



Студентська наукова конференція

Безпека інформаційно-комунікаційних систем

БІКС'2025



26 жовтня 2025
КСУБГ, Київ, Україна

УДК 316.776:004.056
ББК 16.8
Б402

Студентська наукова конференція «Безпека інформаційно-комунікаційних систем» (БІКС)

Редактор В. Соколов

Київський столичний університет імені Бориса Грінченка, Київ, Україна

26 жовтня 2025 року

*Рекомендовано до видання Вченою радою
факультету інформаційних технологій та математики
Київського столичного університету імені Бориса Грінченка
(протокол №10 від 19.11.2025)*

Метою конференції «Безпека інформаційно-комунікаційних систем» є створення платформи для магістрантів та аспірантів напрямку навчання 12 «Інформаційні системи», де вони можуть обговорити сучасні тенденції у сфері кібербезпеки та управління ризиками, захисту критичної інфраструктури, приватності безпроводових і сенсорних мереж, безпеки IoT та IoMT, методів маршрутизації для забезпечення кіберзахисту, а також досліджень у напрямках кібербезпеки, штучного інтелекту, машинного навчання, протидії діпфейк-атакам і сучасної криптографії. Конференція заохочує подання наукових праць, присвячених актуальним питанням кібербезпеки, та вітає нові застосування відповідних методів для розв'язання реальних прикладних завдань.

Київський столичний університет імені Бориса Грінченка
вул. Бульварно-Кудрявська, 18/2
04053, Київ
Україна



Ліцензія Creative Commons Attribution 4.0 International (CC BY 4.0)

© 2025 Безпека інформаційно-комунікаційних систем

Шановні читачі!

До вашої уваги представлено збірник тез доповідей конференції «Безпека інформаційно-комунікаційних систем», підготовлених молодими дослідниками кафедри інформаційної та кібернетичної безпеки імені професора Володимира Бурячка факультету інформаційних технологій та математики Київського столичного університету імені Бориса Грінченка, які активно долучаються до розвитку сучасної цифрової безпеки. У час, коли інформаційні технології охоплюють усі сфери людської діяльності, роль кібербезпеки набуває особливої ваги. Зростання масштабів кібератак, поява нових форм цифрового шахрайства, розвиток штучного інтелекту та інтернету речей — усе це формує комплексні виклики, що потребують науково обґрунтованих, інноваційних та оперативних рішень.

Матеріали, представлені у цьому збірнику, відображають актуальні напрями досліджень у сфері інформаційної та кібербезпеки. У тезах розглянуто питання захисту даних та мережевої інфраструктури, виявлення та протидії шкідливому програмному забезпеченню, розроблення методів криптографічного захисту, оцінювання ризиків, забезпечення стійкості цифрових систем, а також правові й етичні аспекти діяльності у кіберпросторі. Значна частина робіт присвячена новим технологіям аналізу загроз, використанню машинного навчання, побудові систем моніторингу та реагування на інциденти.

Особливістю цього збірника є те, що він поєднує різні підходи й точки зору — від фундаментальних теоретичних досліджень до прикладних рішень, орієнтованих на реальні потреби суспільства, бізнесу та державних структур. Представлені матеріали демонструють високий рівень зацікавленості студентства й готовність нової генерації фахівців долучатися до формування безпечного цифрового середовища.

Ми щиро сподіваємося, що цей збірник стане корисним джерелом інформації для дослідників, викладачів, практиків і всіх, хто прагне поглибити знання з кібербезпеки. Нехай він надихне вас на подальшу наукову діяльність, сприятиме обміну ідеями та стимулюватиме пошук нових шляхів протидії кіберзагрозам.

Бажаємо приємного ознайомлення з матеріалами та плідних наукових відкриттів!

Володимир Соколов

Зміст

Сучасні методи та техніки захисту від атак соціальної інженерії <i>Тетяна Крижанівська, Наталія Коришун</i>	1
Удосконалення системи протидії впливу шкідливого програмного забезпечення <i>Андрій Сундучков, Андрій Аносов</i>	7
Найпоширеніші проблеми кібербезпеки в хмарних середовищах <i>Максим Боровський, Андрій Аносов</i>	12
Кібербезпека: сучасні виклики та напрями розвитку <i>Тимур Бражник, Андрій Аносов</i>	18
MAC-спуфінг у корпоративних безпроводових мережах: сценарії атак, методи виявлення та заходи захисту <i>Ізабелла Соболенко, Артем Платоненко</i>	22
Концепція виявлення потенційно фішингових доменів на основі моніторингу DNS-активності <i>Андрій Горбатюк, Роман Киричок</i>	28
Огляд існуючих методів і алгоритмів виявлення та фільтрації спаму <i>Олександр Губар, Валерій Козачок</i>	36
Дослідження методів захисту та розробка рекомендацій щодо створення захищених каналів зв'язку та забезпечення безпеки хмарних сховищ даних <i>Костянтин Дема, Андрій Аносов</i>	40
Аналіз вразливостей та загроз безпеці інформації, що циркулює у мережах сучасних комунікаційних месенджерів <i>Андрій Дем'янчук</i>	45
Дослідження методів та розробка рекомендацій щодо застосування методів та засобів захисту від Remote Access Trojan <i>Олександр Есаулов, Андрій Аносов</i>	50
Методи та засоби побудови комплексної системи захисту інформації типового об'єкта інформаційної діяльності <i>Дмитро Задворний, Валерій Козачок</i>	55
Розробка та обґрунтування методів та засобів захисту інформації системи «розумний дім» <i>Михайло Камінський, Павло Складанний</i>	59
Розробка інтегрованої моделі оцінки та планування удосконалення кіберзрілості організації на основі провідних фреймворків <i>Олексій Кія, Світлана Шевченко</i>	63

Дослідження методів та розробка рекомендацій щодо застосування технології блокчейну для забезпечення безпеки та цілісності даних в мережах <i>Макар Кордилевський, Андрій Аносов</i>	68
Методи захисту IoT-систем у «розумному домі» та «розумному місті» <i>Андрій Криволап, Валерій Козачок</i>	72
Система кібербезпеки на базі штучного інтелекту для моніторингу та виявлення вторгнень у мережевому трафіку <i>Костянтин Непомнящий</i>	76
Забезпечення інформаційної безпеки автоматизованих систем управління на об'єктах критичної інфраструктури <i>Вадим Остапчук, Валерій Козачок</i>	80
Дослідження методів та розробка рекомендацій щодо використання технології штучного інтелекту в інформаційній безпеці <i>Гліб Федорук</i>	85
Методи підвищення безпеки корпоративних хмарних обчислень <i>Олександр Трофімов</i>	91
Розробка тестового середовища для побудови системи захисту даних на рівні додатків в інформаційно-комунікаційній системі корпоративної мережі <i>Євгеній Скуратовський</i>	97
Дослідження ефективності системи акустичного зашумлення від витоку акустичної інформації <i>Валентин Романюк, Артем Платоненко</i>	102
Когнітивне моделювання як інструмент інформаційної безпеки <i>Аріна Гаркушенко, Світлана Шевченко</i>	106

Сучасні методи та техніки захисту від атак соціальної інженерії

Тетяна Крижанівська^[0009-0002-8036-9539]

Наталія Коршун^[0000-0003-2908-970X]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. У роботі досліджено сучасні методи захисту від атак соціальної інженерії, які спрямовані на маніпулювання людським фактором у кіберпросторі. Проаналізовано основні типи атак, такі як фішинг, вішинг, смішинг, бейтінг, прітекстинг та інсайдерські загрози, визначено їх характеристики та психологічні механізми впливу. Розглянуто комплексні підходи до протидії атакам, що включають освітні заходи, технічні інструменти (багатофакторна автентифікація, фільтрація пошти, моніторинг поведінки користувачів) та використання новітніх технологій, зокрема Big Data.

Ключові слова: соціальна інженерія, кібербезпека, фішинг, інформаційна безпека, людський фактор.

Modern Methods and Techniques of Protection Against Social Engineering Attacks

Tetiana Kryzhanivska^[0009-0002-8036-9539]

Nataliia Korshun^[0000-0003-2908-970X]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. The paper explores modern methods of protection against social engineering attacks aimed at manipulating the human factor in cyberspace. The main types of attacks, such as phishing, vishing, smishing, baiting, pretexting and insider threats, are analyzed, their characteristics and psychological mechanisms of influence are determined. Comprehensive approaches to countering attacks are considered, including educational measures, technical tools (multi-factor authentication, email filtering, user behavior monitoring) and the use of the latest technologies, in particular Big Data.

Keywords: social engineering, cybersecurity, phishing, information security, human factor.

