021*. European Union Agency for Cybersecurity. URL: <https://www.enisa.europa.eu/publications/threat-landscape-for-supply-chain-attacks>
2. IBM. *X-Force Threat Intelligence Index 2024*. IBM Corporation. URL: <https://www.ibm.com/reports/threat-intelligence>
3. Verizon. *Data Breach Investigations Report 2023*. Verizon Communications. URL: <https://www.verizon.com/business/resources/reports/dbir/2023/masters-guide-data-breach-investigations-report-2023.pdf>
4. CISA. *Defending Against Software Supply Chain Attacks*. Cybersecurity & Infrastructure Security Agency. URL: <https://www.cisa.gov/news-events/news/defending-against-software-supply-chain-attacks>
5. NIST. *Secure Software Development Framework (SSDF)*, 2022. National Institute of Standards and Technology. URL: <https://csrc.nist.gov/publications/detail/sp/800-218/final>
6. CrowdStrike. *Global Threat Report 2024*. CrowdStrike Holdings, Inc. URL: <https://www.crowdstrike.com/global-threat-report/>
7. Gartner. *Top Cybersecurity Trends for 2024*. Gartner, Inc. URL: <https://www.gartner.com/en/newsroom/press-releases/2023-10-16-gartner-identifies-top-cybersecurity-trends-for-2024>

АКТУАЛЬНІ ПИТАННЯ ЗАБЕЗПЕЧЕННЯ КІБЕРБЕЗПЕКИ: ЗАХИСТ КРИТИЧНОЇ ІНФРАСТРУКТУРИ

*Снігур М.М.,
магістрант факультету інформаційних систем та технологій,
ПВНЗ «Європейський університет»*

Критична інфраструктура (КІ) — це сукупність об'єктів, систем і мереж, які є життєво важливими для функціонування держави, економіки та суспільства [1]. Йдеться про енергетичні системи, транспорт, телекомунікації, охорону здоров'я, банківську інфраструктуру, водопостачання, оборону та органи державної влади. Будь-які перебої в роботі цих об'єктів можуть спричинити катастрофічні наслідки як для безпеки держави, так і для життя громадян [2].

Останні десятиліття цифрова трансформація охопила більшість сфер критичної інфраструктури. Впровадження ІТ-рішень, автоматизованих систем керування (наприклад, SCADA), віддалений доступ до систем управління, а також підключення до Інтернету зробили об'єкти КІ потенційними мішенями для кібератак. Проблема поглиблюється через застаріле програмне забезпечення, відсутність регулярних оновлень безпеки, використання стандартних паролів і недостатній рівень підготовки персоналу. Сучасні кіберзлочинці, у тому числі ті, що підтримуються державами, активно досліджують і використовують уразливості таких систем.

Одним із яскравих прикладів стала кібератака BlackEnergy на українську енергосистему у 2015 році, що призвела до відключення електроенергії для сотень тисяч громадян [3]. Інший приклад — атака на систему охорони здоров'я в Ірландії у 2021 році, яка паралізувала роботу лікарень [4]. Атаки на транспортну інфраструктуру можуть зупинити логістику, а на банківські системи — викликати паніку серед населення та вплинути на національну економіку.

У світі активно формуються підходи до захисту КІ. У США реалізується модель NIST Cybersecurity Framework, яка визначає п'ять функцій кіберзахисту: ідентифікація, захист, виявлення, реагування, відновлення [5]. У Європейському Союзі діє Директива NIS2, що встановлює мінімальні вимоги для операторів критичної інфраструктури [6]. В Україні прийнято Закон «Про основні засади забезпечення кібербезпеки», створено Державний центр кіберзахисту при Держспецзв'язку, а також затверджено національну стратегію кібербезпеки. Проте імплементація стандартів потребує послідовного та ресурсоємного підходу.

Захист критичної інфраструктури вимагає системного підходу. По-перше, необхідне проведення регулярного аудиту безпеки для виявлення слабких місць. По-друге, слід впроваджувати сегментацію мереж, багаторівневу автентифікацію, системи виявлення вторгнень (IDS/IPS), резервне копіювання. Важливою є підготовка персоналу — як технічного, так і управлінського рівня. Крім того, кожен об'єкт КІ повинен мати план реагування на кіберінциденти (Incident Response Plan), що дозволить оперативно діяти у випадку атаки.

Кіберзахист критичної інфраструктури — це не лише технічне питання, а комплексна проблема, що вимагає взаємодії між державою, бізнесом та

громадянським суспільством. В умовах зростання гібридних загроз, зокрема й з боку держав-агресорів, захист таких об'єктів стає ключовим елементом національної безпеки. Підвищення кіберстійкості можливе лише за умови об'єднання сучасних технологій, професійної підготовки кадрів та адекватного законодавчого регулювання.

Список використаної літератури

1. Закон України «Про основні засади забезпечення кібербезпеки України». Посилання: <https://zakon.rada.gov.ua/laws/show/2163-19>
2. Національна стратегія кібербезпеки України 2021 року. Посилання: <https://www.president.gov.ua/documents/4472021-40013>
3. Атака BlackEnergy на українську енергосистему (CERT-UA, 2015)
Посилання: <https://cert.gov.ua/article/17>
4. Атака на систему охорони здоров'я Ірландії (HSE, 2021). Посилання <https://www.hse.ie/eng/services/news/media/pressrel/cyber-attack.html>
5. NIST Cybersecurity Framework (USA). Посилання: <https://www.nist.gov/cyberframework>
6. Директива NIS2 (Європейський Союз)
Посилання: <https://digital-strategy.ec.europa.eu/en/policies/nis2-directive>
7. ENISA Threat Landscape Reports (2022–2023). Посилання: <https://www.enisa.europa.eu/publications>

ТЕХНОЛОГІЯ ЗАБЕЗПЕЧЕННЯ БЕЗПЕКИ ДОКУМЕНТІВ В СИСТЕМАХ ЕЛЕКТРОННОГО ДОКУМЕНТООБІГУ

*Стрілець Є.А .,
магістрант факультету інформаційних систем та технологій,
ПВНЗ «Європейський університет»*

У сучасних інформаційних системах електронний документообіг (ЕДО) став одним із ключових інструментів цифрової трансформації управлінських, юридичних та бізнес-процесів. Перехід від традиційного паперового документообігу до цифрового значно підвищив швидкість комунікацій, спростив зберігання та пошук інформації, а також дозволив автоматизувати процеси обробки даних. Водночас зростає і залежність від інформаційних технологій, що обумовлює необхідність гарантованого захисту електронних документів від несанкціонованого доступу, фальсифікації, втрати або знищення. Саме тому проблема забезпечення безпеки документів в системах ЕДО набула особливої актуальності.

Системи електронного документообігу оперують даними, які часто мають високу конфіденційність та юридичну значущість. Тому одним із ключових напрямів є впровадження ефективних технічних механізмів захисту інформації, які реалізують такі базові принципи: конфіденційність, цілісність, автентичність, доступність і невідомість. З метою досягнення цих цілей доцільно застосовувати сукупність технологій, що доповнюють одна одну.

Одним з фундаментальних методів захисту в ЕДО є використання криптографічних засобів. Сучасні алгоритми симетричного (AES) та асиметричного

(RSA, ECC) шифрування дозволяють ефективно захищати передані та збережені дані. Для забезпечення юридичної значущості документів застосовується кваліфікований електронний підпис (КЕП), який гарантує автентичність відправника та цілісність документа. У поєднанні з інфраструктурою відкритих ключів (PKI) система ЕДО може забезпечити високий рівень довіри між суб'єктами документообігу.

Окрім криптографії, важливу роль у побудові захищеної системи ЕДО відіграють системи управління доступом. Вони реалізуються за допомогою ролевих моделей, які дозволяють точно обмежити повноваження користувачів на читання, зміну або видалення документів. Для фіксації всіх операцій з документами впроваджується механізм аудиту та логування, що дозволяє відслідковувати всі дії в системі та оперативно виявляти потенційні порушення.

Особливу увагу в сучасних системах приділяють безпечній передачі документів каналами зв'язку. Для цього застосовуються протоколи захищеної комунікації (TLS, HTTPS), що забезпечують шифрування даних під час транспортування та захист від атак типу «man-in-the-middle». Крім того, в окремих сценаріях використання електронного документообігу можливе впровадження технологій блокчейн, які дозволяють фіксувати зміни документів у вигляді транзакцій у розподіленому реєстрі. Це забезпечує додатковий рівень прозорості та захисту від фальсифікацій.

Таким чином, технологія забезпечення безпеки документів у системах електронного документообігу повинна включати комплекс заходів технічного характеру: криптографію, електронний підпис, контроль доступу, аудит, захищені канали зв'язку та, за потреби, технології розподілених реєстрів. Такий інтегрований підхід дозволяє створити надійне середовище для збереження, обробки та обміну документами, що відповідає сучасним викликам у сфері кібербезпеки.

Список використаної літератури

1. Закон України «Про електронні документи та електронний документообіг» №851-IV від 22.05.2003.
2. Закон України «Про захист інформації в інформаційно-телекомунікаційних системах» №80/94-ВР від 05.07.1994.
3. ISO/IEC 27001:2022 — Information Security, Cybersecurity and Privacy Protection — Information Security Management Systems — Requirements.
4. Кузьмінський А.О., Сидоренко О.В. Системи електронного документообігу: захист та безпека. – Харків: ХНУРЕ, 2020.
5. Яровий, Р., Улічев, О., Скляренко, О., & Паширін, В. (2024). Моделювання мультиагентних систем захисту інформаційних ресурсів. *Herald of Khmelnytskyi National University. Technical Sciences*, 337(3(2)), 278-284. <https://doi.org/10.31891/2307-5732-2024-337-3-42>

ІННОВАЦІЙНІ МЕТОДИ ЗАХИСТУ КОНФІДЕНЦІЙНОЇ ІНФОРМАЦІЇ В ЕКОСИСТЕМІ SALESFORCE: ВИКЛИКИ ТА РІШЕННЯ

*Сушинський О.Є. ,
доктор технічних наук, доцент,
завідувач кафедри комп'ютерних наук та програмної інженерії*
*Коцун В.І. ,
кандидат технічних наук, доцент,
доцент кафедри комп'ютерних наук та програмної інженерії
ПВНЗ «Європейський університет»*

В умовах стрімкого розвитку цифрових технологій та зростаючої кількості кіберзагроз, захист конфіденційної інформації став критично важливим аспектом функціонування будь-якої організації. Екосистема_Salesforce, як одна з провідних платформ для управління взаємовідносинами з клієнтами, потребує особливої уваги до питань безпеки. Сучасні виклики у сфері захисту даних вимагають впровадження інноваційних методів та комплексних рішень для забезпечення конфіденційності інформації.

Організації, що використовують Salesforce, стикаються з низкою серйозних викликів у сфері захисту даних. Зростаюча складність кібератак, використання передових технологій зловмисниками та постійна еволюція методів соціальної інженерії створюють серйозні загрози для конфіденційної інформації. Особливу небезпеку становлять цільові атаки, спрямовані на отримання доступу до критично важливих даних клієнтів та бізнес-процесів. Додатковим викликом є необхідність забезпечення відповідності численним регуляторним вимогам та стандартам захисту даних, таким як GDPR та інші локальні нормативні акти.

Salesforce постійно впроваджує передові технологічні рішення для захисту конфіденційної інформації. Shield Platform Encryption забезпечує шифрування даних під час передачі, використовуючи найсучасніші криптографічні алгоритми (Рис.1.1) [1,2]. Event Monitoring дозволяє відстежувати та аналізувати всі дії користувачів у системі в режимі реального часу, що допомагає виявляти підозрілу активність та потенційні загрози. На рис.1.2 представлено процес Shield Platform Encryption [1,2]. Впровадження біометричної автентифікації та вдосконалених методів багатофакторної автентифікації значно підвищує рівень захисту від несанкціонованого доступу.

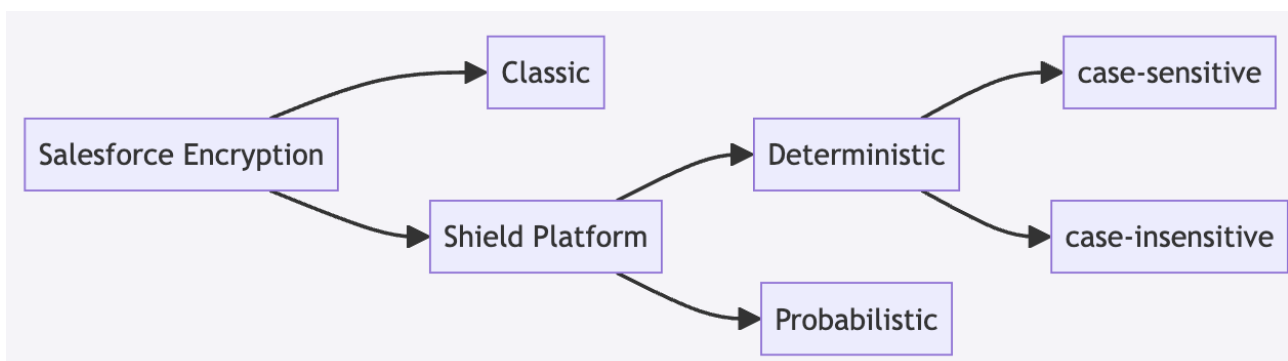


Рис.1.1. Тип шифрування Salesforce [1].

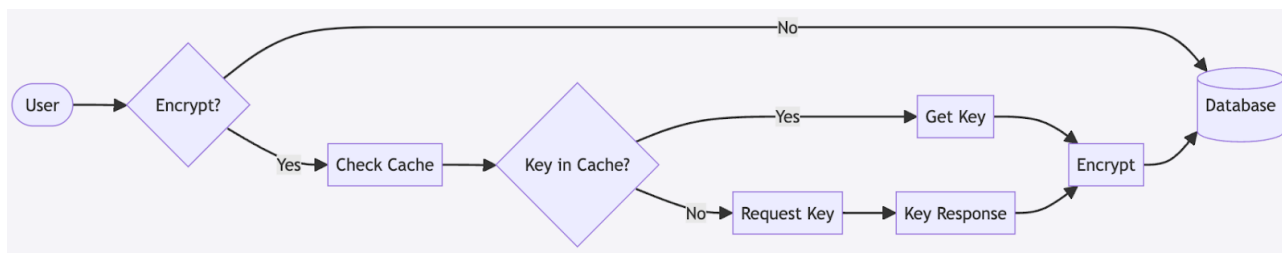


Рис.1.2. Процес Shield Platform Encryption [2].