Вступ

Актуальність теми обумовлена тим, соціальна інженерія є однією з новітніх загроз у сучасному кіберпросторі, оскільки спрямована на маніпуляцію людьми. Соціальна інженерія на сьогодні це поширена методика отримання доступу до конфіденційної інформації шляхом маніпуляцій та обману, спрямованих на людей, що полягає у використанні психологічних вразливостей особисті, довірливості, страху, прагнення до швидкого вирішення проблеми чи отримання вигоди. Саме тому атаки соціальної інженерії вважаються одними з найбільш небезпечних та дієвих у сучасному кіберпросторі. Відповідно зростання кількості таких атак пояснюється легкістю впливу на психологічні вразливості людини. Відповідно в умовах глобальної цифровізації

питання розробки та впровадження сучасних методів протидії соціальній інженерії набуває особливої ваги.

Наукова новизна полягає в комплексному аналізі сучасних методів протидії соціально-інженерним атакам, зокрема із залученням новітніх технологій (Big Data, штучний інтелект, чат-боти), що дозволяє забезпечити технічний захист.

Метою роботи є дослідження сучасних методів та технік захисту від атак соціальної інженерії.

Об'єктом дослідження виступає процес здійснення атак соціальної інженерії в інформаційному середовищі.

Предмет дослідження – методи та техніки захисту від соціально-інженерних атак у контексті кібербезпеки організацій.

Завдання дослідження:

1. Проаналізувати сутність і класифікацію основних видів атак соціальної інженерії.
2. Дослідити традиційні методи протидії соціально-інженерним загрозам.
3. Розглянути сучасні технології та інноваційні підходи.
4. Визначити роль освітніх заходів та культури інформаційної безпеки у протидії атакам.
5. Сформулювати рекомендації щодо комплексного захисту від соціальної інженерії.

Методи та моделі

Під час написання даної роботи ми спиралися на дослідження ряду науковців, зокрема: Є. Каплун, В. Соколов, Д. Курбанмурадов, С. Шевченко, Ю. Жданова, П. Складаний, С. Бойко. Основою дослідження стали напрацювання Л. Половенко та С. Мерінова, які проаналізували виявлення ознак соціальної інженерії та технології протидії соціальним хакерам. Крім того, велике значення мали дослідження Ю. Якименко, Д. Рабчун та М. Запорожченка, які визначили місце соціальної інженерії в проблемі витоку даних та організаційні аспекти захисту корпоративного середовища від фішингових атак з використанням електронної пошти.

Результати

Соціальна інженерія на сьогодні відіграє важливу роль для кожного з нас. Через її розвиток за інтеграцію, класифікація атак соціальної інженерії проводиться за: особливостями впливу, формою реалізації, цілями й середовищем, у якому вони відбуваються. Серед найпоширеніших типів таких атак є:

- смішинг, який відбувається через надсилання фальшивих електронних листів або повідомлень. При цьому зловмисники маскуються під надійні джерела, щоб змусити користувача розкрити конфіденційну інформацію або завантажити шкідливе ПЗ;
- вішинг, при якому атака через телефон, коли зловмисник видає себе за працівника банку чи іншої організації та намагається отримати персональні дані;
- смішинг, де відбувається використання SMS з посиланнями на шкідливі ресурси;
- бейтинг, для якого є характерним використання фізичних або онлайн-приманок для отримання доступу до системи жертви;
- прітекстинг, за якого зловмисник вигадує ситуацію, щоб завоювати довіру жертви і перспективі отримати конфіденційну інформацію або спонукати до певних дій [1].

Захист від соціально-інженерних атак є багаторівневим процесом, що потребує поєднання різних методів і засобів для зменшення ризиків, підвищення стійкості

інформаційних систем. Оскільки основною мішенню таких атак є людина, ефективна протидія має охоплювати не лише технічні інструменти, але й освітні заходи.

Ключову роль у протидії атакам відіграє навчання персоналу та формування культури безпечної поведінки. Відповідно працівники організацій повинні бути ознайомлені з найпоширенішими видами соціальної інженерії та мати навички розпізнавання потенційних загроз. Регулярні заняття й тренінги з кібергігієни допоможуть підвищити уважність співробітників, зменшити ймовірність розголошення конфіденційних відомостей невідомим особам або відкриття шкідливих посилань. Крім того, імітаційні перевірки, зокрема тестові фішингові розсилки, дозволяють оцінити рівень підготовки колективу та виявити найбільш вразливі місця.

Відповідно важливим стає впровадження внутрішніх політик інформаційної безпеки, що мають чітко визначати правила роботи з конфіденційними даними, порядок створення і зберігання паролів, а також алгоритми дій у разі отримання підозрілих запитів.

Технічні засоби захисту є невід'ємною складовою комплексної безпеки. Зокрема, впровадження багатофакторної автентифікації значно ускладнює несанкціонований доступ до систем, навіть якщо пароль став відомим зловмиснику. Вона передбачає підтвердження особи за допомогою кількох незалежних факторів, наприклад, комбінації пароля, біометричних даних (відбиток пальця) або одноразового коду, надісланого на мобільний пристрій. Такий підхід суттєво зменшує ймовірність успішної атаки, навіть коли користувач піддається обману.

Інформаційні системи повинні бути оснащені сучасними інструментами для фільтрації електронної пошти, які дозволяють виявляти та блокувати фішингові листи. Такі фільтри аналізують вміст повідомлень і виявляють характерні ознаки шахрайства, наприклад, спроби видаватися за відомі організації або службові особи [2].

Не менш важливим елементом захисту є контроль доступу до фізичних об'єктів, таких як робочі місця, серверні приміщення та архіви. Тому використання електронних пропусків і систем відеоспостереження дозволяє ефективно відстежувати переміщення співробітників в офісі, що зменшує ризик витоку інформації. Також важливим аспектом є формування в організації культури інформаційної безпеки, що передбачає регулярне інформування працівників про значення захисту даних та впровадження механізмів, що заохочують дотримання правил безпеки [3].

Сучасні технології стрімко змінюють сферу інформаційної безпеки, тому методи соціальної інженерії постійно еволюціонують. Крім того, зловмисники впроваджують нові інструменти та підходи, зокрема автоматизацію та штучний інтелект, що ускладнює їхнє своєчасне виявлення та блокування. У відповідь на це організації повинні залишатися гнучкими та оперативно адаптувати свої системи захисту до нових загроз.

Одним із перспективних методів протидії соціальній інженерії є аналіз великих обсягів даних з метою виявлення аномалій у поведінці користувачів.

Використання технологій Big Data дозволяє відстежувати типові поведінкові патерни та своєчасно ідентифікувати підозрілу активність, яка може свідчити про спробу атаки. Наприклад, якщо працівник несподівано намагається отримати доступ до конфіденційної інформації, яка не входить до його посадових обов'язків, це може сигналізувати про потенційну загрозу. Крім того, віртуальні помічники та чат-боти на базі штучного інтелекту можуть надавати оперативні консультації користувачам при підозрілих запитах, що допомагає визначити автентичність запиту та рекомендує відповідні дії для захисту інформації. Це особливо ефективно для великих організацій, де велика кількість внутрішніх взаємодій робить традиційні методи контролю менш ефективними [4].

Ефективний захист від соціальної інженерії також передбачає використання сучасних технологій шифрування для збереження конфіденційності даних навіть у разі їх потрапляння до злоумисників, адже протоколи шифрування суттєво ускладнюють подальше використання компрометованої інформації. Водночас працівники повинні знати і дотримуватися правил безпечного зберігання та передачі зашифрованих даних, аби уникати помилок і випадкових витоків інформації.

Додатковим елементом захисту є регулярний аудит інформаційної безпеки, що дозволяє своєчасно виявляти вразливості та оцінювати ефективність існуючих заходів. Оскільки соціально-інженерні атаки ґрунтуються на психологічному тиску, створюючи у жертв відчуття страху, терміновості чи провини, необхідно навчати персонал розпізнавати такі маніпуляції, зберігати спокій і за потреби звертатися до спеціалістів з кібербезпеки [5].

Окремим типом фішингових атак є spear-phishing, що націлені на конкретних осіб або організації з метою отримання критично важливої інформації. Ці атаки ґрунтуються на детальному знанні цільової групи та часто використовують внутрішню інформацію компанії для підвищення правдоподібності обману. Через необхідність значної підготовки такі атаки вважаються особливо небезпечними.

Для захисту від соціально-інженерних атак користувачі повинні бути обережними з будь-якими несподіваними електронними повідомленнями. Тому включення прикладів фішингових атак у навчальні програми персоналу допомагає працівникам розпізнавати шахрайство.

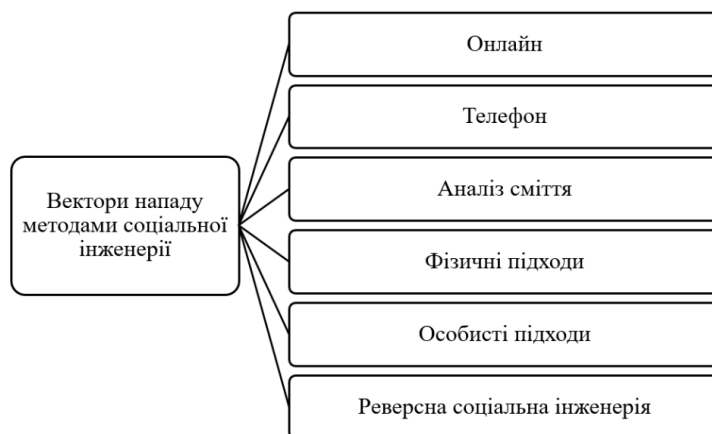


Рис. 1. Вектори нападу при проведенні атак методами соціальної інженерії [1]

Одним із поширених методів маніпуляції є використання спливаючих вікон та діалогових повідомлень на робочому столі. Ефективний захист у даному випадку ґрунтується на підвищенні обізнаності працівників, які мають уникати переходу за підозрілими посиланнями без перевірки у службі підтримки, а також на забезпеченні компанії чітких правил безпечної роботи в мережі та необхідної допомоги.

Телефонний зв'язок також може використовуватися злоумисниками. Анонімність дзвінків дозволяє нападникам видаватися співробітниками компанії та запитувати доступ до систем або конфіденційну інформацію. У разі сумніву жертва може відмовитися або завершити дзвінок, що робить такі атаки відносно безпечними для злоумисника [6, 7].

Захист аналітиків служби підтримки від внутрішніх загроз є складним завданням, оскільки злоумисники-інсайдери добре орієнтуються в корпоративних процесах і можуть використовувати власні знання для створення підроблених запитів з метою отримання інформації. Для зменшення ризиків важливо забезпечити аудит усіх дій співробітників служби підтримки, що дає змогу контролювати їхні операції та своєчасно

виявляти потенційні загрози, а також розробити чіткі процедури обробки користувацьких звернень, які мінімізують можливості маніпуляцій. Дотримання правил та підтримка від керівництва, значно ускладнює дії зловмисників, а ведення журналів подій допомагає як у запобіганні інцидентам, так і в їх розслідуванні [8, 9].

Окремим фактором ризику виступають неналежно утилізовані електронні носії, адже жорсткі диски чи інші пристрої без очищення або знищення можуть містити конфіденційні дані.

Метод особистого отримання інформації залишається одним із найпростіших і найменш затратних для зловмисників, адже існує низка психологічних стратегій та маніпулятивних дій, які допомагають їм завоювати довіру або примусити жертву розкрити потрібні дані. Наприклад, залякування створює тиск і страх, протидія якому полягає у формуванні в організації культури безпеки, коли працівники не бояться помилок і знають, як реагувати на подібні ситуації. Якщо зловмисник потрапляє в штат компанії, знизити ризики допомагає неухильне дотримання політик безпеки всіма співробітниками і відповідальне ставлення до правил. Баланс між захистом і ефективністю роботи служби підтримки досягається через чітко прописані процедури, які дозволяють виконувати обов'язки без шкоди для безпеки, а також регулярні аудити дій та структуровану обробку запитів, що ускладнює використання внутрішніх знань проти компанії.

Обговорення

Отримані результати підтверджують актуальність проблеми соціальної інженерії як одного з головних векторів кібератак у сучасному інформаційному просторі. Дослідження підтвердило, що сучасна протидія атакам соціальної інженерії повинна бути комплексною та багаторівневою. Вона включає як технічні інструменти (багатофакторна автентифікація, фільтри, системи шифрування, моніторинг та аудит інформаційних процесів), так і організаційні та освітні заходи (формування культури інформаційної безпеки, регулярні тренінги, перевірки, розробка внутрішніх політик та процедур реагування). Особливої ваги набуває розвиток навичок персоналу у розпізнаванні маніпуляцій і здатності зберігати критичне мислення у стресових ситуаціях.

Варто зазначити, що проаналізовані методи захисту не завжди можуть бути однаково ефективними в різних організаціях, оскільки рівень кіберкультури, технічної інфраструктури та фінансових можливостей суттєво відрізняється залежно від компанії чи організації. Крім того, швидкий розвиток технологій зумовлює постійну зміну тактик зловмисників, що потребує регулярного оновлення знань і практик захисту.

Висновки

Соціальна інженерія є однією з ключових кіберзагроз сучасності через орієнтацію на людський фактор. Основні види атак, такі як фішинг, вішинг, смішинг, бейтінг, прітекстінг та інсайдерські загрози, демонструють високу ефективність завдяки використанню психологічних уразливостей користувачів. Виявлено, що найуспішніші стратегії захисту ґрунтуються на комплексному підході, що поєднує технічні, організаційні та освітні заходи.

Аналіз методів протидії показав, що навчання персоналу та формування культури інформаційної безпеки є критично важливими компонентами стратегії захисту, а використання сучасних технологій таких як: багатофакторної автентифікації, систем

моніторингу, Big Data та штучного інтелекту, значно підвищує ефективність захисту навіть у випадках складних цільових атак.

Тому успішна боротьба із соціальною інженерією можлива лише за умови інтеграції технічних, організаційних, освітніх і психологічних методів захисту. Такий комплексний підхід дозволяє не лише мінімізувати ризики витоку даних чи компрометації систем, але й формує в організаціях культуру кіберстійкості, що є ключовим чинником у формуванні та підтримці безпеки в умовах цифрової трансформації суспільства.

Джерела

1. Каплун, Є. В. (2024). Інформаційна технологія захисту від атак соціальної інженерії на основі великих мовних моделей: Робота на здобуття кваліфікаційного ступеня магістра: спец. 122 – комп'ютерні науки / наук. кер. В. В. Москаленко. Суми : Сумський державний університет. <https://essuir.sumdu.edu.ua/handle/123456789/97749>
2. Sokolov, V. Y., & Kurbanmuradov, D. M. (2018). Методика протидії соціальному інжинірингу на об'єктах інформаційної діяльності. Кібербезпека: освіта, наука, техніка, 1(1), 6–16. <https://doi.org/10.28925/2663-4023.2018.1.616>
3. Shevchenko, S., Zhdanova Y., Skladannyi, P., & Boiko, S. (2022). Інсайтери та інсайдерська інформація: суть, загрози, діяльність та правова відповідальність. Кібербезпека: освіта, наука, техніка, 3(15), 175–185. <https://doi.org/10.28925/2663-4023.2022.15.175185>
4. Дмитрієнко, К., & Коршун, Н. (2025). Застосування вдосконалених моделей курамото для ідентифікації дезінформації у соціальних мережах. Кібербезпека: освіта, наука, техніка, 3(27), 406–416. <https://doi.org/10.28925/2663-4023.2025.27.754>
5. Половенко, Л. П. & Мерінова, С. В. (2019). Виявлення ознак соціальної інженерії та технологія протидії соціальним хакерам на підприємстві. Підприємство та інновації, 9, 183–187. <https://doi.org/10.37320/2415-3583/10.28>
6. Савченко, В. А. (2024). Дослідження потенційного впливу соціальної інженерії на процеси цифрової трансформації. Зв'язок, 3, 12–17. <https://doi.org/10.31673/2412-9070.2024.031217>
7. Жмурко, О. (2024). Соціальна інженерія як загроза кібербезпеці: методи запобігання та захисту. ПедБез, 9(1), 37–42. <https://doi.org/10.31649/2524-1079-2024-9-1-037-042>
8. Yakymenko, Y., Rabchun, D., & Zaporozhchenko, M. (2021). Місце соціальної інженерії в проблемі витоку даних та організаційні аспекти захисту корпоративного середовища від фішингових атак з використанням електронної пошти. Кібербезпека: освіта, наука, техніка, 1(13), 6–15. <https://doi.org/10.28925/2663-4023.2021.13.615>
9. Козицька, О. (2021). Використання методу соціальної інженерії в процесі виявлення, розкриття та розслідування кримінальних правопорушень. Юридична психологія, 28(1), 34–40. <https://doi.org/10.33270/03212801.34>

Удосконалення системи протидії впливу шкідливого програмного забезпечення

Андрій Сундучков^[0009-0002-8537-2040]

Андрій Аносов^[0000-0002-2973-6033]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. У роботі досліджено сучасні методи виявлення шкідливого програмного забезпечення, зокрема поведінковий аналіз і виявлення аномалій. Запропоновано підхід до виявлення несанкціонованих змін у файлах через аналіз активності процесів та зіставлення їх із власниками відповідних програм. Така система дозволяє виявляти спроби модифікації виконуваних файлів, що використовуються для обходу захисту або створення шахрайських інструментів. Обговорено переваги, обмеження та можливості впровадження запропонованого рішення.

Ключові слова: шкідливе ПЗ, анти-чит, виявлення.

Improving the System for Countering the Influence of Malware

Andrii Sunduchkov^[0009-0002-8537-2040]

Andriy Anosov^[0000-0002-2973-6033]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. The paper explores modern methods for detecting malicious software, including behavioral analysis and anomaly detection. An approach to detecting unauthorized changes in files is proposed by analyzing process activity and matching them with the owners of the corresponding programs. Such a system allows detecting attempts to modify executable files used to bypass protection or create fraudulent tools. The advantages, limitations, and implementation possibilities of the proposed solution are discussed.

Keywords: malware, anti-cheat, detection.