Інтелектуальні системи здатні аналізувати величезні обсяги даних про активність користувачів, виявляти аномалії та передбачати потенційні загрози. Einstein Security Insights, інтегрований у платформу Salesforce, використовує передові алгоритми для автоматичного виявлення та запобігання підозрілим діям, що можуть загрожувати безпеці даних.

Впровадження концепції «Security by Design» забезпечує врахування вимог безпеки на всіх етапах розробки та впровадження рішень на платформі Salesforce [3]. Створення центрів компетенції з безпеки, регулярні тренінги персоналу та розробка детальних політик безпеки допомагають формувати культуру відповідального ставлення до захисту інформації.

Впровадження автоматизованих систем моніторингу, реагування на інциденти та управління доступом дозволяє значно знизити ризики людського фактору та підвищити швидкість реакції на потенційні загрози. Security Center надає централізований контроль над усіма аспектами безпеки, автоматизуючи рутинні операції та забезпечуючи комплексний огляд стану захищеності системи.

Використання технології Single Sign-On (SSO) у поєднанні з передовими методами верифікації особи забезпечує зручний та безпечний доступ до ресурсів системи. Важливим аспектом є також впровадження принципу найменших привілеїв та регулярний перегляд прав доступу.

Особлива увага приділяється захисту конфіденційної інформації при інтеграції Salesforce з іншими системами та під час обміну даними з партнерами. Використання захищених API, шифрування даних при передачі та впровадження суворих політик контролю доступу до інтеграційних інтерфейсів забезпечують безпеку інформації за межами основної системи. Важливим аспектом є також моніторинг та аудит усіх операцій з обміну даними.

Майбутнє захисту конфіденційної інформації в екосистемі Salesforce пов'язане з подальшим розвитком технологій квантової криптографії, вдосконаленням методів біометричної автентифікації та розширенням можливостей штучного інтелекту в сфері безпеки. Особлива увага приділятиметься розробці рішень для захисту даних в умовах розподілених обчислень та використання хмарних технологій.

Отже, захист конфіденційної інформації в екосистемі Salesforce вимагає комплексного підходу, що поєднує інноваційні технологічні рішення, передові методи організації процесів безпеки та постійне вдосконалення існуючих механізмів захисту. Успішне впровадження описаних методів та рішень дозволяє організаціям

ефективно протистояти сучасним кіберзагрозам та забезпечувати надійний захист конфіденційної інформації.

Список використаної літератури

1. Anderson, J., & Thompson, E. (2024). «Next-Generation Security Measures in Salesforce: From Classic to Shield Platform Encryption.» *Information Security Management Handbook*, Vol. 7, 412-436.
2. Chen, Y., Williams, P., & Kumar, R. (2023). «Modern Approaches to Data Protection in CRM Systems: Focus on Salesforce Security Features.» *Cloud Security Journal*, 8(2), 178-195.
3. Gupta, A., & Dhillon, G. (2024). «Enterprise Security Architecture for Cloud Platforms: A Salesforce Implementation Guide.» *Journal of Information Security and Applications*, 75(1), 103-121.
4. Richardson, M., & Smith, K. (2023). «Innovation in Cloud Security: Salesforce Platform Security Controls and Compliance.» *International Journal of Cloud Computing and Services Science*, 12(3), 245-267.
5. Петренко, О.В., Ковальчук, Л.М. (2023). «Сучасні методи захисту корпоративних даних у хмарних CRM-системах.» *Кібербезпека: освіта, наука, техніка*, 5(2), 89-102.

ВИКОРИСТАННЯ ТЕХНОЛОГІЇ ZERO TRUST ДЛЯ ПІДВИЩЕННЯ КІБЕРБЕЗПЕКИ БІЗНЕС-ПРОЦЕСІВ В УМОВАХ ЦИФРОВИХ ЗАГРОЗ

*Терещук Г.М.,
асистент кафедри інформаційних систем та технологій,
Київський національний університет імені Тараса Шевченка
Скляренко П.А.,
аспірант,
ПВНЗ «Європейський університет»*

У сучасних умовах цифрової трансформації підприємств, коли інформаційні ресурси виступають ключовим активом бізнесу, забезпечення їхнього захисту є критично важливим завданням [3]. Традиційні моделі безпеки, засновані на довірі до внутрішньої мережі, втрачають свою ефективність через динамічну зміну робочого середовища, розвиток віддаленого доступу, використання хмарних сервісів та зростання кількості кіберзагроз. У зв'язку з цим усе більшої популярності набуває концепція Zero Trust — стратегія, що базується на принципі «нікому не довіряй, завжди перевіряй» [1].

Концепція Zero Trust передбачає, що жоден користувач або пристрій не вважається автоматично довіреним, навіть якщо він знаходиться всередині корпоративної мережі. В основі підходу лежать механізми багатофакторної автентифікації, мінімізації прав доступу (Least Privilege), контекстного контролю, мікросегментації мережі та безперервного моніторингу активності. Застосування цих елементів дає змогу зменшити ризики несанкціонованого доступу, витоку даних і саботажу критичних бізнес-функцій [2].

Для бізнес-процесів важливою перевагою Zero Trust є гнучке управління доступом до ресурсів відповідно до ролей співробітників, поточного місця перебування, пристрою, з якого здійснюється вхід, та інших параметрів. Це дозволяє

знизити вразливість системи до атак з боку скомпрометованих акаунтів або внутрішніх користувачів.

Застосування Zero Trust у корпоративному середовищі також передбачає використання Software-Defined Perimeter (SDP) — віртуального периметру, який динамічно обмежує доступ до цифрових активів на основі попередньо визначених політик. У поєднанні з інструментами Endpoint Detection and Response (EDR) та Extended Detection and Response (XDR) це дає змогу ефективно відслідковувати аномальні дії в системі, оперативно реагувати на інциденти та підвищувати загальну кіберстійкість організації.

Одним із ключових чинників ефективності Zero Trust у сфері бізнесу є інтеграція з аналітикою на основі штучного інтелекту та машинного навчання. Такі системи здатні ідентифікувати атипову поведінку користувачів, прогнозувати потенційні загрози на основі історичних даних і автоматично застосовувати заходи реагування.

Суттєвою перевагою концепції Zero Trust у корпоративному секторі є її адаптивність до мультихмарних середовищ, гібридної інфраструктури, а також можливість поетапної реалізації без необхідності повного демонтажу наявних систем безпеки. Водночас серед викликів впровадження слід відзначити складність інтеграції в застарілі системи, високі витрати на впровадження та обслуговування, а також необхідність навчання персоналу основам нової моделі безпеки.

Інституційна підтримка впровадження Zero Trust на рівні корпоративного управління передбачає розробку детальних журналів подій, аудитів безпеки, регулярних тестувань політик доступу та ознайомлення працівників з оновленими стандартами кібергігієни.

Для ефективного впровадження Zero Trust необхідно поетапно здійснювати модернізацію існуючих інформаційних систем, забезпечуючи їхню інтеграцію з новими підходами. Особливу увагу слід приділяти розмежуванню доступу на основі ролей та контексту використання даних. Важливо також створювати детальні журнали подій для аналізу та аудиту можливих порушень політик безпеки. Інтеграція системи штучного інтелекту та машинного навчання допомагає в автоматичному виявленні аномалій, що може бути використано для миттєвого реагування на кіберзагрози.

Zero Trust є не просто набором технологій, а комплексною стратегією, що дозволяє значно підвищити рівень корпоративної безпеки. Завдяки її впровадженню можливо суттєво знизити ризики витоку конфіденційних даних, покращити контроль доступу до критично важливих систем та мінімізувати вплив потенційних атак. У сучасних умовах, коли кібератаки стають все більш складними та масованими, саме Zero Trust є тією парадигмою, яка допоможе забезпечити надійний захист від зовнішніх та внутрішніх загроз.

Ключовим компонентом моделі Zero Trust є її здатність адаптуватися до сучасних викликів кіберзахисту. В умовах гібридної війни критично важливим є захист мережевих систем від кібершпигунства, саботажу та атак на критичну інфраструктуру. Zero Trust дозволяє виявляти та блокувати спроби проникнення на ранніх етапах, забезпечуючи тим самим високий рівень стійкості до атак супротивника. Крім того, важливою складовою успішного впровадження Zero Trust є

його інтеграція з існуючими системами захисту інформації. Завдяки концепції Zero Trust ці системи можуть бути оптимізовані та доповнені новими механізмами безпеки, що дозволяє значно підвищити загальний рівень захисту.

Отже, концепція Zero Trust пропонує ефективну відповідь на зростаючі виклики цифрової безпеки в бізнес-середовищі. Вона забезпечує контрольований, адаптивний і надійний доступ до цифрових активів, дозволяє мінімізувати вразливості в інформаційних системах підприємств, підвищує рівень захисту критичних бізнес-процесів та створює основу для безпечної цифрової трансформації. Урахування принципів Zero Trust є необхідною умовою формування стійкої інформаційної архітектури сучасних компаній.

Список використаної літератури

1. Ворохоб М. В., Киричок Р. В., Яскевич В. О., Добришин Ю. Є., Сидоренко С. М. Сучасні перспективи застосування концепції Zero Trust при побудові політики інформаційної безпеки підприємства. // *Кібербезпека: освіта, наука, техніка*. – 2023. – №1(21). – С. 223–233.
2. Ковальчук А. М., Лисенко І. В., Григоренко О. С. Концепція нульової довіри: сучасні методи забезпечення кібербезпеки корпоративних мереж. // *Вісник Львівського державного університету безпеки життєдіяльності*. – 2024. – №29. – С. 45–5
3. Фролов І., Колодінська Я. Менеджмент кібербезпеки ІТ-продуктів. Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка». 2024. Т. 3, № 23. С. 310–317.

ШТУЧНИЙ ІНТЕЛЕКТ ПРОТИ ШТУЧНОГО ІНТЕЛЕКТУ: ПОТЕНЦІАЛ І РИЗИКИ ВИКОРИСТАННЯ ГЕНЕРАТИВНИХ МОДЕЛЕЙ У КІБЕРАТАКАХ ТА СИСТЕМАХ ЗАХИСТУ

Тритиниченко О.О
магістрант,
ПВНЗ «Європейський університет»

Сучасні досягнення в галузі штучного інтелекту (ШІ), зокрема генеративні моделі (LLM, GAN тощо), значно впливають на сферу кібербезпеки. З одного боку, вони відкривають нові можливості для автоматизації виявлення загроз, аналізу вразливостей та моніторингу систем у режимі реального часу. З іншого – самі ці технології стають інструментом у руках зловмисників.

У 2023–2024 роках були зафіксовані численні випадки використання генеративних моделей для створення шкідливого програмного коду, фішингових повідомлень та deepfake-матеріалів. Прикладом є поява моделей WormGPT та FraudGPT, які дозволяють формувати фішингові листи, обходити захист CAPTCHA, а також генерувати шкідливі сценарії для соціальної інженерії.

Водночас моделі, подібні до ChatGPT, використовуються з метою захисту: для аналізу логів, побудови сценаріїв реагування на інциденти та симуляції потенційних векторів атак. Вони також дозволяють формувати ХАІ (пояснюваний ШІ) підходи у виявленні аномалій, що критично важливо для систем фінансового моніторингу.

Особливої уваги заслуговують кейси з використанням deepfake для

компрометації бізнес-процесів. У 2024 році один із банків у Гонконзі втратив понад 25 млн доларів США після того, як deepfake-відео директора переконало співробітника переказати кошти на шахрайські рахунки. Такі атаки демонструють зростаючу загрозу для корпоративної безпеки.

Окрім цього, важливо відзначити, що використання генеративних моделей у кібербезпеці супроводжується проблемами інтерпретації та прозорості. Пояснюваність моделей (Explainable AI, XAI) набуває ключового значення в контексті довіри до автоматизованих систем прийняття рішень. У випадках, коли моделі прогнозують загрозу або збої у роботі, фахівці мають розуміти логіку таких прогнозів.

Сучасні дослідження зосереджені на застосуванні методів SHAP та LIME для аналізу рішень глибоких моделей. Наприклад, SHAP дозволяє зрозуміти, які саме параметри мережевого трафіку вплинули на класифікацію пакета як потенційно шкідливого. Це особливо важливо при моніторингу фінансових операцій або кіберзагроз у державних інформаційних системах.

Не менш актуальним є питання підготовки кадрів у сфері використання ШІ в кіберзахисті. В українських університетах поступово інтегруються дисципліни, що охоплюють питання безпечної розробки, етичного використання ШІ та протидії його зловживанню. Підготовка таких фахівців є запорукою формування сталої кібербезпеки у найближчому майбутньому.

Таким чином, генеративні моделі стають як інструментом загроз, так і засобом протидії у сфері кібербезпеки. Їхнє етичне, контрольоване та пояснюване використання — ключ до побудови надійного кіберзахисту. Україні, як державі, що знаходиться в епіцентрі кібервійни, необхідно не лише впроваджувати інновації, але й формувати стратегії їхнього відповідального застосування.