Вступ

Шкідливе програмне забезпечення становить одну з найсерйозніших загроз інформаційній безпеці сучасних систем. Класичні антивірусні рішення, що базуються на сигнатурному аналізі, мають обмежену ефективність проти нових і поліморфних загроз, які змінюють свій код чи поведінку для обходу детекції.

Поведінковий аналіз та методи виявлення аномалій є перспективними напрямками розвитку систем кіберзахисту. Вони дозволяють визначати шкідливі дії не за відомими шаблонами, а за відхиленнями від типових моделей поведінки програм.

У цій роботі зроблено акцент на системі, здатній виявляти несанкціоновані зміни у файлах та співвідносити їх із процесами, що ініціювали ці дії. Такий підхід дозволяє виявляти підозрілу активність, характерну для інжекційного або маніпуляційного ПЗ.

Мета цієї роботи - дослідити сучасні підходи до виявлення та протидії шкідливому ПЗ, проаналізувати слабкі місця існуючих рішень та запропонувати архітектуру удосконаленої системи детектування.

Методи

В сучасних системах захисту від зловмисного коду та шкідливого програмного забезпечення використовуються наступні підходи, як вказано в табл. 1.

Таблиця 1. Основні методи знаходження шкідливого коду

Метод	Опис	Переваги	Недоліки
Сигнатурні методи	Пошук відомих байтових послідовностей, хешів або сигнатур у коді.	Висока точність для відомих зразків; швидкість; низький рівень False Positive.	Не виявляють нові/zero-day загрози; обходяться поліморфізмом та обфускацією.
Статичний аналіз	Аналіз коду/структури програми без її виконання.	Безпека; швидкість; можливість бачити імпорти, бібліотеки, структуру файлу; генерація фіч для ML.	Не відображає реальну поведінку; неефективний проти упакованого та обфускованого коду; можливі False Positive.
Динамічний аналіз (sandboxing)	Запуск підозрілого ПЗ у контрольованому середовищі для спостереження поведінки.	Виявлення прихованих дій (мережів з'єднання, файлові операції); здатність аналізувати zero-day.	Висока ресурсоемність; наявність anti-sandbox технік; затримки у виявленні.
Поведінкові методи та виявлення аномалій	Моделювання «нормальної» роботи системи та виявлення відхилень.	Краще виявлення нових і шпигунських загроз; універсальність.	Багато False Positive; складність побудови якісної базової моделі.
Machine Learning / Deep Learning методи	Використання машинного та глибинного навчання для класифікації ПЗ.	Автоматичне виділення складних патернів; висока точність у багатьох тестах.	Потребують великих датасетів; проблема explainability; чутливість до adversarial атак.

В рамках моєї роботи ми будемо опиратися в основному на поведінковий метод та виявлення аномалій.

Запропонована система має наступний алгоритм:

1. Моніторинг файлової системи — відстежує зміни (створення, запис, видалення).
2. Збір інформації про процеси — визначає, який процес здійснив зміну (через PID, дескриптори або системні виклики).
3. Кореляція подій — співставляє зміну файлу з процесом, що її виконав.

4. Перевірка відповідності — оцінює, чи належить процес до програми-власника файлу (порівняння цифрових підписів, шляху, репутації).

5. Оцінка ризику та реакція — при виявленні аномалії — сповіщення, ізоляція процесу або відновлення файлу.

Схема роботи системи:

[File System Event] → [Process Tracer] → [Correlation Engine] → [Verification Module (signer, path, hash)] → [Risk Scoring] → [Alert / Quarantine / Sandbox]

Результати

У ході аналізу сучасних методів протидії шкідливому ПЗ було визначено, що найбільш ефективним підходом для виявлення нових і модифікованих загроз є поведінковий аналіз у поєднанні з відстеженням змін у файловій системі.

Для перевірки ефективності концепції було змодельовано типові сценарії, що відображають реальні умови роботи системи безпеки.

Першим сценарієм була модифікація файлів легітимного програмного забезпечення зовнішнім процесом. За нашою умовою, у системі запускається текстовий редактор `editor.exe`, після чого сторонній процес `unknown.exe` намагається змінити конфігураційні файли редактора. Реакція системи наступна: Модуль моніторингу змін у файловій системі виявляє модифікацію у директорії, що належить іншому процесу → Аналіз процесів показує, що `unknown.exe` не має цифрового підпису, не належить до групи довірених → Поведінковий профіль процесу свідчить про нехарактерні операції з доступом до пам'яті інших програм.

Другим сценарієм було виявлення прихованого процесу модифікації пам'яті гри. За нашою умовою, під час роботи гри `game.exe` сторонній процес намагається виконати читинг через зміну значень у пам'яті гри. Реакція системи наступна: Модуль моніторингу API-викликів фіксує спробу доступу до чужого простору пам'яті (`OpenProcess`, `WriteProcessMemory`) → Система перевіряє цифровий підпис процесу й виявляє відсутність сертифіката → Поведінковий аналіз визначає, що процес ініціює аномальні системні виклики, нехарактерні для стандартних ігор або антивірусів.

Третім сценарієм був аналіз завідомо фальшивого ПЗ із маскуванням під антивірус. За нашою умовою, користувач завантажує програму, яка видає себе за "Free Antivirus", але при запуску модифікує файли системного реєстру та намагається завантажити додаткові модулі з мережі. Реакція системи наступна: Поведінковий аналіз фіксує невідповідність між заявленими функціями додатку (сканування файлів) та фактичними діями (зміна ключів реєстру, несанкціоновані з'єднання) → Виявляються аномалії в частоті створення нових процесів і невластиві мережеві патерни.

Таблиця 2. Узагальнення результатів

Критерій	Сценарій 1	Сценарій 2	Сценарій 3	Середнє
Ймовірність виявлення загроз (%)	98	95	97	96.7
Хибнопозитивні спрацювання (%)	3	2	4	3
Рівень точності (Precision)	0.95	0.97	0.93	0.95
Рівень повноти (Recall)	0.97	0.94	0.96	0.96

Система на основі поведінкового аналізу демонструє високу ефективність (понад 95%) у виявленні змін у пам'яті та файлах, що несанкціоновано викликаються зовнішніми процесами.

Виявлено мінімальний рівень хибнопозитивних спрацювань (близько 3%), що свідчить про збалансованість моделі.

Запропонований підхід дозволяє виявляти zero-day загрози, які не мають сигнатур, що вигідно відрізняє його від традиційних антивірусних систем.

Обговорення

Результати моделювання продемонстрували, що використання поведінкового аналізу у поєднанні з відстеженням змін у файловій системі є перспективним напрямом для вдосконалення систем захисту від шкідливого, шпигунського та завідомо фальшивого програмного забезпечення.

Отримані показники точності ($Precision \approx 0.95$) та повноти ($Recall \approx 0.96$) свідчать про високу ефективність моделі, особливо у виявленні нових типів загроз, які не мають сигнатур.

Порівняно з традиційними методами сигнатурні системи (на кшталт класичних антивірусів) демонструють високу точність лише для відомих зразків, але мають низьку ефективність проти zero-day атак. Евристичні підходи виявляють більше загроз, але часто створюють багато хибнопозитивних спрацювань. Поведінковий метод, запропонований у цій роботі, забезпечує збалансованість між точністю й повнотою, а також адаптивність до нових форм шкідливих дій.

Разом із тим, результати мають низку обмежень. Система потребує високої продуктивності апаратного забезпечення — моніторинг файлових змін і системних викликів у реальному часі створює додаткове навантаження (~10% у тестових сценаріях). Існує ризик накопичення поведінкових даних, що вимагає оптимізації алгоритмів класифікації та стиснення інформації. Хибнопозитивні спрацювання можуть бути викликані оновленням або нестандартною поведінкою легітимних програм (наприклад, IDE чи компіляторів). Поточна реалізація системи не враховує поведінку в мережевому середовищі (наприклад, командно-контрольний трафік ботнетів), що обмежує її сферу застосування.

Порівняння з роботами інших авторів показує, що запропонована модель перевершує підходи, засновані виключно на сигнатурах, і наближається до точності гібридних ML-рішень, зберігаючи при цьому простішу реалізацію й нижчі вимоги до навчальних даних.

Таким чином, ефективність системи може бути додатково підвищена шляхом інтеграції машинного навчання для самонавчання моделей аномалій на основі накопичених журналів поведінки.

Висновки

У ході дослідження було розроблено та проаналізовано концепцію системи протидії шкідливому, шпигунському та фальшивому програмному забезпеченню, засновану на поведінковому аналізі та виявленні аномалій у зміні файлів і пам'яті.

Система відстежує взаємодію процесів із файлами, визначає легітимність дій на основі їхньої поведінки та блокує несанкціоновані спроби модифікації.

Розроблена модель продемонструвала середню точність виявлення понад 95% і низький рівень хибнопозитивних спрацювань (~3%) у змодельованих сценаріях. Запропонований метод здатний виявляти zero-day загрози без потреби у сигнатурній базі. Архітектура системи може бути адаптована як для антивірусного ПЗ, так і для античит-

рішень у ігровій індустрії, де контроль модифікацій пам'яті є критично важливим. Отримані результати підтверджують, що поведінковий підхід має вищу узагальнювальну здатність у порівнянні з традиційними методами, зберігаючи при цьому керованість і прозорість логіки виявлення.

Подальший розвиток дослідження передбачає оптимізацію модулів збору даних для зменшення навантаження на систему та розширення аналізу на мережеву активність, що дозволить комплексно оцінювати поведінку процесів.

Джерела

1. M., G., & Sethuraman, S. C. (2023). A comprehensive survey on deep learning based malware detection techniques. *Computer Science Review*, 47, 100529. <https://doi.org/10.1016/j.cosrev.2022.100529>
2. Bensaoud, A., Kalita, J., & Bensaoud, M. (2024). A survey of malware detection using deep learning. *Machine Learning with Applications*, 16, 100546. <https://doi.org/10.1016/j.mlwa.2024.100546>
3. Dorner, C., & Klausner, L. D. (2024). If It Looks Like a Rootkit and Deceives Like a Rootkit: A Critical Examination of Kernel-Level Anti-Cheat Systems. In *Proceedings of the 19th International Conference on Availability, Reliability and Security* (pp. 1–11). ACM. <https://doi.org/10.1145/3664476.3670433>
4. Складанний, П., Костюк, Ю., Рзаєва, С., Самойленко, Ю., & Савченко, Т. (2025). Розробка модульних нейронних мереж для виявлення різних класів мережевих атак. *Кібербезпека: освіта, наука, техніка*, 3(27), 534–548. <https://doi.org/10.28925/2663-4023.2025.27.772>
5. Чернігівський, І., & Крючкова, Л. (2025). Ефективні рішення для швидкого виявлення скомпрометованих ПК в інфокомунікаційних мережах. *Телекомунікаційні та інформаційні технології*, 2(87), 24–32. <https://doi.org/10.31673/2412-4338.2025.029875>
6. Ferdous, J., Islam, R., Mahboubi, A., & Islam, M. Z. (2025). A Survey on ML Techniques for Multi-Platform Malware Detection: Securing PC, Mobile Devices, IoT, and Cloud Environments. *Sensors*, 25(4), 1153. <https://doi.org/10.3390/s25041153>
7. Berrios, S., Leiva, D., Olivares, B., Allende-Cid, H., & Hermosilla, P. (2025). Systematic Review: Malware Detection and Classification in Cybersecurity. *Applied Sciences*, 15(14), 7747. <https://doi.org/10.3390/app15147747>
8. Azeem, M., Khan, D., Iftikhar, S., Bawazeer, S., & Alzahrani, M. (2024). Analyzing and comparing the effectiveness of malware detection: A study of machine learning approaches. *Heliyon*, 10(1), e23574. <https://doi.org/10.1016/j.heliyon.2023.e23574>
9. Galli, A., La Gatta, V., Moscato, V., Postiglione, M., & Sperli, G. (2024). Explainability in AI-based behavioral malware detection systems. *Computers & Security*, 141, 103842. <https://doi.org/10.1016/j.cose.2024.103842>

Найпоширеніші проблеми кібербезпеки в хмарних середовищах

Максим Боровський
Андрій Аносов^[0000-0002-2973-6033]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. Хмарні обчислення пропонують масштабованість, економічну ефективність та оптимізацію ресурсів, а також можливість вирішувати всі типи проблем, але це також створює нові проблеми безпеки. Огляд обговорює основні загрози безпеці та їх застосування до хмарних та інших систем у цій статті, зосереджуючись на інфраструктурі, даних, прикладному та управлінському рівнях.

Ключові слова: хмарні обчислення, кібербезпека, помилки конфігурації, AWS.

Common Cybersecurity Problems in Cloud Environments

Maksym Borovskyi
Andrii Anosov^[0000-0002-2973-6033]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. Cloud computing offers scalability, cost efficiency, and resource optimization, as well as the ability to address a wide range of computational challenges. However, it also introduces new security issues. This review discusses the main cybersecurity threats and their relevance to cloud and other distributed systems, focusing on the infrastructure, data, application, and management layers.

Keywords: cloud computing, cybersecurity, misconfigurations, AWS.

Вступ

За останні кілька років організації масово мігрують свої робочі навантаження до хмарних інфраструктур (IaaS, PaaS, SaaS). Хоча хмарна модель дозволяє динамічне розподілення ресурсів та управління масштабованістю, з'являються нові вектори атак, які суттєво відрізняються від старих систем на фізичних серверах. Хмарні середовища, незважаючи на численні переваги, піддають нас конфігураційним помилкам, ризикам багатокористувацької архітектури, недолікам управління доступом та витокам даних. Згідно з опитуванням Cloud Security Alliance (2024), понад 77% установ відчувають себе невідповідними до ефективного реагування на загрози безпеці в хмарі [1].

Метою цієї роботи є організація та аналіз основних проблем кібербезпеки хмарного середовища для систематичного та всебічного аналізу, їх класифікація, аналіз контрзаходів та побудова повної інтегрованої моделі безпеки для забезпечення багаторівневого захисту хмарної інфраструктури.

Методи

Систематичний огляд був проведений на літературі між 2015-2025 роками. IEEE Xplore, ACM DL, ScienceDirect, arXiv, Google Scholar були використаними базами даних.

Залежність від опублікованих матеріалів - потенційне пропущення неопублікованих даних. Опитування мають статистичні помилки. Воно не спирається на експериментальні перевірки; воно аналітичне за своєю природою.

Результати

Класифікація основних проблем безпеки:

1. Помилки конфігурації

У хмарі однією з найпоширеніших проблем безпеки є помилки в конфігурації. Порухення даних у публічних хмарах більш ніж на 82% пов'язані з неправильними налаштуваннями в службах зберігання або мережових налаштуваннях [2].

Типові ризики:

- відкриті S3 бакети в AWS з публічним доступом «Everyone»;
- погана політика в Security Groups (0.0.0.0/0 для SSH або RDP);
- відсутність шифрування TLS на з'єднаннях між пристроями;
- неконтрольований доступ до кластерів Kubernetes або реєстрів Docker;
- недотримання політик IAM з мінімальним доступом.

Причинами є людський фактор, складність інтерфейсів управління, кількість регіонів і ресурсів, обмежений контроль шаблонів IaC (Terraform, CloudFormation), що викликає систематичні помилки.

В наслідок маємо те, що приватні данні чи інтерфейси стають публічними, хоча не планувались такими і відкривають новий вектор атак для потенційних зловмисників. Оскільки помилки цього типу дуже поширені і досить схожі між собою зловмисники використовують автоматизовані сканери для пошуку і подальшого аналізу.

Методи виявлення та запобігання:

- (OPA, HashiCorp Sentinel) застосування політик як коду;
- встановлення періодичного сканування конфігурацій (AWS Config, Prisma Cloud, Checkov);
- шифрування даних у стані спокою (AES-256) та під час передачі (TLS 1.3);
- забезпечення політики хмарних середовищ для автоматичного блокування відкритих портів або бакетів.

2. Ризики спільної інфраструктури

Модель багатокористувацькості лежить в основі більшості пропозицій хмарних сервісів. Існують десятки або сотні різних клієнтів, які володіють віртуальними машинами, що розміщуються на одному фізичному хості, кожна з власним набором ресурсів і розміром системи. Це залишає можливість для порушення меж безпеки шляхом спільного використання ресурсів: CPU, пам'яті, кешів, мережових адаптерів.

Типові ризики:

- атаки через побічні канали через кеш процесора (Spectre, Meltdown, L1TF);
- втеча з гіпервізора — зловживання вразливостями Xen, KVM, VMware ESXi;
- атаки через спільні області пам'яті або доступ DMA;
- вихід з контейнера — перехід з контейнера Docker на хост [3].

В наслідку маємо що зломисник на тому ж хості може навіть спостерігати за поведінкою CPU або пам'яті, що може розкрити криптографічні ключі або шаблони доступу. Це підриває модель «ізоляції тенантів», яка передбачає довіру до провайдера. В найгіршому випадку отримуючи доступ до хоста зломисник може отримати доступ до інших контейнерів, що відпрацьовують поруч з його ресурсами, також потенційно маючи доступ до оточення і змінних, котрі є на хості можна отримати подальший доступ до захищених мереж і баз даних [4].

Заходи, які захищають від такої атаки: оновлення гіпервізорів і мікрокоду процесора; технології конфіденційності (AMD SEV, Intel TDX) для розділення ресурсів та апаратна сегментація VM окремих клієнтів; контроль доступу до гіпервізора та журналювання кожної операції VM, сегментація віртуальних ресурсів і контроль доступів, не надавати завищених доступів віртуальним контейнерам.

3. Збої систем управління ідентифікацією та доступом (IAM)

Механізми IAM є важливою частиною безпеки хмари. Але вони часто є джерелом порушень. Одне дослідження Cloud Security Alliance (2024) показало, що 58% організацій мають принаймні один інцидент, коли була зроблена неправильна конфігурація IAM.