Список використаної літератури

1. WormGPT and FraudGPT: What They Are and How They Work. SecurityWeek.
<https://www.securityweek.com/wormgpt-fraudgpt-how-generative-ai-is-being-used-in-cybercrime/>
2. OpenAI. GPT and Security Implications. <https://openai.com/blog/gpt-4-system-card>
3. McKinsey & Company. The future of AI in cybersecurity.
<https://www.mckinsey.com/business-functions/mckinsey-digital/our-insights/the-role-of-ai-in-cybersecurity>

у

ПЕРСПЕКТИВИ ТРАНСФОРМЕР-МОДЕЛЕЙ В АНАЛІЗІ ФОНДОВОГО РИНКУ

*Троян К.М.
аспірант,*

ПВНЗ «Європейський університет», м.Київ

Трансформер-моделі на основі механізму уваги швидко розвинулися від задач обробки природної мови до кількісних фінансів, завдяки здатності моделювати довгострокові, нестационарні залежності в цінових рядах. Ранні емпіричні

дослідження показали, що навіть проста трансформер-модель, використовуючи лише енкодер блок, перевершила мережі з довгостроковою і короткостроковою пам'яттю (long short term memory, LSTM) приблизно на п'ять відсоткових пунктів у точності прогнозування 15-хвилинних рухів S&P 500, що свідчить про те, що механізм уваги може вловлювати перехресні закономірності, які часто не здатні змодельовати рекурентні комірочки [2]. У подальших дослідженнях були запропоновані такі варіанти архітектури моделі, як Informer, Autoformer і Temporal Fusion Transformer (TFT), які замінюють квадратичну увагу на розріджені або частотні механізми, що дозволяють обробляти тисячі часових інтервалів одночасно. Стефанюк і Слепачук [4] застосували Informer до п'ятихвилинних прибутків Bitcoin і виявили, що Informer, навчений на узагальненому середньому абсолютному збитку (generalised mean absolute directional loss, GMADL), показав коефіцієнт Шарпа на 0,7 вище, ніж buy-and-hold, і перевершив дві стратегії, засновані на технічних індикаторах, у чотирьох з шести тестових періодів, тоді як ідентична мережа, навчена на середній квадратичній похибці, деградувала на більш частотних даних [4]. Для акцій модель TFT, доповнена багатовимірними технічними індикаторами та контролем мультиколінеарності, забезпечила нижчу середньоквадратичну похибку, ніж LSTM, GRU або згорнуті базові лінії для семи американських акцій, підтверджуючи, що вбудовування статичних коваріатів і багатогоризонтних виходів допомагає моделі адаптуватися до мінливих режимів [1].

Окрім точкових прогнозів, карти уваги та гейтинг на рівні ознак забезпечують обмежену інтерпретацію: Стефанюк і Слепачук помітили, що GMADL-Informer концентрує увагу на періодах з великим обсягом торгів відразу після відкриття американського ринку, що відповідає зв'язку між волатильністю та обсягом ліквідності; Джайн та ін. заявили, що їхня модель TFT підкреслює ознаки осцилятора імпульсу під час трендових режимів, водночас зменшуючи їхню вагу у фазах повернення до середніх значень. Тим не менш, встановлення причинно-наслідкових зв'язків залишається складним завданням, оскільки механізм уваги не тотожний важливості. Ще однією практичною проблемою є перенавчання: Пропускна здатність трансформер-моделі надлінійно залежить від величини вхідних параметрів, тому дослідники часто застосовують метод відсіювання або відбір ознак на основі взаємної інформації. Пейк та ін. використовували взаємну інформацію для відсіювання надлишкових показників перед навчанням TFT на Bitcoin, Ethereum і Litecoin; їхній скорочений набір ознак прискорив роботу моделі на 30%, покращивши оцінку F1 на два відсоткові пункти [3].

Незважаючи на сильні результати зворотного тестування, залишаються дві прогалини. По-перше, більшість досліджень були відкалібровані для періодів бичачого ринку; стійкість до падінь, таких як у березні 2020 року або криптовалютної зими 2022 року, все ще недостатньо вивчена. По-друге, лише кілька робіт інтегрують комісії та витрати на виконання запиту в навчальну мету. Нещодавній огляд п'ятдесяти двох статей, присвячених часовим рядам Transformer, показав, що 79% перевершили базові показники RNN за точністю прогнозування поза вибіркою і 88% перевершили індикаторні стратегії на основі правил, але лише в дев'яти з них було враховано сліпедж або комісійні [5]. Подальші дослідження мають поєднати врахування втрат з діагностикою перемикання режимів і дослідити

перехресний трансфер активів - попередні дані свідчать про те, що теорії, засвоєні на Bitcoin, не поширюються на альткоїни або акції, вказуючи на приховані мікроструктурні відмінності, які ще належить змодельовати. Отже, Transformer-моделі стали вершиною фінансового прогнозування; у поєднанні з продуманими функціями втрат і ретельно підібраними технічними індикаторами вони послідовно перетворюють прогнозовані прибутки в економічно значущу перевагу, але питання інтерпретативності, режимної стійкості і реального виконання залишаються відкритими дослідницькими проблемами.

Список використаної літератури.

1. Jain, P., Gupta, S., & Sharma, R. (2024). Temporal fusion transformers for enhanced multivariate time-series stock prediction. *International Journal of Advanced Computer Science and Applications*, 15(7), 148–156. <http://dx.doi.org/10.14569/IJACSA.2024.0150713>
2. Nguyen, T., & Kim, J. (2024). Predictive modeling of stock prices using the Transformer model. *Proceedings of the 13th International Conference on Big Data Analytics* (pp. 112–121). <https://doi.org/10.1145/3674029.3674037>
3. Peik, A., Zare Chahooki, M. A., Milani Fard, A., & Agha Sarram, M. (2024). Leveraging time-series categorisation and temporal fusion transformers to improve cryptocurrency price forecasting (arXiv Preprint No. 2412.14529). <https://doi.org/10.48550/arXiv.2412.14529>
4. Stefaniuk, F., & Ślepaczuk, R. (2025). Informer in algorithmic investment strategies on high-frequency Bitcoin data (arXiv Preprint No. 2503.18096). <https://doi.org/10.48550/arXiv.2503.18096>
5. Zhang, L., & Li, H. (2024). Reviews on Transformer-based models for financial time-series forecasting. *Asia-Pacific Conference on Economics*, 56–67. <https://doi.org/10.54254/2755-2721/96/20241351>

ПИТАННЯ КІБЕРБЕЗПЕКИ ПРИ ВПРОВАДЖЕННІ ЦИФРОВИХ РІШЕНЬ У БІЗНЕС-ПРОЦЕСИ В УМОВАХ ВОЄННОГО СТАНУ

**Федік О.І.,
аспірант,
Скляренко О.В.
к.ф.-м.н., доцент,
завідувач кафедри математичних дисциплін та
інноваційного проектування
ПВНЗ «Європейський університет»**

Повномасштабне вторгнення росії на територію України є одним з найбільших військових конфліктів останніх десятиліть, а, отже, і виклики з якими стикаються господарюючі суб'єкти наразі мають зовсім інший масштаб. Востаннє з викликами загальнонаціонального масштабу Україна, як і більшість світу стикались в період пандемії COVID-19. Попри значні економічні збитки дана світова криза сприяла і пришвидшенню темпів цифрових рішень в компаніях світу. Компанії, які раніше лише будували плани щодо впровадження окремих цифрових рішень, аби втримати

свою операційну діяльність та конкурентоспроможність, були змушені реалізовувати цілі комплекси заходів з цифровізації.

Консалтингова компанія McKinsey в 2020 році провела McKinsey Global Survey опитавши 899 респондентів з-поміж керівників С-рівня та старших менеджерів з різних регіонів, галузей та за розміром компаній. В ході опитування більшість керівників зазначали, що така потужна криза, як пандемія COVID-19 змусила їх у пришвидшеному темпі впроваджувати трансформаційні зміни, в тому числі і з точки зору цифровізації. Серед основних кроків було [5]:

- зростання частки співробітників, що працюють віддалено і забезпечення їх технікою;
- збільшення використання передових технологій в операціях;
- збільшення використання передових технологій в прийнятті управлінських рішень;
- пришвидшення міграції даних у хмарні сховища;
- збільшення інвестицій в безпеку даних.

Як бачимо, пандемія COVID-19 окрім значних негативних наслідків для економіки та бізнесу все ж сприяла певним трансформаційним змінам в галузі цифровізації. В свою чергу, з початком повномасштабного вторгнення, впровадження цифрових рішень у бізнес-процеси було логічною реакцією господарюючих суб'єктів на порушені війною економічні, соціальні та ін. зв'язки в країні.

Проте чи можна даний досвід певним чином узагальнити до моделі або ж підходу з точки зору кібербезпеки при впровадженні цифрових рішень в умовах воєнного стану?

Перш за все, варто оцінити поточний стан кібербезпеки в країні та її компоненти на загальнодержавному рівні. Так, відповідно до основних показників стратегічної кібербезпеки за National Security Index станом на кінець 2023 року відзначався досить високий рівень кібербезпеки [1]. Лише показник «Освіта та професійний розвиток» має суттєво нижче значення відносно інших, що на думку дослідників свідчить про необхідність інвестування у підготовку фахівців та навчальні програми.

Наявність кваліфікованих кадрів з кібербезпеки на рівні підприємства або підготовка відповідних фахівців у закладах вищої освіти (ЗВО) є безумовно важливим фактором її забезпечення, проте, на нашу думку, загальний рівень цифрової грамотності суспільства теж не можна ігнорувати. Адже, так чи інакше будь-який співробітник компанії може стати джерелом ризиків для кібербезпеки. Саме тому Л. Арсенович та інші зазначають, що процес ознайомлення з основами кібербезпеки має відбуватися ще на рівні освітнього процесу в ЗВО [4].

Boston Consulting Group на основі дослідження підходів до трансформаційних процесів в компаніях в період пандемії сформувавши загальний підхід до побудови стійкості бізнесу в умовах кризи, що складається з трьох етапів і 6 аспектах стійкості [3]. Серед зазначених у підході BCG аспектів є і посилення кібербезпеки.

На відміну від багатьох інших аспектів даного підходу, вимір кібербезпеки може розглядатися на всіх етапах Відповідь/Реагування, Відновлення, Переосмислення з точки зору реалій умов воєнного стану. В той же час, оскільки

даний підхід було розроблено на основі дій та рекомендацій в період пандемії, він потребує адаптації до умов та реалій воєнного стану, що відображено в таблиці 1.

Таблиця 1

Адаптація підходу BCG з забезпечення стійкості підприємства шляхом цифрової трансформації в період COVID-19 до умов воєнного стану за критерієм кібербезпеки

	Реагування	Відновлення	Переосмислення
	<i>Забезпечення безпеки та безперервності бізнесу</i>	<i>Підготовка до перезапуску та відновлення</i>	<i>Формування бачення конкурентної переваги в новій реальності</i>
Посилення кібербезпеки	Захист локальних операцій при дистанційному доступі	Масштабування систем кібербезпеки та їх розширення	Кібербезпека є одним із пріоритетів керівників С-рівня. Сформовано бачення «центру» безпекових операцій. Передбачено безпекові заходи для ланцюгів постачання та e-commerce платформу, в тому числі і партнерські.
<i>Адаптація</i>		+ Посилення протоколів безпеки та тестування на предмет виявлення вразливостей. + Проведення термінових освітніх заходів для персоналу компанії з питань кіберзахисту	+ Створення довгострокової програми освітніх заходів для персоналу компанії з питань кіберзахисту. + Постійний аналіз потенційних ризиків кібербезпеки для кращого розуміння потенційних бізнес-загроз. + Оцінка підходів до дублювання та зберігання критично важливих даних в хмарних сховищах і т.п. + Планування систем контролю доступу та сповіщень про несанкціоновані дії. + Розробка планів розосередження об'єктів критичної ІТ-інфраструктури підприємства.

Джерело: сформовано авторами на основі [2,3].

Звісно, даний підхід є узагальненим і впровадження цифрових рішень в різних галузях матиме свої особливості [6].

З урахуванням вищевикладеного можна зробити висновок, що впровадження цифрових рішень доцільно здійснювати на основі моделей та підходів, сформованих у періоди попередніх кризових явищ, наприклад, пандемії COVID-19. Водночас, такі рішення потребують глибокої адаптації до специфіки сучасних викликів, обумовлених умовами воєнного стану. Окрім того, реалізація цифрових трансформацій невід’ємно пов’язана з необхідністю забезпечення належного рівня кібербезпеки та підвищення стійкості інформаційно-комунікаційної інфраструктури.

Список використаної літератури

1. Данілейко Л.Ю., Романюк І.І. Оцінка рівня кібербезпеки України в умовах війни. *ResearchGate*. Retrieved from https://www.researchgate.net/publication/388310208_OCINKA_RIVNA_KIBERBEZPEKI_UKRAINI_V_UMOVAN_VIJNI (Accessed on April 16, 2025)
2. Кузьменко О.Ю., Маклюк О.В., Чернишова О.О. Кібербезпека бізнесу під час війни. *Economy and Society*, Issue 44, 2022. DOI: <https://doi.org/10.32782/2524-0072/2022-44-21>. Retrieved from <https://economyandsociety.in.ua/index.php/journal/article/view/1790/1725> (Accessed on April 16, 2025)
3. Close K., Grebe M., Andersen P., Khurana V., Franke M.R., Kalthof R. Digital path to business resilience. *BCG*, July 06, 2020. Retrieved from <https://www.bcg.com/publications/2020/digital-path-to-business-resilience> (Accessed on April 16, 2025)
4. Leonid Arsenovych, Oleksandr Nikolaievsky, Olena Skliarenko, Leonid Lytvynenko, Ivan Kydriavskyi. Organization of Training with the Use of Digital Technologies for Ensuring Cybersecurity in the Educational Space // *WSEAS Transactions on Computer Research*, ISSN / E-ISSN: 1991-8755 / 2415-1521, Volume 12, 2024, p. 524-536, DOI: 10.37394/232018.2024.12.51 , <https://wseas.com/journals/articles.php?id=9988>
5. O'Toole C., Schneider J., Smaje K., LaBerge L. How COVID-19 has pushed companies over the technology tipping point—and transformed business forever. *McKinsey & Company*, October 5, 2020. Retrieved from <https://www.mckinsey.com/capabilities/strategy-and-corporate-finance/our-insights/how-covid-19-has-pushed-companies-over-the-technology-tipping-point-and-transformed-business-forever> (Accessed on April 16, 2025)
6. Фролов І., Колодінська Я. Менеджмент кібербезпеки ІТ-продуктів. Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка». 2024. Т. 3, № 23. С. 310–317.