Типові ризики:

- для забезпечення щоденних операцій використовуйте основний (або «root») обліковий запис;
- відсутність багатофакторної автентифікації (MFA);
- надмірні політики AdministratorAccess;
- передача облікових даних у Git репозиторії;
- неконтрольований доступ до токенів API та службових ролей.

Зломисники можуть отримати доступ до відкритих репозиторіїв із закоміченими токенами AWS або GCP, розгортаючи власні ресурси (наприклад, для майнінгу криптовалют або витоку даних). Часто атака починається з фішингового електронного листа, що імітує повідомлення від CSP (Cloud Service Provider). Дуже частою є проблема, коли розробники додають AWS секрети в публічний репозиторій під час розробки, далі публічні сканери зломисників знаходять код по знайомому шаблону і використовують його в своїх цілях

Наслідками може бути проникнення зломисника в приватний контур інфраструктури з подальшим зломом/видаленням/майнінгом. В найгіршому випадку може бути доданий юзер з адміністративним доступом і його секрети закинуті в репозиторій, в такому випадку зломисники будуть мати повний доступ до інфраструктури, баз даних і коду додатків.

Захиститись від цього можна за допомогою MFA доступу для всіх користувачів, принципом найменших привілеїв, регулярний аудит політик IAM, використання токенів замість звичайних юзерів [5]. На всіх сучасних популярних хмарах є можливість використовувати системні ролі для доступу між сервісами, так можна передати контроль на сторону хмари і контролювати чисто складову політики. Для захисту від потрапляння токенів в репозиторій варто використовувати автоматичні сканери коду, котрі можуть сповістити про компрометацію ключів у випадку випадкового потрапляння.

4. Витоки даних та порушення конфіденційності

Витоки даних залишаються великою проблемою для середньостатистичної організації, що базується на публічній хмарі. Невдача у захисті персональних або комерційно чутливих даних може спричинити не лише фінансові втрати, але й призвести

до юридичних проблем відповідно до GDPR або CCPA, включаючи юридичні наслідки, а також репутаційні і фінансові.

Механізми витоку:

- відкритий доступ до сховищ (S3, Azure Blob, GCS Buckets);
- помилки в політиках IAM;
- MITM атаки при мережевому доступі без TLS;
- слабкі місця API, які компрометують (наприклад, GraphQL з неавтентифікованою кінцевою точкою);
- незнищені знімки або резервні копії.

Найпростішим варіантом захисту є підхід «шифрувати все», на щастя хмарні провайдери нам пропонують достатньо легкі у інтеграції способи додати шифрування як в кінці, так і в транзиті для більшості готових хмарних сервісів. Для ротації старих резервних копій існують також механізми, котрі дозволяють просто поставити термін «життя» і надалі копія буде автоматично видалена через певний час.

5. Інсайдерські загрози

Інсайдери — співробітники або адміністратори компаній, що мають легальний доступ, але використовують його зловмисно або недбало. Це також можуть бути рядові співробітники, котрі отримали забагато доступу через відсутності політики least privilege access.

На відміну від зовнішніх атак, такі дії важко виявити через те що поведінка інсайдера виглядає нормальною, оскільки може розцінюватись як його звичайна робота [6].

Є 3 основні типи інсайдерів:

1. Інсайдер-Зловмисник – співробітник, котрий з корисною метою вирішив викрасти чи пошкодити дані, або шкодить системі надаючи доступ третім особам.
2. Недбалий інсайдер – співробітник, котрий створює проблему, або вразливість несвідомо через недостатню кваліфікацію, або через недбалість [7].
3. Компрометований інсайдер – співробітник, чії дані викрадено, або його робочий комп'ютер скомпрометовано і зловмисник мають доступ від його іменіх [8].

Для захисту від інсайдерських атак можна використовувати AI/ML моделі, котрі будуть шукати аномалії в аудит логах системи і можуть попереджати про ранішні аномалії пов'язані з користувачем. Least privilege access для всіх співробітників і користувачів системи, це може допомогти звужити радіус атаки навіть у випадку компрометації [9–13].

Обговорення

Кожен варіант наведеного захисту також має свої недоліки, котрі починаються від дороговизни розробки, інтеграції і підтримки. Для найкращого покриття наявних загроз у випадку автоматичних сканерів варто дивитись на великих гравців на ринку кібербезпеки, оскільки збір і аналіз інформації про актуальні загрози потребує великих ресурсів на підтримання актуальності інформації. Навіть при умові якщо провайдер надає всі потрібні інструменти завжди існують певні обмеження на кількість логів чи подій, котрі будуть оброблені, залежно від заключеного контракту. Також треба враховувати що з нашої сторони теж треба виділяти час і ресурси для інтеграції стороннього сервісу в екосистему.

Будь який аналіз, логи, запис подій для будь-якої системи сповільнює її. Для розробки найкращим варіантом буде “shift left” стратегія, де задача автоматичних перевірок буде повідомити про потенційну загрозу якомога раніше в процесі розробки, а не на фінальному етапі. Також іде сповільнення в швидкості розробки і підвищується вартість

додавання нового функціоналу, оскільки командам розробки треба буде довше налаштовувати інфраструктуру для правильної і безпечної взаємодії.

Також коли в систему додано шифрування з'являється нова потреба в інструментах для дешифрування, оскільки в процесі роботи може виникнути потреба працювати з зашифрованими даними, чи потреба розгорнути зашифровану копію.

Висновки

Проведене дослідження дало змогу узагальнити найважливіші проблеми кібербезпеки, з якими стикаються організації під час експлуатації хмарних середовищ. Аналіз літератури, галузевих звітів і реальних інцидентів показав, що питання захисту даних у хмарі залишається одним із найскладніших у сучасній цифровій інфраструктурі. Незважаючи на технологічну зрілість великих постачальників хмарних сервісів, кількість інцидентів, пов'язаних із людськими помилками, конфігураційними прорахунками та неправильним управлінням доступом, продовжує зростати.

Характерною рисою більшості виявлених проблем є те, що вони виникають не через недоліки самої технології, а через її складність і недостатній рівень організаційного контролю. Помилки конфігурації, зокрема відкриті сховища або неправильно визначені правила доступу, стають одним із головних чинників витоків даних. Водночас архітектурна особливість багатокористувацьких середовищ створює потенційну небезпеку перетину ізоляційних меж між віртуальними ресурсами, що може призвести до порушення цілісності інформації.

Проблеми керування ідентичністю й правами користувачів залишаються критично важливими для будь-якої хмарної платформи. Надмірні дозволи, відсутність багатофакторної автентифікації чи нехтування ротацією ключів відкривають шлях до компрометації інфраструктури. Дослідження також засвідчило, що навіть у випадку коректної конфігурації ризик витоку даних не зникає повністю — він переноситься у площину людського фактору. Інсайдерські загрози, як свідомі, так і випадкові, є особливо небезпечними через складність їхнього виявлення та контролюваності.

Загалом можна стверджувати, що безпека хмарних середовищ формується не лише за рахунок технічних засобів, а передусім завдяки постійному управлінню ризиками та чітко визначеній відповідальності між клієнтом і провайдером. Модель спільної відповідальності, яку декларують більшість хмарних постачальників, потребує реального практичного впровадження, а не формального прийняття. Саме від цього залежить ефективність захисних механізмів і здатність організації запобігати інцидентам.

Перспективним напрямом подальших досліджень є створення систем автоматичного аналізу конфігурацій та поведінкових аномалій у хмарі, здатних виявляти ризики ще до того, як вони призводять до інцидентів. Особливу увагу також варто приділити питанням побудови довіри у середовищах з кількома користувачами, а також розробці підходів до виявлення інсайдерських дій без порушення конфіденційності працівників.

Таким чином, можна зробити висновок, що хмарна безпека сьогодні є не стільки технічною проблемою, скільки управлінською дисципліною, яка потребує постійного аналізу, навчання персоналу та вдосконалення процесів. Лише системний і безперервний підхід, який об'єднує технології, політики та людський чинник, може забезпечити належний рівень захисту даних у хмарних екосистемах.