КІБЕРЗАХИСТ ОСВІТНІХ СТАРТАПІВ: ВИКЛИКИ ЦИФРОВОЇ БЕЗПЕКИ ТА ІННОВАЦІЙНІ РІШЕННЯ

*Фролов І.М.
викладач кафедри математичних дисциплін
та інноваційного проектування,
Колодінська Я.О.,
старший викладач кафедри математичних дисциплін
та інноваційного проектування
ПВНЗ «Європейський університет»*

Глобалізаційні процеси, пандемія Covid-19, війна в Україні та зміни на ринку праці актуалізували потребу в трансформації освітніх підходів і впровадженні інноваційних рішень. В умовах цифровізації навчального процесу особливої ваги набуває питання кібербезпеки, оскільки освітні IT-стартапи стають вразливими до цифрових загроз через технічну недосконалість і обмежені ресурси. Актуальним напрямом стає аналіз викликів у сфері цифрової безпеки освітніх (EdTech) стартапів і пошук інноваційних механізмів їх захисту в сучасному інформаційному середовищі.

Аналізуючи визначення поняття стартапу в науковій літературі, можна визначити освітній стартап як інноваційний проєкт, ініційований учнями, студентами або викладачами в межах навчального або наукового процесу, що реалізується з метою розробки новітнього освітнього продукту, сервісу чи технології. Такий стартап проходить послідовні етапи — від генерації ідеї до апробації, верифікації та промоції, орієнтуючись на досягнення практичної користі (прибутку, соціального ефекту чи освітньої цінності). Йому притаманні обмежені ресурси, чіткий дедлайн реалізації (від одного місяця до кількох років), а також спрямованість на інновації, масштабованість і потенціал зростання [5].

Сучасний розвиток ринку освітніх стартапів зумовлений інтенсивною еволюцією інформаційно-комунікаційних технологій та загальними процесами цифровізації суспільства, що, у свою чергу, спричиняє зростання попиту на технологічні рішення, спрямовані на впровадження інноваційних форматів освітнього процесу. Серед потенційних напрямів розвитку стартапів у галузі освіти можна виділити наступні [1, 4, 5]:

- технології мобільного навчання;
- технології e-learning;
- корпоративна освіта;
- learning management systems (системи управління навчанням);
- гейміфікація освітнього процесу;
- впровадження VR і AR технологій для доповнення навчального процесу;
- впровадження технологій штучного інтелекту;
- робототехніка.

Перехід до онлайн-навчання в умовах війни дозволив забезпечити безпеку учасників освітнього процесу, але водночас виявив гостру потребу у забезпеченні кіберзахисту освітнього середовища. Освітні заклади, що активно використовують

цифрові платформи, стали привабливою мішенню для кіберзлочинців через зростання вразливостей інформаційної безпеки [3]. У сучасних умовах інформаційно-комп'ютерного протистояння, що включає кібератаки та кампанії дезінформації, кібербезпека набуває критичного значення. Для освітніх стартапів, які функціонують у цифровому середовищі, забезпечення захисту інформації стає необхідною умовою стабільного розвитку та протидії загрозам гібридної війни.

Аналіз наукової літератури дозволив виділити наступні види кіберзагроз для освітніх стартапів [2, 3, 6]:

- загрози технічного характеру, а саме, фішингові атаки (отримання особистих даних через підроблені ресурси), шкідливе програмне забезпечення (віруси, трояни, програми-вимагачі, бекдори), DDoS-атаки (перевантаження системи фальшивим трафіком з метою недоступності платформи), автоматизовані боти (сканування вразливостей та впровадження шкідливого коду), злом баз даних (несанкціонований доступ до внутрішньої інформації), застаріла IT-інфраструктура (ризики через відсутність оновлень та патчів безпеки);

- загрози, пов'язані з людським фактором, серед яких низька цифрова та кіберграмотність користувачів, помилкові дії або халатність персоналу, інсайдерські загрози (ненавмисне або навмисне зловживання доступом), соціальна інженерія (маніпуляції для отримання конфіденційної інформації), дії учнів чи студентів, що призводять до зараження або пошкодження систем;

- загрози для інформаційної безпеки, такі як витік персональних даних (прізвища, адреси, банківські дані, освітня історія), крадіжка інтелектуальної власності (коди, внутрішні розробки, навчальні ресурси), розголошення конфіденційної документації, несанкціонований доступ до електронного документообігу;

- загрози, пов'язані з використанням штучного інтелекту, зокрема, збір та обробка великої кількості особистих даних без належного захисту, ризик упереджених алгоритмів (bias) у прийнятті рішень, можливість кібератак через недосконалість моделей ШІ, етичні порушення та втручання у приватне цифрове середовище учнів;

- інформаційно-соціальні загрози, тобто, маніпуляція свідомістю через ботів та фейкові профілі у соцмережах, таргетована пропаганда та дезінформація, психологічний вплив на студентів через алгоритми рекомендацій, поширення небажаного контенту або інформаційний тиск на користувачів.

Забезпечення кібербезпеки та сталого розвитку вітчизняних освітніх стартапів у сучасному інформаційному середовищі вимагає впровадження комплексних інноваційних рішень, орієнтованих на адаптацію до новітніх технологічних викликів. Одним із ключових підходів є інтеграція концепції Zero Trust Security, яка базується на принципі повної недовіри до всіх запитів незалежно від їхнього джерела. Така модель передбачає обов'язкову автентифікацію, авторизацію та шифрування всіх запитів, що дозволяє суттєво зменшити ризики інсайдерських атак, несанкціонованого доступу та витоку конфіденційної інформації. Ефективним інструментом кіберзахисту також виступає застосування технологій штучного інтелекту для моніторингу поведінкових патернів користувачів, що дозволяє

виявляти аномальні дії в режимі реального часу та попереджати фішингові атаки, спроби вторгнення, а також інциденти соціальної інженерії.

Особливу увагу доцільно приділяти підвищенню рівня цифрової обізнаності учасників освітнього процесу через впровадження гейміфікованих програм кібергігієни — інтерактивних курсів і тренажерів, які дозволяють ефективно формувати безпечну поведінку в цифровому просторі серед викладачів, студентів та розробників. Не менш важливим напрямом є інтеграція багаторівневого шифрування для захисту персональних даних і освітнього контенту в середовищах хмарного зберігання та в базах даних. Регулярне проведення незалежного аудиту інформаційної безпеки, зокрема, із залученням автоматизованих платформ, сприяє своєчасному виявленню вразливостей та адаптації захисних механізмів до змін у кіберландшафті.

З урахуванням активного використання штучного інтелекту в EdTech-продуктах особливої актуальності набуває дотримання етичних фреймворків, які включають забезпечення прозорості алгоритмів, обмеження доступу до персональних даних та уникнення дискримінаційних рішень. Перспективним напрямом є також застосування технологій блокчейн для захисту результатів навчання, сертифікатів та особистих досягнень користувачів, що забезпечує їх незмінність, достовірність і підвищує рівень довіри до освітньої платформи. У зв'язку з активною взаємодією користувачів у соціальних компонентах освітніх платформ, важливо розробляти рішення для виявлення фейкових облікових записів, бот-активності та інформаційного впливу, який може мати психологічні наслідки.

На стратегічному рівні необхідним є інтегрування принципів кібербезпеки в бізнес-модель освітнього стартапу, що дозволяє не лише забезпечити стійкість продукту до кіберзагроз, а й підвищити інвестиційну привабливість, сформувати довіру з боку споживачів та дотримання вимог регуляторних органів. Такий підхід сприяє формуванню сталого, конкурентоспроможного та безпечного освітнього середовища в умовах цифрової трансформації.

Підсумовуючи, в умовах стрімкої цифровізації освітнього середовища та зростання кіберзагроз, питання забезпечення кібербезпеки освітніх стартапів набуває критичного значення. Враховуючи технічні вразливості, людський фактор та інтенсивне використання інноваційних технологій, зокрема, штучного інтелекту, необхідно впроваджувати комплексні, технологічно обґрунтовані й етично виважені заходи захисту. Запропоновані інноваційні рішення дозволяють не лише мінімізувати ризики, а й сформувати довіру до освітніх продуктів, забезпечити їх стабільність, сталий розвиток і конкурентоспроможність у сучасному інформаційному просторі.

Список використаної літератури

1. Пищик О. Інтеграція цифрових технологій у сучасній освіті: аналіз викликів та можливостей. *Актуальні проблеми в системі освіти: загальноосвітній заклад середньої освіти – доуніверситетська підготовка – заклад вищої освіти*. 2024. № 4. С. 368–377.
2. Фролов І., Колодінська Я. Менеджмент кібербезпеки ІТ-продуктів. *Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка»*. 2024. Т. 3, № 23. С. 310–317.

3. Зінченко О. Проблеми забезпечення кібербезпеки дистанційної освіти в Україні в умовах війни. *Інтеграція системи освіти України в європейський освітній простір : VI Міжнар. науково-практ. конф., м. Київ, 25 жовт. 2024 р. Київ, 2024. С. 172–173.*
4. Скляренко О.В., Ягодзінський С.М., Ніколаєвський О.Ю., Невзоров А.В. Цифрові інтерактивні технології навчання як невід’ємна складова сучасного освітнього процесу // *Науковий журнал «Інноваційна педагогіка», Випуск 68, Т.2. – 2024. - с. 250-257.*
5. Куркіна С., Журавльов В. Стартани в дистанційному навчанні: сучасний світовий досвід. *Наукові записки. Серія: Педагогічні науки. 2022. № 207. С. 180–184.*
6. Шевчук М. Сучасні виклики і загрози в сфері інформаційної безпеки держави. *Актуальні проблеми вітчизняної юриспруденції. 2024. № 6. С. 160–166.*

МЕТОДИ БЕЗПЕЧНОГО ЗБЕРІГАННЯ КОНФІДЕНЦІЙНИХ ДАНИХ CRM/ERP-СИСТЕМ

*Чорненко М.В.,
аспірант,
ПВНЗ «Європейський університет»*

CRM/ERP-системи нині є невід’ємною частиною бізнес-процесів абсолютної більшості підприємств, що пояснюється необхідністю автоматизації, оптимізації та контролю рутинної виробничої діяльності. Такі системи глибоко інтегровані в бізнес-процеси компанії і, як наслідок, часто містять багато конфіденційної інформації, витік якої може спричинити фінансові, репутаційні збитки підприємству та завдати шкоди його клієнтам і постачальникам. На тлі регулярних кібератак і витоків даних, які, на жаль, давно стали звичним явищем, нехтування цими ризиками потенційно може мати високу ціну, що підкреслює необхідність коректного зберігання та поводження з конфіденційними даними. Варто зазначити, що не існує єдиного універсального методу, здатного гарантувати абсолютну цифрову безпеку. Отже, надійний захист завжди має бути багаторівневим і комплексним, тобто охоплювати певний спектр методів, конфігурація якого залежить від потреб компанії, соціального проєкту та ін.

Наразі для комплексного захисту чутливої інформації в CRM/ERP-системах застосовують кілька методів[1-3], наведених далі.

Шифрування даних передбачає перетворення даних у недоступний для сторонніх формат. Основними видами є: шифрування даних на рівні бази даних, яке, зазвичай, застосовується для зберігання фінансової та клієнтської інформації; шифрування під час передачі даних, що використовується для захисту інформації, яка передається між клієнтським пристроєм і сервером; асиметричне шифрування, яке застосовується для захисту ключів автентифікації. Зазвичай, ці методи шифрування інтегровані в ядро CRM/ERP-системи і не потребують значних змін, окрім налаштувань відповідно до вимог користувача.

Токенізація даних являє собою метод обробки конфіденційної інформації, який передбачає заміну чутливих даних на унікальні ідентифікатори — токени, що не мають внутрішнього значення поза межами відповідної системи. Отримані токени можуть бути використані автентифікованим користувачем для доступу до первинних

даних, що зберігаються в окремому захищеному сховищі. Такий підхід значно знижує ризик компрометації конфіденційної інформації у разі несанкціонованого доступу до основної бази даних.

Контроль доступу та рольова модель обмежує доступ до даних компанії з урахуванням посадових обов'язків користувача та зберігає журнал його активності. Як результат, користувач має доступ виключно до даних, які необхідні йому для безпосередньої службової діяльності, що знижує ризик витоку. При цьому користувач генерує журнал своєї активності, що може бути використаний для подальшого аналізу на предмет підозрілої активності.

Багатофакторна автентифікація забезпечує додаткові рівні перевірки особи при входженні в систему, що зменшує ризики, пов'язані з витоком облікових даних, необхідних для доступу в систему.

Протидія соціальній інженерії включає тренінги з метою поширення знань про потенційні методи зламу і викрадення даних (фішингові розсилки, підозрілі сайти й посилання та ін.), а отже, протидії атакам через користувачів.

Моніторингові системи захисту використовуються з метою виявлення підозрілої активності, що може бути маркером зламу системи. Основними системами такого типу є: DLP (системи запобігання витоку даних), IDS/IPS (сфокусовані на виявленні спроб несанкціонованого доступу), AI/ML (спрямовані на автоматичне виявлення аномалій у поведінці користувача).

Отже, для збереження конфіденційності даних слід впроваджувати комплексну безпекову стратегію, що включатиме шифрування / токенизацію; контроль доступу; багатофакторну автентифікацію; аудит і моніторинг робочих процесів задля мінімізації асоційованих із соціальною інженерією ризиків та забезпечуватиме системний моніторинг кібератак і захист від них.

Список використаної літератури

1. Microsoft. *Dynamics 365 Security Whitepaper*. 2023. URL: <https://learn.microsoft.com> (дата звернення: 24.04.2025).
2. OWASP Foundation. *OWASP Top 10: The Ten Most Critical Web Application Security Risks*. URL: <https://owasp.org/www-project-top-ten/> (дата звернення: 24.04.2025).
3. Cloud Security Alliance (CSA). *Security Guidance for Critical Areas of Focus in Cloud Computing v4.0*. 2021. URL: <https://cloudsecurityalliance.org> (дата звернення: 24.04.2025).

ПОРІВНЯЛЬНИЙ АНАЛІЗ МІЖНАРОДНИХ СТАНДАРТІВ УПРАВЛІННЯ КІБЕРРИЗИКАМИ В КОРПОРАТИВНОМУ СЕКТОРІ

Шаповалов Б. Д.
студент кафедри безпеки інформаційних технологій
Коробейнікова Т. І.
кандидат технічних наук
доцент кафедри безпеки інформаційних технологій
Національний університет «Львівська політехніка»

Сучасні кіберзагрози вимагають від корпоративного сектору системного підходу до управління кіберризиками. Мета дослідження – провести порівняльний аналіз ключових міжнародних стандартів кібербезпеки (NIST CSF 2.0, ISO 27001/27002, SOC2, GDPR) для оцінки їхньої застосовності до задач корпоративного управління кіберризиками, підвищення стійкості організацій та відповідності законодавчим нормам.

NIST Cybersecurity Framework 2.0, розроблений Національним інститутом стандартів і технологій США, є гнучкою рамковою програмою для управління кіберризиками. Оновлена у 2024 році версія включає шість функцій: ідентифікація, захист, виявлення, реагування, відновлення та керування. Цей стандарт підходить для організацій будь-якого типу – від корпорацій до некомерційних структур – завдяки акценту на інтеграцію кібербезпеки в загальне управління ризиками [1]. NIST CSF 2.0 супроводжується практичними ресурсами, такими як посібники та приклади впровадження, що сприяє гармонізації підходів до кіберстійкості на глобальному рівні.

ISO 27001 та ISO 27002, створені Міжнародною організацією зі стандартизації, є визнаними стандартами для побудови системи управління інформаційною безпекою (СУІБ). ISO 27001 передбачає сертифікацію, що дозволяє організаціям демонструвати відповідність регуляторним вимогам перед партнерами та клієнтами. ISO 27002 доповнює його практичними рекомендаціями щодо впровадження засобів контролю [2–3]. Перевагою є структурований підхід до виявлення вразливостей, однак сертифікація вимагає значних ресурсів і часу, а безперервний моніторинг необхідний для врахування нових ризиків.

SOC2, розроблений Американським інститутом сертифікованих бухгалтерів (AICPA), орієнтований на аудит сторонніх постачальників і перевірку безпеки даних клієнтів. Стандарт включає понад 60 вимог відповідності та передбачає тривалий аудит (до року), результатом якого є звіт про стан кібербезпеки [4]. SOC2 особливо актуальний для фінансового сектору, де високі стандарти відповідності є критичними, але його складність може бути бар'єром для менших організацій.

Загальний регламент про захист даних (GDPR), прийнятий у 2016 році в ЄС, регулює обробку персональних даних громадян Євросоюзу. Він стосується як європейських компаній, так і тих, що працюють із даними резидентів ЄС. GDPR включає 99 статей, що охоплюють права споживачів, політики захисту даних та вимоги до повідомлення про витоки (72 години). Штрафи за порушення (до 20 млн євро або 4% глобального доходу) стимулюють компанії до впровадження суворих

заходів безпеки [5]. GDPR гармонізує кібербезпеку з міжнародними стандартами, але вимагає значних зусиль для дотримання.

Практичне застосування розглянутих стандартів дозволяє організаціям адаптувати їх до власних потреб. Наприклад, NIST CSF 2.0 може бути основою для швидкого реагування на інциденти, ISO 27001 – для побудови довгострокової стратегії кібербезпеки, SOC2 – для співпраці з постачальниками, а GDPR – для роботи на європейському ринку [6–7]. Такий підхід сприяє не лише захисту даних, а й підвищенню конкурентоспроможності компаній у глобальному середовищі.

Порівняльний аналіз показав, що NIST CSF 2.0 вирізняється адаптивністю та простотою інтеграції, ISO 27001/27002 – формалізацією процесів і міжнародною визнаністю, SOC2 – глибиною аудиту сторонніх сторін, а GDPR – акцентом на захисті даних і регуляторній відповідності. Вибір стандарту залежить від розміру організації, галузі, доступних ресурсів та законодавчого контексту. Комбіноване використання цих стандартів може оптимізувати корпоративне управління кіберризиками, підвищити стійкість до загроз і забезпечити виконання адміністративних та правових вимог. Це дає поштовх до розвитку ефективних стратегій кібербезпеки в умовах глобальних викликів.

Список використаної літератури

1. NIST. *Cybersecurity Framework 2.0*. [Gaithersburg], 2024. 55 p.
2. ISO. *ISO 27001: Information Security Management*. [Geneva], 2022. 32 p.
3. ISO. *ISO 27002: Code of Practice for Information Security Controls*. [Geneva], 2022. 48p.
4. AICPA. *SOC2 – Service Organization Control Type 2*. [Washington], 2023. 45 p.
5. European Union. *General Data Protection Regulation (GDPR)*. [Brussels], 2016. 88 p.
6. Korobeinikova T., Tachenko I., Chekhmestruk R., Mykhaylov P., Romanyuk O., Romanyuk S. *A General Method of Risk Estimation*. 2023 13th International Conference on Advanced Computer Information Technologies (ACIT) (Wroclaw, Poland, 2023), pp. 410–413.
7. Шаповалов Б. Д. *Управління ризиками: процес та стратегії реагування*. Матеріали Міжнародної науково-практичної конференції «Кібербезпека в сучасному світі: актуальні виклики» (2024), с. 267–270.

ЗАХИСТ ЦИФРОВОГО АКТИВУ В УМОВАХ ЦИФРОВІЗАЦІЇ ТА ВПРОВАДЖЕННЯ ІНСТРУМЕНТІВ ШТУЧНОГО ІНТЕЛЕКТУ В УПРАВЛІННЯ БІЗНЕС ПРОЦЕСАМИ

*Шевчук А.А.,
кандидат економічних наук, докторант
ПВНЗ«Європейський університет»*

В умовах цифрової трансформації вирішальне значення для підтримання конкурентоспроможності й стійкості бізнесу мають сучасні інструменти ІІІ. Сучасні наукові дослідження та публікації висвітлюють можливості та загрози пов'язані з інтеграцією штучного інтелекту в управління бізнес процесами ,що потребує

різних підходів до розв'язання проблем у цій сфері та подолання наявних викликів. Використання штучного інтелекту та інтегрування в усі сфери, що сприяє ефективності бізнесу. Зокрема доводять, що штучний інтелект це дієвий бізнес-інструмент, який використовує нейронні мережі та пришвидшує час на виконання завдань, формування інформації, прогноз показників та прийняття оптимальних рішень, що дасть змогу замінити цілі відділи і призведе до появи нової посади в компаніях -R&D-менеджер[1]. Важливо зазначити, що переваги від розвитку штучного інтелекту вагомі та будуть ще більшими, так як існує прямий зв'язок інформаційних технологій та розвитку бізнесу при ефективному управлінні на основі якісної інформації. В своїх дослідженнях Шевчук А.А. [2] зазначає, що в сучасних умовах інформація є інноваційним цифровим активом, тобто дані мають цінність для управління бізнес процесами та потребують надійної системи захисту, що сприятиме результативності бізнесу та розвитку цифрової економіки. Ефект від використання інструментів ШІ дає нові можливості для оптимізації управління бізнес-процесами та отримати значні конкурентні переваги на ринку. Особливо важливим є визначення ефекту від впровадження інструментів ШІ в управління бізнес-процесами в умовах формування цифрової економіки. Адже цифрова економіка стає відособленою системою, яку можна назвати ефективною системою управління, економічна сутність якої полягає в динамічному процесі перетворення та інтеграції існуючих методів обліку, аналізу, планування, контролю і координації в єдину систему отримання, обробки інформації для прийняття на її основі ефективних управлінських рішень для розвитку бізнесу та економіки країни в цілому. Важливо, що впровадження ШІ та переведення інформації в мережі змінює її структурне значення та стає цифровим активом, яка є основою управління та розвитку цифрової економіки на основі можливостей Інтернету. Розглянемо ключові показники ефективності (KPIs), які оцінюють впровадження інструментів ШІ в управління бізнес-процесами щодо оптимізації в умовах формування цифрової економіки(рис.1).



Рис.1. Показники ефективності (KPIs), які оцінюють результативність впровадження інструментів ШІ в управління бізнес-процесами

Для удосконалення організаційних і методичних положень функціонування системи управління бізнес процесами щодо їх оптимізації із застосуванням штучного інтелекту необхідно перевести інформацію в цифрову форму . Адже даний процес призведе до створення єдиної цифрової інфраструктури та стане основою переходу до цифрової економіки. На створення мережевої (підтримуючої) інфраструктури як основи формування цифрової економіки спрямовано дослідження науковців[1]. Розвиток цифрової платформи сприятиме використанню ресурсів її користувачів замість власних ресурсів для управління бізнес процесами, тобто використовується мережева модель створення цінності даних, які стають цифровим активом та замінять лінійну модель , що підвищить ефективність управлінських процесів в бізнес середовищі в умовах формування мережевої економіки. Підтверджено значний вплив цифрових трансформацій на управління бізнес процесами з використанням ШІ в умовах формування цифрової економіки та визначає проблеми, які пов'язані з використанням мережі[3].Тобто розвиток цифрової економіки тісно пов'язаний із застосуванням цифрових технологій, шляхом створення та обміну цифрових даних ,які мають цінність і є цифровим активом. Цифровий інноваційний актив становить інформаційне забезпечення управління є комерційною таємницею бізнесу, включає особисті дані ,що використовується для управління бізнес процесами .Важливо , що витік інформації може спричинити серйозні проблеми для бізнесу. Тобто втрата цифрового активу може призвести до отримання бізнесом додаткових ризиків, а не бізнес можливостей ,що потребує захисту від кібератак, які несуть ризики[4] . Управління бізнес-процесами здійснюється на основі цифрового активу, адже використанням ШІ в бізнесі виходить за рамки інфраструктури бізнесу через мобільні, хмарні та гібридні застосунки отримання та передачі даних про бізнес компанію підрядникам та клієнтам. Дана ситуація потребує захисту від кібератак цифрового активу, який набув ділової цінності для управління і вирішення бізнес проблем .Для кібербезпеки доречно в умовах цифрової трансформації використовувати штучний інтелект, який став важливим сучасним інструментом в контексті захисту цифрового активу[4].

Встановлено, що використання інтелектуального ШІ в захисті цифрового активу є практикою в умовах формування цифрової економіки. Адже прогнозуючий ШІ може оптимізувати виявлення загроз і створювати ефективні рішення з кібербезпеки, що тісно пов'язані з ландшафтом загроз, які мають тенденцію постійно змінюватися та створювати небезпеку для цифрового активу. Системи захисту ґрунтовані на інтелектуальному прогнозуючому ШІ можуть само контролюватися , самонавчатися, застосовувати аналіз у непередбачуваних ситуаціях та приймати рішення на основі власних спостережень, що є важливим в умовах формування цифрової економіки.

Список використаної літератури

1.Гевчук А.В, Шевчук А.А.(2023) Мережева (підтримуюча) інфраструктура і штучний інтелект в управлінні бізнес процесами - основа формування цифрової економіки. Цифрова економіка та економічна безпека . 8 (08), 207-212. URL:<https://doi.org/10.32782/dees.8-34>

2.Шевчук А.А. (2024)Інноваційний цифровий актив - основа управління бізнес процесами в умовах цифрової трансформації економіки. XLV Міжнародна науково-практична конференція «Сучасні аспекти модернізації науки: стан, проблеми, тенденції розвитку» (м. Олександрополіс (Греція) 07 червня 2024 року).

3.Шевчук А.А. (2024). Вплив цифрових трансформацій на управління бізнес процесами з використанням ШІ в умовах формування цифрової економіки: переваги та ризики. Інфраструктура ринку. В. 78/2024. URL: <http://www.market-infr.od.ua/uk/78-2024>.

4.Шевчук А.А. (2024). Використання ШІ для захисту цифрового активу щодо управління бізнес процесами в умовах цифрових трансформацій: еволюція та майбутнє. Економічний просвіт: № 191 С.244-249. DOI: <http://prostir.pdaba.dp.ua/index.php/journal/article/view/1564>.

СУЧАСНИЙ ЛАНДШАФТ КІБЕРЗАГРОЗ ТА МЕТОДИ ЗАХИСТУ

Шугалій Д.А.

аспірант,

Яровий Р.О.

к.т.н, доцент, доцент кафедри кібербезпеки

та захисту інформації

ПВНЗ «Європейський університет»

У сучасних умовах стрімкого розвитку цифрових технологій та глобальної цифровізації, кібербезпека стає одним із ключових компонентів національної безпеки, стійкості бізнесу та захисту критичної інфраструктури. Авторами окреслено основні виклики, з якими стикаються інформаційні системи, а також проаналізовано актуальні загрози, типи атак і відповідні методи захисту. Цифровізація державних і приватних сервісів, особливо у критично важливих сферах, таких як енергетика, зв'язок, транспорт і фінанси, значно підвищує їхню вразливість до кіберзагроз. Переважна більшість сучасних бізнес-процесів здійснюється у цифровому середовищі, що зумовлює необхідність переходу до принципово нових стратегій забезпечення інформаційної безпеки.

Статистичні дані, зокрема звіт Verizon DBIR 2024, свідчать про значне зростання кількості інцидентів — понад 30 000 випадків кіберзагроз у 2023 році, з яких понад 10 600 підтверджених витоків даних, що демонструє ескалацію як у кількісному, так і якісному аспекті атак. Геополітичний чинник відіграє дедалі важливішу роль у ландшафті загроз: кібератаки стали невід'ємним елементом гібридної війни. Зокрема, інцидент 12 грудня 2023 року, пов'язаний з атакою на «Київстар», призвів до відключення зв'язку та збоїв систем оповіщення, а також атаки на державні реєстри Міністерства юстиції у грудні 2024 року та онлайн-сервіси «Укрзалізниці» у березні 2025 року, які спричинили порушення роботи критичних систем. Ці приклади свідчать про високу ефективність атак на цифрову інфраструктуру в умовах конфлікту, що потребує перегляду концептуальних засад кіберзахисту.