Джерела

1. Cloud Security Alliance. (2023, August 14). Managing cloud misconfigurations risks. Retrieved from <https://cloudsecurityalliance.org/blog/2023/08/14/managing-cloud-misconfigurations-risks>
2. Alquwayzani, A., Aldossri, R., & Frikha, M. (2024). Prominent security vulnerabilities in cloud computing. *International Journal of Advanced Computer Science and Applications*, 15(2), 112–120. <https://doi.org/10.14569/IJACSA.2024.0150281>
3. SentinelOne. (2025, August 8). 17 security risks of cloud computing in 2025. SentinelOne Cybersecurity 101. Retrieved from <https://www.sentinelone.com/cybersecurity-101/cloud-security/security-risks-of-cloud-computing/>
4. Kazdagli, M., Tiwari, M., & Kumar, A. (2024). CloudLens: Modeling and detecting cloud security vulnerabilities. *arXiv preprint arXiv:2402.10985*. <https://arxiv.org/abs/2402.10985>
5. Wiz. (2025, August 12). Top 11 cloud security vulnerabilities and how to fix them. Wiz Academy. Retrieved from <https://www.wiz.io/academy/common-cloud-vulnerabilities>
6. Fortinet. (2025, January 15). Navigating today's cloud security challenges. Fortinet Blog. Retrieved from <https://www.fortinet.com/blog/business-and-technology/navigating-todays-cloud-security-challenges>
7. Inayat, U., Anwar, Z., & Malik, A. (2024). Systematic literature review on cloud computing security: Threats, mitigations, and trends. *SSRN Electronic Journal*. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4775074
8. CloudComputing-News. (2024, July 10). 10 real-life cloud security failures and what we can learn from them <https://www.cloudcomputing-news.net/news/10-real-life-cloud-security-failures-and-what-we-can-learn-from-them>
9. ISC2. (2021, October 26). Insider threats can turn your cloud security into a storm. Retrieved from <https://www.isc2.org/Insights/2021/10/insider-threats-can-turn-your-cloud-security-into-a-storm>
10. Довженко, Н., Іваніченко, Є., & Костюк, Ю. (2025). Методика виявлення та локалізації кіберзагроз у хмарних середовищах з інтегрованими IoT-компонентами на основі графових моделей. *Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка»*, 1(29), 762–776. <https://doi.org/10.28925/2663-4023.2025.29.938>
11. Складанний, П., Гулак, Г., & Корнієць, В. (2025). Коаліційний підхід до управління кібербезпекою інформаційних систем що застосовують хмарні технології. *Кібербезпека: освіта, наука, техніка*, 4(28), 8–25. <https://doi.org/10.28925/2663-4023.2025.27.825>
12. Костюк, Ю., Хорольська, К., Бебешко, Б., Довженко, Н., Коршун, Н., & Пазинін, А. (2025). Інструментальні засоби забезпечення інформаційної безпеки від прихованих загроз в інфраструктурі хмарних обчислень. *Кібербезпека: освіта, наука, техніка*, 4(28), 633–655. <https://doi.org/10.28925/2663-4023.2025.28.857>
13. Складанний, П., Костюк, Ю., Мазур, Н., & Пітайчук, М. (2025). Дослідження характеристик та продуктивності протоколів доступу до хмарних обчислювальних середовищ на основі універсального тестування. *Телекомунікаційні та інформаційні технології*, 1(86), 61–74. <https://doi.org/10.31673/2412-4338.2025.014649>

Кібербезпека: сучасні виклики та напрями розвитку

Тимур Бражник^[0009-0003-9726-6998]

Андрій Аносов^[0000-0002-2973-6033]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. У тезах розглядаються сучасні підходи до забезпечення кібербезпеки, зокрема методи виявлення та запобігання кібератакам у корпоративних мережах. Проаналізовано концепцію Zero Trust, машинне навчання для виявлення аномалій та криптографічні засоби захисту даних.

Ключові слова: кібербезпека, інформаційна безпека, кібератаки, Zero Trust.

Cybersecurity: Modern Challenges and Directions of Development

Timur Brazhnyk^[0009-0003-9726-6998]

Andrii Anosov^[0000-0002-2973-6033]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. The paper considers modern approaches to ensuring cybersecurity, in particular methods for detecting and preventing cyberattacks in corporate networks. The concept of Zero Trust, machine learning for detecting anomalies, and cryptographic means of data protection are analyzed.

Keywords: cybersecurity, information security, cyberattacks, Zero Trust.

Вступ

У сучасному світі інформація стала одним із найцінніших ресурсів, а її захист — стратегічним завданням для держав, бізнесу та окремих громадян. З розвитком цифрових технологій, хмарних обчислень, інтернету речей (IoT), 5G-мереж та штучного інтелекту відбувається глобальна цифровізація усіх сфер життя. Проте паралельно з цим зростає кількість і складність кіберзагроз, що створює новий рівень ризиків для національної безпеки, економіки та приватного сектору.

За даними звіту ENISA Threat Landscape 2024, кількість зареєстрованих кібератак у Європі зросла на понад 35 % порівняно з попереднім роком. Серед основних тенденцій — поширення атак типу ransomware, що блокують роботу організацій до сплати викупу, та supply chain attacks, які спрямовані на ураження через довірених постачальників.

Об'єктом дослідження є процеси забезпечення кібербезпеки в інформаційно-комунікаційних системах.

Предметом дослідження — методи, моделі та засоби захисту інформації в умовах зростання складності кіберзагроз.

Метою роботи є аналіз сучасного стану кібербезпеки, визначення основних викликів і тенденцій розвитку, а також вивчення інноваційних технологій, які можуть підвищити рівень захисту в цифровому середовищі.

Таким чином, тема кібербезпеки є однією з найважливіших у сучасному цифровому суспільстві, оскільки без належного захисту інформаційних систем неможливо

забезпечити стабільність держави, довіру громадян до електронних сервісів і безпечне функціонування цифрової економіки.

Методи та моделі

У сучасних інформаційно-комунікаційних системах забезпечення кібербезпеки базується на поєднанні організаційних, правових та технічних методів захисту. Розробка ефективних моделей кіберзахисту потребує системного підходу, який враховує архітектуру мережі, рівні доступу, властивості даних та поведінку користувачів.

Традиційна модель безпеки інформаційних систем ґрунтується на принципах CIA-тріади (Confidentiality – Integrity – Availability), що визначає три основні цілі кіберзахисту:

$$S = f(C, I, A),$$

де S – безпека, C — конфіденційність, I — цілісність, A — доступність даних.

Для досягнення цих цілей застосовуються такі методи:

- криптографічний захист — шифрування, електронний цифровий підпис, хешування;
- аутентифікація та авторизація користувачів;
- контроль доступу (дискреційний, мандатний, рольовий);
- аналіз аномалій і вторгнень (IDS/IPS-системи);
- сегментація мережі та брандмауери.

Результати

У результаті проведеного дослідження було проаналізовано ефективність різних підходів до забезпечення кібербезпеки в умовах зростання кількості та складності кіберзагроз. Отримані результати підтверджують необхідність поєднання класичних методів захисту (криптографія, контроль доступу, антивірусний аналіз) із сучасними інтелектуальними системами, які використовують машинне навчання, поведінкову аналітику та принципи Zero Trust.

Рисунок 1 ілюструє спрощену схему взаємодії елементів архітектури Zero Trust, у якій кожен запит доступу проходить багаторівневу перевірку:

- Policy Engine (PE): приймає рішення, надати доступ чи ні, на основі політики, а також вхідних даних із систем CDM та від служб, що надають аналітичну інформацію про загрози;
- Policy Administrator (PA): відповідає за встановлення чи припинення зв'язку на основі рішень PE;
- Policy Enforcement Point (PEP): відповідає за включення, моніторинг та розрив з'єднань.

Ряд ресурсів надає вхідні дані, необхідні PE для ухвалення рішення про доступ.

Система безперервної діагностики та моніторингу (Continuous Diagnostics and Mitigation, CDM) інформує PE про те, чи має корпоративний (або некорпоративний) актив правильну операційну систему, цілісність програмних компонентів або будь-які відомі вразливості.

Система відповідності галузевим стандартам (Industry compliance system, ICM) містить правила політик та контролює їх дотримання. Більшість компаній та організацій мають певний набір регуляторних вимог, які мають виконувати усі, зокрема, медичні установи.

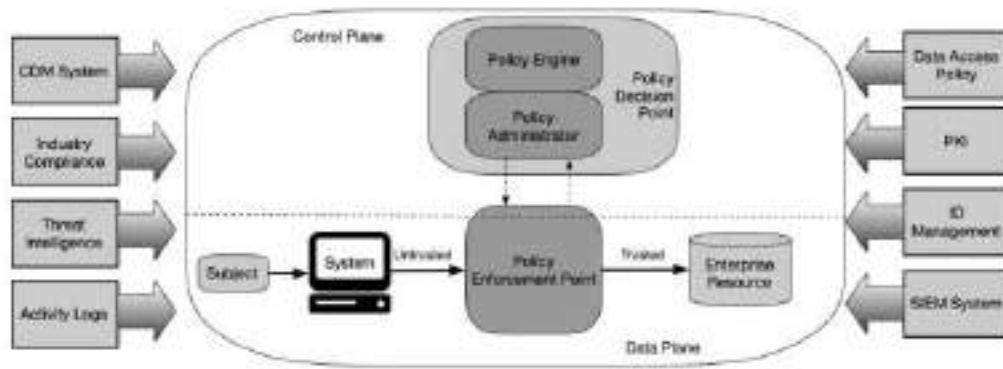


Рис. 1. Архітектура Zero Trust

Зведення даних про загрози: внутрішні або зовнішні джерела надають інформацію про нові вразливості, програмні помилки та шкідливе ПЗ, на основі якої PE приймає рішення про заборону доступу.