Фінансові та репутаційні наслідки таких атак є значними. Наприклад, після атаки на «Київстар» компанія змушена була виділити понад 90 мільйонів доларів США на відновлення інфраструктури та посилення заходів безпеки. Окрім того, несанкціонований доступ до баз персональних даних державних установ становить серйозну загрозу приватності громадян, підвищує ризики фішингових кампаній, компрометації акаунтів і соціальної інженерії.

Кіберзагрози класифікуються за різними векторами: malware (у т.ч. ransomware, spyware, трояни, RAT), фішинг (включаючи spear phishing, smishing, vishing), DDoS-атаки, атаки на ланцюги постачання. За даними ENISA Threat Landscape 2024, атаки на доступність стали однією з головних загроз для фінансового сектору ЄС, а частка атак на supply chain зросла на 68 % у 2024 році. Особливої актуальності набувають прогресивні загрози, пов'язані з використанням штучного інтелекту для автоматизації атак, генерації deepfake-контенту, підготовки гіперперсоналізованих фішингових повідомлень, а також загрози, пов'язані з неправильною конфігурацією хмарних середовищ, шкідливими мобільними застосунками та вразливістю IoT-пристроїв.

Система кіберзахисту має ґрунтуватися на принципах конфіденційності, цілісності та доступності (CIA), а також застосовувати розширені механізми, зокрема багатофакторну аутентифікацію, принцип найменших привілеїв, сегментацію мережі, непримовність дій користувачів. Технічна реалізація включає використання сучасних засобів кіберзахисту: міжмережевих екранів (firewalls), систем виявлення та запобігання вторгненням (IDS/IPS), веб-аплікаційних фаєрволів (WAF), віртуальних приватних мереж (VPN), платформ захисту кінцевих точок (EDR/NGAV), SIEM-систем і центрів моніторингу (SOC). ENISA наголошує на доцільності автоматизації реагування та інтеграції механізмів Threat Intelligence для забезпечення адаптивної стійкості систем.

Організаційні заходи відіграють не менш важливу роль — це формування політик безпеки, підготовка планів реагування на інциденти, проведення регулярних тренінгів для персоналу, аудит інформаційних систем і практики red/blue teaming. Людський фактор залишається визначальним: згідно з DBIR 2024, близько 68 % інцидентів пов'язані з помилками користувачів або соціальною інженерією. У зв'язку з цим інвестиції в навчальні програми, моделювання фішингових атак і регулярний аудит облікових записів є критично важливими для зниження ризиків.

Серед ключових тенденцій майбутнього розвитку сфери кібербезпеки виокремлюються перехід до пост-квантової криптографії (PQC), імплементація принципу Security by Design на всіх етапах життєвого циклу розробки програмного забезпечення, впровадження систем автоматизованого реагування (SOAR), а також поглиблена аналітика даних з використанням штучного інтелекту. На тлі глобального дефіциту кадрів у сфері кібербезпеки особливого значення набувають програми підготовки фахівців і розвиток інструментів автоматизації, що дозволяють компенсувати нестачу людських ресурсів у критичних сегментах захисту інформаційних систем.

Таким чином, у контексті зростаючих цифрових загроз, гібридних воєнних дій та стрімкої еволюції кіберзлочинності, забезпечення ефективної кібербезпеки набуває системного та багаторівневого значення. Сучасні виклики вимагають

переходу від реактивного до проактивного підходу, орієнтованого на попередження інцидентів, комплексне управління ризиками та інтеграцію безпеки на всіх етапах функціонування ІТ-систем. Використання сучасних технологій, таких як багатофакторна аутентифікація, SIEM/SOC-платформи, штучний інтелект і автоматизовані засоби реагування, має поєднуватися з організаційними заходами — розробкою політик безпеки, навчанням персоналу, внутрішнім аудитом та міжсекторальною співпрацею. Лише цілісний, міждисциплінарний підхід здатен забезпечити належний рівень стійкості цифрового середовища як для державних структур, так і для бізнесу.

Список використаної літератури

1. Verizon, 2024 Data Breach Investigations Report – Executive Summary, Apr. 2024. [SpyCloud](#)
2. Verizon, Key Insights from DBIR 2024, SpyCloud, Jul. 2024. [SpyCloud](#)
3. SANS Institute, Tackling Modern Human Risks in Cybersecurity, Insights from the Verizon DBIR 2024. [SANS Institute](#)
4. Reuters, Ukraine's Kyivstar allocated \$90 million to deal with cyberattack aftermath, May 20 2024. [Reuters](#)
5. ENISA, Threat Landscape: Finance Sector, Feb. 2025. [ENISA](#)
6. AP News, Europe's cybersecurity chief: disruptive attacks doubled in 2024, May 2024. [AP News](#)

ІННОВАЦІЙНІ МЕТОДИ КОНТРОЛЮ ДОСТУПУ ПЕРСОНАЛУ В ІНФОРМАЦІЙНИХ СИСТЕМАХ ПІДПРИЄМСТВ НА ОСНОВІ БЛОКЧЕЙН-ТЕХНОЛОГІЙ

*Ямнич А. Б.
аспірант кафедри безпеки інформаційних технологій
Коробейнікова Т. І.
к.т.н., доцент
доцент кафедри безпеки інформаційних технологій
Національний університет «Львівська політехніка»*

Розвиток цифрових технологій та поширення віддалених форм роботи зумовлюють зростаючу потребу у вдосконаленні механізмів контролю доступу персоналу до інформаційних систем підприємств. У сучасних умовах кібербезпеки традиційні методи автентифікації та авторизації виявляють низку недоліків, зокрема високу вразливість до атак, залежність від централізованих систем управління та складність масштабування. Одним із перспективних напрямів вирішення цих проблем є використання блокчейн-технологій, які забезпечують децентралізований контроль доступу та підвищують рівень безпеки інформаційних ресурсів.

Застосування блокчейну в управлінні доступом дозволяє реалізувати механізми ідентифікації та автентифікації, що не залежать від єдиного центру перевірки прав доступу. У цьому підході дані про дозволи зберігаються у розподіленій мережі, що унеможливорює їх підробку або несанкціоновану зміну.

Крім того, використання смарт-контрактів автоматизує процеси перевірки та надання доступу, знижуючи потребу у втручанні адміністратора та зменшуючи ризик людського фактора.

Однією з ключових переваг інтеграції блокчейну є можливість гнучкого керування правами доступу відповідно до ролей користувачів та рівня їхніх повноважень. Це особливо важливо для підприємств, які працюють у середовищі з високими вимогами до безпеки даних. Використання блокчейну дозволяє вести прозорий облік всіх операцій доступу, що забезпечує їхню перевірюваність та відстежуваність.

У базовій реалізації запропонованого технологічного ланцюжка система складається з кількох ключових компонентів: користувача, який ініціює запит на доступ; ресурсу, до якого надається доступ; сервісу авторизації; блокчейн-мережі, де зберігається інформація про політики доступу; і смарт-контрактів, що автоматизують перевірку прав. У цьому ланцюжку користувач формує запит, який через сервіс авторизації перевіряється за допомогою смарт-контракту, що містить правила доступу. У разі підтвердження відповідності запиту – доступ до ресурсу дозволяється.

Цей підхід дозволяє усунути посередників, централізовані сервери авторизації, які можуть бути вразливими до атак або працювати нестабільно. Усі перевірки виконуються на основі консенсусу в мережі, що забезпечує високий рівень достовірності даних. Крім того, блокчейн-фіксація кожного запиту створює надійний журнал подій, який може бути використаний для аудиту та розслідувань інцидентів безпеки.

Однак, навіть така модель не є остаточною, адже вона не враховує всіх особливостей реального корпоративного середовища, зокрема потреби у врахуванні додаткових контекстних факторів доступу, підтримці гнучких схем делегування прав, або інтеграції з уже наявними ІТ-системами. Саме тому доцільним є подальше вдосконалення базової архітектури.

Модифікований варіант технологічного ланцюжка передбачає розширення функціональності базової моделі за рахунок інтеграції додаткових компонентів, що враховують контекст запиту доступу. Зокрема, пропонується використання сервісу контекстного аналізу, який перевіряє параметри середовища (геолокацію, тип пристрою, час доступу, попередню активність користувача тощо) і на їх основі формує зважене рішення щодо надання прав.

Також у модифікованій архітектурі реалізується підтримка механізмів делегування доступу на основі політик, визначених у смарт-контрактах. Це дозволяє тимчасово або умовно передавати права іншому користувачу з чітким контролем меж дозволеного. Такий підхід особливо актуальний для організацій зі складною структурою, де часто виникає потреба в гнучкому управлінні доступом між підрозділами або співробітниками на різних рівнях.

Додатково, модифікований ланцюжок передбачає кращу інтеграцію з корпоративними системами управління, зокрема через застосування API та відповідних протоколів обміну. Це дає змогу впроваджувати блокчейн-рішення без необхідності повної заміни існуючих інфраструктур, що знижує витрати та спрощує перехід до нової моделі безпеки.

Висновки: запропонований модифікований технологічний ланцюжок контролю доступу на основі блокчейн-технологій дозволяє вирішити низку проблем, характерних для традиційних централізованих систем. Завдяки використанню смарт-контрактів, токенів доступу та розподіленого реєстру досягається підвищена безпека та надійність управління доступом. Інтеграція з методами поведінкової аналітики та мультипідпису сприяє гнучкому адаптуванню політик контролю доступу в реальному часі, що зменшує ризики компрометації даних.

Практичне впровадження цієї моделі в корпоративні інформаційні системи підприємств забезпечує розширену функціональність та підвищує довіру до механізмів автентифікації та авторизації. Подальші дослідження можуть бути спрямовані на вдосконалення методів збереження криптографічних ключів та інтеграцію з децентралізованими ідентифікаційними системами, що дозволить ще більше підвищити рівень безпеки та ефективність контролю доступу.

Список використаної літератури

1. Kurii and I. Oprisky, «Analysis and Comparison of the NIST SP 800-53 and ISO/IEC 27001: 2013» *NIST Spec. Publ.*, pp. 10-13, 2022.
2. T. Korobeinikova, I. Tachenko, R. Chekhmestruk, P. Mykhaylov, O. Romanyuk and S. Romanyuk, «A General Method of Risk Estimation,» 2023 13th International Conference on Advanced Computer Information Technologies (ACIT), Wrocław, Poland, 2023, pp. 410-413, doi: 10.1109/ACIT58437.2023.10275626
3. A. Yamnych, T. Korobeinikova «A Model For Personnel Access Control To Information Resources Of Industrial Enterprises Based On RBAC And Blockchain Technology» / *Herald of Khmelnytskyi National University. Technical sciences*, vol. 343, no. 6(1), pp. 380–386, Nov. 2024, doi: 10.31891/2307-5732-2024-343-6-56.

МОДЕЛЮВАННЯ ЗАГРОЗ ДЛЯ СИСТЕМ УПРАВЛІННЯ НАВЧАННЯМ, ПІДСИЛЕНИХ ШТУЧНИМ ІНТЕЛЕКТОМ

*Янішевський В. І.,
аспірант,
ПВНЗ «Європейський університет»*

Упродовж останнього десятиліття стрімкий розвиток штучного інтелекту (ШІ) радикально трансформував освітній ландшафт, змінивши підходи до викладання, навчання та адміністрування освітнього процесу. Одним із ключових напрямів цієї трансформації стало впровадження інтелектуальних систем управління навчанням (Learning Management Systems, LMS), які забезпечують персоналізований підхід до навчання, автоматизацію рутинних завдань та інтерактивний зворотний зв'язок завдяки використанню технологій інтелектуальної аналітики.

Застосування таких можливостей ШІ, як обробка природної мови, машинне навчання та адаптивні алгоритми, зробило платформи LMS більш гнучкими й адаптивними до потреб окремих користувачів. Проте така інноваційність несе й нові загрози — зокрема у сфері кібербезпеки. Ці загрози вже не обмежуються

класичними вразливостями програмного забезпечення, а охоплюють специфічні ризики, пов'язані з маніпуляцією алгоритмами ШІ, порушенням конфіденційності даних та можливим зловживанням аналітичними інструментами.

У контексті активної цифровізації освітніх процесів безпечність використання LMS, підсилених ШІ, набуває критичного значення. Потенційні наслідки — від несанкціонованого доступу до навчальної інформації до підриву довіри студентів та викладачів, а також дискредитації академічної доброчесності.

Визначення проблеми

Попри численні переваги, впровадження ШІ в LMS суттєво розширює поверхню потенційних кіберзагроз. Традиційні підходи до моделювання загроз не завжди в змозі ідентифікувати специфічні атаки, спрямовані на моделі машинного навчання — такі як отруєння даних (data poisoning), інверсія моделі (model inversion) чи змагальні вхідні атаки (adversarial examples). Таким чином, виникає нагальна потреба у створенні адаптованих до ШІ методологій аналізу для систем, що інтегрують інтелектуальні компоненти.

Метою дослідження є подолання прогалини між функціональністю ШІ та безпечним проектуванням освітніх цифрових систем шляхом застосування сучасних підходів до моделювання загроз у контексті LMS, вдосконалених штучним інтелектом.

Структура моделювання загроз: STRIDE

Модель STRIDE, розроблена Microsoft, використовується як основна структура класифікації загроз. Він класифікує загрози на шість типів:

- Спуфінг (викривлення особистих даних)
- Втручання (несанкціонована зміна даних)
- Відмова (відмова від дій)
- Розкриття інформації (витік даних)
- Відмова в обслуговуванні (DoS) (збій системи)
- Підвищення привілеїв (несанкціоноване отримання привілеїв)

Кожен компонент у системі оцінюється за допомогою категорій STRIDE, доповнених спеціальними можливостями ШІ (наприклад, отруєння моделі під час втручання або протилежне введення під час DoS)[1].

Вектори атак ШІ в контексті STRIDE: деталізація та приклади

При моделюванні загроз у системах управління навчанням (LMS), підсилених компонентами штучного інтелекту, важливо не лише класифікувати традиційні ризики, а й враховувати специфічні вектори атак, характерні для машинного навчання. Ці атаки, хоча й мають нетрадиційну природу, добре вкладаються в рамки моделі STRIDE, якщо адаптувати категорії під контекст ШІ. Нижче подано аналіз чотирьох основних загроз:

1. Отруєння даних (Data Poisoning)

Суть атаки: Зловмисник свідомо вводить шкідливі або маніпулятивні дані в набір для навчання ШІ-моделі, що призводить до викривлення її поведінки. У результаті модель починає видавати некоректні або навіть небезпечні результати.

Зв'язок із STRIDE: Це класична форма втручання — зловмисник змінює цілісність даних, які формують логіку прийняття рішень ШІ.

Приклад: У LMS студент подає завдання, яке містить шкідливі патерни у вигляді навчальних закладів (наприклад, специфічні комбінації слів у есе), які спричиняють адаптацію моделі перевірки текстів, внаслідок чого ШІ починає хибно вважати певні шаблони академічним плагіатом або, навпаки, пропускає його.

2. Змагальні вхідні дані (Adversarial Inputs)

Суть атаки: ШІ-системи виявляють вразливість до незначних модифікацій вхідних даних, які для людини непомітні, але здатні радикально змінити поведінку моделі.

Зв'язок із STRIDE: неправильна класифікація або аварійна реакція системи може спричинити її відмову або блокування доступу. Атака маскує справжні наміри користувача або фальсифікує введення.

Приклад: У системі LMS, що автоматично оцінює відповіді студентів, спеціально змінений текст (наприклад, невидимі символи, синоніми зі зсунутим контекстом) може змусити ШІ оцінити правильну відповідь як помилкову або навпаки. Це створює вразливість для зловживань при автоматичній перевірці.

3. Інверсія моделі (Model Inversion)

Суть атаки: Шляхом багаторазових запитів до моделі, зловмисник може реконструювати фрагменти навчального набору, які містять конфіденційну інформацію.

Зв'язок із STRIDE: В результаті такої атаки відбувається витік даних, що порушує конфіденційність користувачів LMS, зокрема студентів.

Приклад: Уявімо, що викладач завантажує до LMS модуль адаптивного тестування на базі ШІ, який тренується на реальних відповідях студентів. Зловмисник (наприклад, технічно обізнаний студент) може багаторазово звертатися до API моделі, моделюючи запити, які дозволяють йому частково відновити початкові відповіді інших студентів — включаючи особисті думки чи конфіденційну інформацію.

Список використаної літератури

1. Liu, X. et al. (2021). *Privacy and Security Issues in Deep Learning: A Survey*. *IEEE Access*, 9, 4566–4593. <https://doi.org/10.1109/ACCESS.2020.3045078>
2. OWASP Foundation. (2023). *OWASP Threat Dragon Documentation*. <https://owasp.org/www-project-threat-dragon/>
3. Papernot, N., McDaniel, P., Sinha, A., & Wellman, M. P. (2018). *SoK: Security and Privacy in Machine Learning*. *Proceedings of the 2018 IEEE European Symposium on Security and Privacy*. <https://arxiv.org/abs/1804.10117>
4. Biggio, B., & Roli, F. (2018). *Wild Patterns: Ten Years After the Rise of Adversarial Machine Learning*. *Pattern Recognition*, 84, 317–331. <https://doi.org/10.1016/j.patcog.2018.07.023>

ВПЛИВ ВИБОРУ МОВИ ПРОГРАМУВАННЯ НА СТІЙКІСТЬ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ ДО КІБЕРАТАК

*Яровий Р.О.,
кандидат технічних наук, доцент,
доцент кафедри кібербезпеки та захисту інформації,
Рихальський О.Ю.,
викладач кафедри комп'ютерних наук та програмної інженерії
ПВНЗ «Європейський університет»,
ст. викладач «DAN.IT Education».*

Вибір мови програмування [1,4] суттєво впливає на потенційну вразливість програмного забезпечення до кіберзагроз. Вона визначає набір доступних засобів для розробника, включно з можливостями керування пам'яттю, обробкою виключень та інтеграцією із зовнішніми бібліотеками. Деякі мови надають більш високий рівень абстракції, що ускладнює створення небезпечного коду, тоді як інші надають розробнику повний контроль над ресурсами, що може призвести до критичних помилок. Саме тому правильний вибір мови має стратегічне значення для формування стійкості ПЗ до атак.

Мови з автоматичним керуванням пам'яттю, такі як Java, Rust знижують ризик помилок, пов'язаних із переповненням буфера. Оскільки більшість атак на рівні пам'яті базуються на прямому доступі до масивів або небезпечних маніпуляціях із вказівниками, автоматичне керування пам'яттю дозволяє уникнути таких проблем. У мовах як Java або Rust розробник не має прямого доступу до вказівників, і більшість дій, що можуть призвести до переповнення буфера, контролюються віртуальною машиною або компілятором. Це не гарантує повної безпеки, але значно зменшує ймовірність типових помилок.

Натомість мови низького рівня, наприклад C або C++, дають більше контролю, але вимагають ретельнішого дотримання правил безпеки. Вони дозволяють працювати безпосередньо з пам'яттю, маніпулювати покажчиками, вручну виділяти й звільняти ресурси. Цей підхід відкриває широкі можливості для оптимізації, але одночасно створює серйозні ризики. Без належного знання стандартів безпеки і постійної перевірки коду такі мови часто стають причиною вразливостей, що експлуатуються хакерами для виконання довільного коду або витоку даних.

Rust [2,3] пропонує інноваційний підхід до управління пам'яттю без використання збирача сміття, забезпечуючи високий рівень безпеки на рівні компіляції. Rust використовує систему запозичень і часу життя, які гарантують безпеку пам'яті ще на етапі компіляції. Завдяки цьому багато типових помилок, властивих C або C++, стають неможливими. Мова забороняє доступ до звільненої пам'яті, використання неініціалізованих змінних та інші небезпечні операції, що суттєво підвищує стійкість створеного ПЗ до експлуатаційних атак.

Типізовані мови дозволяють виявляти потенційні помилки ще до запуску програми, що підвищує її стійкість до атак. Статична типізація дає змогу виявляти багато логічних помилок на етапі компіляції, зокрема невідповідності типів, неправильне використання змінних або невраховані гілки умов. Завдяки цьому

зменшується кількість вразливостей, які можуть бути використані під час виконання програми.

Підтримка сучасних криптографічних бібліотек і безпечних API є критичною характеристикою безпечної мови програмування. Мова повинна мати якісні бібліотеки для шифрування, хешування, цифрових підписів та інших засобів безпечної обробки даних. При цьому важливо, щоб ці бібліотеки реалізовували сучасні стандарти криптографії, проходили незалежний аудит і мали надійну документацію. Без цього навіть добре написана програма залишиться вразливою через використання застарілих або небезпечних алгоритмів.

Статичний аналіз коду та інструменти перевірки безпеки часто залежать від особливостей конкретної мови. Ефективність аналізу коду на предмет вразливостей залежить від того, наскільки добре інструменти аналізу інтегруються з мовою. Наприклад, для Java існує багато зрілих рішень для перевірки на рівні байткоду, а для Rust – спеціалізовані лінери та інструменти перевірки властивостей. У мовах з динамічною типізацією інструменти аналізу мають меншу точність, що збільшує ризик пропущених вразливостей.

Вибір мови також визначає доступність спільноти, документації та практик безпечного програмування. Спільнота навколо мови програмування часто формує культуру безпеки: наявність прикладів, гайдів і поширених практик програмування значно полегшує створення безпечного ПЗ. Наприклад, Rust-спільнота активно просуває принцип «безпека за замовчуванням».

Деякі мови мають вбудовані механізми перевірки вхідних даних, що запобігає поширеним атакам типу SQL-ін'єкцій [1,3]. У деяких мовах реалізовано концепції типобезпечних запитів до баз даних або шаблонів для генерації HTML, які мінімізують ризик впровадження шкідливого коду. Наприклад, в Haskell ускладнено створення SQL-ін'єкції через уніфіковану обробку типів, а у сучасних web-фреймворках на Scala чи Rust – шаблони перевіряються ще до компіляції. Це значно зменшує ризик критичних атак на рівні інтерфейсу.

Загалом, мова програмування не гарантує повної захищеності [1], але правильний вибір суттєво зменшує поверхню атаки. Жодна мова програмування не є абсолютно безпечною, адже остаточна безпека залежить від архітектури, реалізації і людського фактору. Проте мови, які мають вбудовані механізми контролю, строгі компілятори і підтримку безпечних практик, допомагають мінімізувати кількість помилок і зробити систему більш захищеною до атак. Таким чином, мова – це не панацея, але важливий інструмент у створенні стійкого програмного забезпечення.

Список використаної літератури

1. Kannan S. *Your Programming Language and Its Impact on the Cybersecurity of Your Application.* <https://tarides.com/blog/2023-08-17-your-programming-language-and-its-impact-on-the-cybersecurity-of-your-application/>
2. What is memory safety and why does it matter? <https://www.memorysafety.org/docs/memory-safety/>
3. Sible J., Svoboda D. *Rust Software Security: A Current State Assessment* <https://insights.sei.cmu.edu/blog/rust-software-security-a-current-state-assessment/>
4. The 'Viral' Secure Programming Language That's Taking Over Tech –WIRED. <https://www.wired.com/story/rust-secure-programming-language-memory-safe/>

КІБЕРЗАХИСТ КРИТИЧНИХ ОБ'ЄКТІВ ЗАЛІЗНИЧНОЇ ІНФРАСТРУКТУРИ

*Яровий Я.Р.
студент,*

Український державний університет залізничного транспорту

Сучасна залізнична інфраструктура, як одна з основних компонентів критичних об'єктів транспортної системи, все частіше стає мішенню кібератак через зростаючий рівень цифровізації, інтеграцію ІТ та ОТ-систем, використання застарілих контролерів, а також тривалий життєвий цикл обладнання. Оскільки залізничні системи працюють у безперервному режимі (24×7×365) і обслуговують як пасажирські, так і вантажні перевезення, атаки на інфраструктуру можуть мати катастрофічні наслідки, включаючи транспортний колапс, шкоду для життя та значні економічні втрати.

Основні загрози охоплюють як технічні, так і організаційні вектори атак. Уразливості Eurobalise, зокрема відсутність вбудованого криптозахисту, створюють ризик для цілісності даних передачі (Lim et al., 2017). Наявність слабо сегментованих мереж між ІТ та ОТ-середовищами дозволяє атакам розповсюджуватися через так звані «бічні двері». Поширеною є також експлуатація вразливостей у ПЛК і RTU, зокрема з CVE, які вже відомі, але не були усунуті через брак оновлень або підтримки.

Цифровізація транспортної інфраструктури, включно з впровадженням IoT та ПоТ-пристроїв, розширила поверхню атаки, залучаючи як державних акторів, так і кримінальні угруповання (ENISA, 2022). Окремо слід зазначити атаки на Wi-Fi-інфраструктуру вокзалів та спроби саботажу систем сигналізації в країнах ЄС, що підкреслюють актуальність досліджень у сфері залізничної кібербезпеки.

Нормативна база формує основу для побудови кіберзахисту в галузі. У Європейському Союзі ключовими є ENISA-доповіді (2022), технічна специфікація CLC/TS 50701:2023, а також Делегований Регламент ЄС 2022/30/EU, який накладає вимоги на проектування кіберзахищеного обладнання. У США актуальним є Security Directive 1580/82 від TSA (2023). Для уніфікації підходів до захисту ОТ доцільним є застосування рекомендацій IEC 62443 та NIST SP 800-82 Rev. 3.

Технічні заходи включають впровадження сегментації мережі через DMZ між ІТ та ОТ-зонами, використання багатофакторної аутентифікації для доступу до інженерних інтерфейсів, а також впровадження IDS/IPS для ОТ-середовищ. Важливою є також інтеграція логів з SCADA-систем до SIEM із підтримкою UEBA (User and Entity Behavior Analytics), що дає змогу своєчасно виявляти аномалії.

Криптографічний захист передбачає використання легковагових алгоритмів шифрування (SotoCruz et al., 2025) для обмежених у ресурсах пристроїв, таких як баліси, а також створення захищених каналів VPN для віддаленого доступу.

Організаційні заходи охоплюють інцидент-менеджмент (визначення SLA, процедури ескалації, координація з CISA та ENISA), навчання персоналу на всіх рівнях, створення кіберполігонів (cyber ranges) та проведення регулярних аудитів відповідно до IEC 62443 та CLC/TS 50701. Використання LLM (великомасштабних мовних моделей) для автоматизованого документування та перевірки відповідності

(BCA → PCA) є інноваційним рішенням для підвищення ефективності оцінки безпеки ОТ-систем.

Побудова цілісної моделі управління ризиками неможлива без комплексного поєднання технічних, нормативних і організаційних засобів. Підвищення кіберстійкості критичної інфраструктури передбачає не лише адаптацію до поточних загроз, але й активне впровадження інтелектуальних технологій, нормативну еволюцію та розвиток міжвідомчої співпраці. Рекомендовано розвивати регіональні кібернавчальні платформи, що спеціалізуються на транспортній галузі, як частину загальнодержавної стратегії кіберзахисту.

Список використаної літератури:

1. *European Union Agency for Cybersecurity (ENISA), Railway Cybersecurity: Good Practices in Cyber Risk Management, ENISA, 2022.*
2. *European Committee for Electrotechnical Standardization (CENELEC), CLC/TS 50701:2023 Railway Applications – Cybersecurity, 2023.*
3. *European Union, Commission Delegated Regulation (EU) 2022/30 of 16 December 2021, Official Journal L 5, 2022.*
4. *U.S. Department of Homeland Security, TSA Security Directive 1580/82202201C, effective 1 July 2024.*
5. *International Electrotechnical Commission (IEC), IEC 62443 Series: Industrial Communication Networks – Network and System Security, 2018.*
6. *ENISA, Railway Cybersecurity – Study on Implementation of Cybersecurity Measures in the Railway Sector, 2020.*
7. *ENISA, ENISA Threat Landscape 2022 – Transport Sector, November 2022.*
8. *Alstom S.A., “Towards the first railway cybersecurity international standard,” press release, 13 March 2024.*
9. *C. Schlehuber, HandsOn CLC/TS 50701 (Railway Applications – Cybersecurity), ERA Workshop, 2023.*
10. *H. W. Lim et al., “Data Integrity Threats and Countermeasures in Railway Spot Transmission Systems,” arXiv:1709.05935, 2017.*
11. *J. SotoCruz et al., “A Survey of Efficient Lightweight Cryptography for Power-Constrained Microcontrollers,” Technologies 13(1), 3 (2025).*
12. *NIST, SP 800-82 Rev. 3: Guide to OT Security, September 2023.*

ANALYSIS OF THE USE OF SIEM SYSTEMS FOR THREAT MONITORING IN CRITICAL INFRASTRUCTURE

V. O. Bilyi

*Ph.D. in Physics and Mathematic,
Senior Lecturer at Special Department No. 1,*

M. S. Kamenev

Student,

ISZZI of National Technical University of Ukraine «Igor Sikorsky KPI»

With the rapid development of digital technologies, the amount of data processed in the IT systems of enterprises and public institutions is increasing. This complicates the

tasks of cybersecurity specialists, as manual threat monitoring and response become ineffective. This issue is especially relevant for Ukraine's critical infrastructure, which, under full-scale invasion, remains under constant pressure from cybercriminals and hostile state-sponsored hacker groups. To counter these threats, SIEM systems are used to provide centralized collection, analysis, and correlation of security events, enabling timely detection and blocking of attacks.

SIEM (Security Information and Event Management) systems combine the functionalities of Security Event Management (SEM) and Security Information Management (SIM). They collect data from logs of network equipment, servers, antivirus software, firewalls, analyze events, and generate alerts about potential threats. The core functions of such systems include centralized collection and processing of security events, real-time detection and analysis of anomalous activity, as well as risk assessment and threat management.

SIEM systems significantly enhance the effectiveness of cybersecurity for critical infrastructure. One of their key advantages is the optimization of event processing, allowing a reduction in the number of events requiring analysis and focusing on the most critical incidents. Additionally, SIEM systems identify threats by tracking deviations in user and process behavior that may indicate security breaches. Automated alerting ensures timely notification of administrators about potential threats, and deep data analysis enables event normalization, correlation, and detailed security compliance reporting.

The use of SIEM in critical infrastructure involves continuous monitoring of attacks targeting energy, transportation, and public administration facilities, enabling early-stage threat detection. An important aspect is identifying anomalous actions within systems that may indicate attempted network compromise or unauthorized access. Furthermore, SIEM systems provide automated incident response through integration with protection tools such as firewalls, antivirus software, and access control systems. They also allow for the generation of detailed reports on security events, helping analyze attack trends and improve cybersecurity strategies.

For example, among the most popular SIEM systems are Splunk, IBM QRadar, and ArcSight. Splunk is known for its high scalability and extensive integration capabilities, but it comes with a high cost. IBM QRadar provides in-depth threat analytics but requires significant resources for deployment. ArcSight, in turn, is a flexible solution with an advanced event correlation engine, though it is complex to configure.

Real-world cases demonstrate the effectiveness of SIEM systems in preventing serious cyber incidents, such as attacks on government servers or banking institutions. The use of advanced SIEM solutions like Splunk, IBM QRadar, and ArcSight provides centralized threat control and management in public institutions and large enterprises. However, challenges arise during the implementation of these systems, including high costs, configuration complexity, and the need for highly qualified personnel. Nevertheless, practical experience shows that effective use of SIEM helps detect complex threats, such as multi-stage attacks or long-term data theft campaigns, which could go unnoticed without deep analysis.

Conclusion: Modern cybercriminals use sophisticated attack methods, exploiting weaknesses in security systems. Without contextual analysis, individual events may remain unnoticed, posing significant risks to critical infrastructure. SIEM systems enable

effective monitoring and event analysis, allowing rapid threat detection and prevention of serious incidents. The further development of such technologies, including integration with artificial intelligence systems, will significantly increase the efficiency of cybersecurity.

References:

1. ArcSight. *SIEM for Threat Detection and Response. Micro Focus: [Website]. URL: <https://www.microfocus.com/en-us/cyberres/secops/arcsight-siem>*
2. IBM Corporation. *IBM QRadar SIEM Overview: [Website]. URL: <https://www.ibm.com/security/security-intelligence/qradar>*
3. Splunk Inc. *[Website]. URL: <https://www.splunk.com/>*

AI-ENHANCED PHISHING SIMULATIONS: OPPORTUNITIES AND RISKS

A.S. Glova¹

V.I. Kotsun²

¹PhD candidate

²Candidate of Technical Sciences, Associate Professor, Head of the Department of Mathematics and Computer Sciences, PHEE «European University»

The evolving landscape of cybersecurity is increasingly shaped by the capabilities of artificial intelligence (AI). One critical application is the enhancement of phishing simulation training programs. By leveraging AI, simulations can dynamically generate realistic and context-specific phishing attacks, offering more sophisticated and adaptive learning environments [1]. However, while AI-enhanced phishing simulations present significant opportunities to improve cybersecurity awareness and defense mechanisms, they also introduce new risks, including ethical concerns and the potential misuse of AI technologies [2].

Opportunities Offered by AI-Enhanced Simulations. Recent studies have shown that AI can significantly improve the realism and effectiveness of phishing simulations by analyzing user behavior and crafting personalized phishing emails. Machine learning models can predict user vulnerabilities and dynamically adjust the complexity and content of phishing scenarios, ensuring that users are continually challenged and engaged [3].

Natural language generation (NLG) techniques allow AI to create phishing emails that mimic the writing style and communication patterns of trusted contacts, making simulations more convincing and training more impactful [4]. Furthermore, reinforcement learning algorithms can adapt training sequences based on real-time user responses, enhancing learning outcomes over traditional static methods.

Risks and Ethical Concerns. Despite these advancements, AI-enhanced phishing simulations are not without risks. One major concern is overpersonalization, where highly tailored simulations could potentially lead to breaches of employee privacy or psychological distress [5]. Ethical boundaries must be established to ensure that simulations remain educational rather than punitive.

Another risk is the potential for adversarial use. If cybercriminals gain access to AI tools designed for simulations, they could create more dangerous and sophisticated phishing attacks that are harder to detect [6]. Thus, organizations must implement strict controls and monitoring over AI models used in security training.

Additionally, the bias in training data can impact AI model behavior. If the datasets used to train AI models are not diverse, certain groups of employees may be unfairly targeted or overlooked in simulations, reducing overall training effectiveness and possibly fostering workplace discrimination [2].

Future Directions. To maximize the benefits and mitigate the risks of AI-enhanced phishing simulations, several strategies are recommended:

- Ethical AI frameworks must guide the development and deployment of simulation programs.
- Continuous model evaluation and bias detection mechanisms should be employed.
- Organizations should maintain transparency with employees about the use of AI in training processes.
- Partnerships between cybersecurity professionals, psychologists, and ethicists are essential to design effective, safe, and respectful training programs [8].

Conclusion. AI-enhanced phishing simulations represent a powerful advancement in cybersecurity training, offering unprecedented realism, adaptability, and user engagement. However, their implementation must be carefully managed to address privacy, ethical, and adversarial concerns. A balanced approach, grounded in ethical practices and continuous oversight, will ensure that AI serves as a force for strengthening, rather than undermining, organizational cybersecurity defenses.

References

1. Chen, F., Wu, T., Nguyen, V., Wang, S., Hu, H., Abuadbba, A., & Rudolph, C. (2024). *Adapting to Cyber Threats: A Phishing Evolution Network (PEN) Framework for Phishing Generation and Analyzing Evolution Patterns using Large Language Models*. arXiv. <https://arxiv.org/abs/2411.11389>
2. Owen-Jackson, C. (2024). *Navigating the ethics of AI in cybersecurity*. IBM Security Intelligence. <https://securityintelligence.com/articles/navigating-ethics-ai-cybersecurity/>
3. Keepnet Labs. (2025). *How Keepnet's AI-Powered Phishing Simulator Delivers Hyper-Personalized Security Awareness*. <https://keepnetlabs.com/blog/how-keepnets-ai-powered-phishing-simulator-delivers-hyper-personalized-security-awareness>
4. Heiding, F., Lermen, S., Kao, A., Schneier, B., & Vishwanath, A. (2024). *Evaluating Large Language Models' Capability to Launch Fully Automated Spear Phishing Campaigns: Validated on Human Subjects*. arXiv. <https://arxiv.org/abs/2412.00586>
5. PhishDeck. (n.d.). *Keeping Phishing Tests Ethical*. <https://phishdeck.com/blog/keeping-phishing-tests-ethical>
6. Begou, N., Vinoy, J., Duda, A., & Korczynski, M. (2023). *Exploring the Dark Side of AI: Advanced Phishing Attack Design and Deployment Using ChatGPT*. arXiv. <https://arxiv.org/abs/2309.10463>

ADVERSARIAL ATTACKS ON CYBERSECURITY WITH MACHINE LEARNING IN DEFENCE

Taras Melko¹

Volodymyr Kotsun²

¹Ph.D Candidate, Department of Computer Science and Software Engineering

²Ph.D Tech, Associate Professor of the Department of Computer Science and Software Engineering

PHEE «European University»

In recent years, the utilization of machine learning—the process through which computers can be programmed to recognize and learn the repetition of certain patterns, resulting in time-based forecasts—has steadily increased. Machine learning is an essential component of artificial intelligence (AI), which is used in computer vision, voice recognition, robot technology, self-driving vehicles, and even in simple applications like spam email and malware detection [1].

The employment of machine learning (ML) for not only detecting but also mitigating cyber threats raises concern for adversarial attacks on cybersecurity systems. The refinement of consequent attacks utilizes the weaknesses of ML-enabled system defenses, endangering systems such as intrusion recognition, malware classification, and anomaly detection. They comprise yet another perilous challenge. This is where malicious users actively attempt to mislead an ML model by injecting noise into the data. These ‘adversarial examples’ have been shown to completely alter the decisions of the model with little to no changes in input data, creating problems, especially for intrusion detection systems (IDS).

Adversarial attacks on cybersecurity AI systems exploit vulnerabilities in models used for threat detection, such as intrusion detection systems and malware classifiers, by manipulating inputs to cause misclassification or evade detection. These attacks, including data poisoning and evasion tactics, undermine AI reliability, with attackers injecting corrupted data into training sets or crafting subtle perturbations to bypass real-time defenses. For instance, data poisoning can skew an AI’s learning process, rendering it ineffective against new threats [2], while evasion attacks have been shown to trick AI-based antivirus systems [3]. The dual-use nature of AI amplifies this challenge, as attackers leverage similar technologies to enhance attack sophistication, such as AI-driven DDoS attacks reaching 398 million requests per second [5]. Addressing these vulnerabilities requires robust defenses, including secure training data and adversarial resilience.

There are protective measures such as Adversarial Training and Robust Optimization. Adversarial Training stands as one of the most favored approaches used to improve the resilience of ML models from attacks. The core concept is to enhance the training set’s pre-processing by integrating adversarial examples, which are data points that closely resemble the original samples but are deceptively altered in a manner that

is bound to confuse the model, thereby resulting in misclassifications. The model learns to classify both normal and altered data correctly, and therefore, it can tolerate such manipulations later on. Another approach, Robust Optimization, defines models with resistance to a great deal of possible perturbations within the data. Rather than reacting to specific attacks as in adversarial training, this method aims to ensure that a model's decisions remain accurate within the worst-case perturbation scenarios that exist within a given radius, such as L_2 or L_∞ norm variance [5]. Still, developing robust models is quite a challenging task because of the constantly changing nature of hacking techniques.

AI's integration into cybersecurity enables richer tools to fight against digital threats, however, challenges and limitations still exist. From resource allocation, ethics and regulation gaps, organizations will need to adopt AI with care while actively engaging human endeavors working on AI's weaknesses, putting in place strong fences to protect sensitive operations. With the progress of time, the changing face of threat, and continuous research, inter-industry, intra-academia, and inter-policy collaboration will help mitigate these challenges towards allowing AI to become a more positive force in cybersecurity.

References

1. *Adversarial Machine Learning* (2019), <https://cltc.berkeley.edu/aml>
2. Sharma, S., Kumar, R., & Singh, K. (2023). Role of artificial intelligence in cybersecurity: A comprehensive review. *Computers & Security*, 132, 103321. <https://doi.org/10.1016/j.cose.2023.103321>
3. eSecurity Planet. (2024, January 15). AI in cybersecurity: Benefits, challenges, and use cases. <https://www.esecurityplanet.com/trends/ai-in-cybersecurity/>
4. Palo Alto Networks. (2023, October 25). The rise of AI-powered DDoS attacks. <https://www.paloaltonetworks.com/blog/2023/10/ai-powered-ddos-attacks>
5. Joseph, A. D., Nelson, B., Rubinstein, B. I. P., & Tygar, J. D. (2019). *Adversarial machine learning*. Morgan & Claypool Publishers.

АКТУАЛЬНІ ПИТАННЯ ЗАБЕЗПЕЧЕННЯ КІБЕРБЕЗПЕКИ ТА ЗАХИСТУ ІНФОРМАЦІЇ

Матеріали

*XI Міжнародної науково-практичної конференції «Актуальні питання
забезпечення кібербезпеки та захисту інформації»
24 квітня 2025 р.*

Підписано до друку 16.05.2025. Формат 60x84¹/₁₆.

Папір офсетний. Гарнітура Times New Roman.

Ум. друк. арк. 14.

Зам. № 83.

Друк: поліграфкомбінат Європейського університету.

03115, Україна, Київ-115, вул. Депутатська, 15/17.

Тел.: +38(044) 503-33-96

Реєстраційне свідоцтво ДК №3833 від 14.07.2010 р.