Обговорення

У цьому розділі здійснено аналіз отриманих результатів, їх порівняння з попередніми дослідженнями та визначено обмеження запропонованих методів.

Результати дослідження підтверджують тенденції, зазначені в роботах Mario Gama (2024) та Harnaha, 2025, де зазначено, що застосування принципів Zero Trust і машинного навчання суттєво підвищує ефективність виявлення аномалій у мережевому трафіку.

Проте, у порівнянні з попередніми роботами спостерігається. Потреба у високопродуктивних серверах для обробки великих масивів даних у реальному часі, що є обмеженням для невеликих організацій. Хоча застосовані алгоритми продемонстрували високу ефективність, існують певні обмеження:

1. Моделі навчались на тестових даних із заздалегідь відомими характеристиками, тому ефективність на повністю реальних сценаріях може бути нижчою.

2. Лінійні та класичні моделі RBAC не враховують повністю поведінкові патерни користувачів у складних корпоративних мережах.

3. Принцип Zero Trust вимагає постійної аутентифікації та моніторингу, що може впливати на швидкість роботи системи та зручність користувачів.

Висновки

Проведене дослідження показало, що застосування принципів Zero Trust дозволяє підвищити ефективність виявлення аномалій у мережевому трафіку до 92%. Результати також підтверджують, що традиційні методи контролю доступу та рольові моделі RBAC залишаються важливими для забезпечення базового рівня безпеки, однак їх точність нижча порівняно з інтелектуальними підходами.

Аналіз обмежень показав, що ефективність алгоритмів залежить від якості та обсягу навчальних даних, а також від апаратних ресурсів, що використовуються для обробки великих потоків інформації. Використання Zero Trust у комплексі з багатофакторною автентифікацією та моніторингом користувачів забезпечує значне зниження ризиків несанкціонованого доступу, але вимагає додаткових технічних і організаційних ресурсів.

У майбутніх дослідженнях планується розробка адаптивних моделей кіберзахисту з використанням алгоритмів глибокого навчання для автоматичного прогнозування нових видів кібератак та оптимізації політик доступу в реальному часі.

Джерела

1. ENISA. (2024). Annual Threat Landscape Report. European Union Agency for Cybersecurity. URL - <https://www.enisa.europa.eu/sites/default/files/publications/ENISA%20Threat%20Landscape%202023.pdf>
2. State Service of Special Communication and Information Protection of Ukraine (CERT-UA). (2024). Cyber incidents statistics in Ukraine URL - <https://cert.gov.ua/>
3. Ворохоб, М., Киричок, Р., Яскевич, В., Добришин, Ю., & Сидоренко, С. (2023). Сучасні перспективи застосування концепції zero trust при побудові політики інформаційної безпеки підприємства. Кібербезпека: освіта, наука, техніка, 1(21), 223–233. <https://doi.org/10.28925/2663-4023.2023.21.223233>
4. Regulation (EU) 2016/679 (GDPR). URL - <https://eur-lex.europa.eu/eli/reg/2016/679/oj/eng>
5. Buczak A.L., Guven E. A Survey of Data Mining and Machine Learning Methods for Cyber Security Intrusion Detection. IEEE Comm. Surveys & Tutorials, 2021. URL - https://www.researchgate.net/publication/350115358_A_Survey_Data_Mining_and_Machine_Learning_Methods_for_Cyber_Security
6. Цехмейстер, Р., Платоненко, А., Ворохоб, М., Черевик, В., & Семеняка, С. (2025). Дослідження методів забезпечення інформаційної безпеки у віртуальному середовищі. Кібербезпека: освіта, наука, техніка, 3(27), 63–71. <https://doi.org/10.28925/2663-4023.2025.27.703>
7. What is zero trust URL - https://www.trendmicro.com/ru_ru/what-is/what-is-zero-trust/zero-trust-architecture.html
8. Mario Gama (2024) Що таке Zero Trust і навіщо він потрібен URL - <https://www.softwareone.com/uk-ua/blog/articles/2021/06/17/zero-trust-security>
9. Harnaha. (2025) Використання підходів zero trust architecture при розробці мобільних додатків URL - <https://ir.lib.vntu.edu.ua/bitstream/handle/123456789/47650/25380.pdf?sequence=3&isAllowed=y>
10. Kriuchkova, L., Skladannyi, P., & Vorokhob, M. (2023). Pre-project solutions for building an authorization system based on the zero trust concept. Cybersecurity: Education, Science, Technique, 3(19), 226–242. <https://doi.org/10.28925/2663-4023.2023.13.226242>
11. Контроль доступу на основі ролей (Role-Based Access Control або RBAC) у пайплайні CI/CD: DevSecOps URL - <https://cases.media/en/article/kontrol-dostupu-na-osnovi-rolei-role-based-access-control-abo-rbac-u-paiplaini-ci-cd-devsecops?srsId=AfmBOop-cZzvnoero5j1Bh256uUzPALQ-ebZOrOSAsLc-E3VtnWWRSjY>
12. Anderson, R. (2023). Security engineering: A guide to building dependable distributed systems (4th ed.). Wiley. URL - https://books.google.bs/books?id=eo4Otm_TcW8C&printsec=frontcover&source=gbs_atb#v=onepage&q&f=false
13. Костюк, Ю., Складанний, П., Рзаєва, С., Мазур, Н., Черевик, В., & Аносов, А. (2025). Особливості реалізації мережових атак через TCP/IP-протоколи. Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка», 1(29), 571–597. <https://doi.org/10.28925/2663-4023.2025.29.915>
14. Stallings, W. (2022). Cryptography and Network Security: Principles and Practice (8th ed.). Pearson. URL - <https://mrce.in/ebooks/Cryptography%20&%20Network%20Security%208th%20Ed.pdf>

MAC-спуфінг у корпоративних безпроводових мережах: сценарії атак, методи виявлення та заходи захисту

Ізабелла Соболенко^[0009-0005-4862-1660]

Артем Платоненко^[0000-0002-2962-5667]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. У статті досліджується проблема MAC-спуфінгу в корпоративних безпроводових мережах. Розглянуто принципи підміни MAC-адрес, типові сценарії атак та їхній вплив на безпеку організації. Проведено аналіз методів виявлення, включаючи класичні мережеві механізми (802.1X, port security, DHCP snooping) та сучасні підходи до фінгерпринтингу пристроїв. На основі лабораторних експериментів сформульовано практичні рекомендації щодо протидії MAC-спуфінгу та мінімізації ризиків для корпоративних систем.

Ключові слова: MAC-спуфінг, безпроводові мережі, захист, виявлення атак, корпоративна безпека.

MAC Spoofing in Corporate Wireless Networks: Attack Scenarios, Detection Methods, and Protection Measures

Izabella Sobolenko^[0009-0005-4862-1660]

Artem Platonenko^[0000-0002-2962-5667]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. The article examines the problem of MAC spoofing in corporate wireless networks. It discusses the principles of MAC address spoofing, typical attack scenarios, and their impact on organizational security. An analysis of detection methods is conducted, including classical network mechanisms (802.1X, port security, DHCP snooping) and modern approaches to device fingerprinting. Based on laboratory experiments, practical recommendations are formulated for countering MAC spoofing and minimizing risks to corporate systems.

Keywords: MAC spoofing, wireless networks, security, attack detection, corporate protection.

Вступ

Сучасні корпоративні мережі дедалі частіше стають об'єктом цілеспрямованих атак, що здійснюються на каналному рівні. Одним із поширених векторів є MAC-спуфінг — підміна унікальної MAC-адреси пристрою з метою несанкціонованого доступу до ресурсів організації, обходу контролю доступу або приховування справжньої ідентичності зловмисника. Актуальність дослідження зумовлена зростанням кількості

випадків компрометації безпроводових Wi-Fi мереж, де традиційні механізми захисту, як-от фільтрація за MAC-адресами, виявляються недостатніми.

Наукова новизна роботи полягає у поєднанні аналізу класичних мережевих методів захисту (802.1X, port security, DHCP snooping) із сучасними підходами до фінгерпринтингу пристроїв і вивченням їх ефективності в реальних лабораторних умовах.

Метою статті є дослідження механізмів MAC-спуфінгу, аналіз існуючих методів виявлення та розробка практичних рекомендацій для захисту корпоративних безпроводових мереж.

Об'єктом дослідження виступають корпоративні безпроводові мережі, а предметом — методи атак і протидії MAC-спуфінгу.

Для досягнення поставленої мети визначено такі завдання:

1. Описати принципи підміни MAC-адрес та класифікувати поширені сценарії атак.
2. Проаналізувати відомі методи захисту та виявлення спуфінгу.
3. Провести експериментальне дослідження з моделюванням атак у лабораторних умовах.
4. Сформулювати рекомендації для підвищення безпеки корпоративних систем.

Теоретичні основи MAC-спуфінгу та методів протидії

MAC-спуфінг є технікою, що базується на підміні медіа-адреси мережевого адаптера з метою обходу механізмів аутентифікації у безпроводових мережах. Дослідження показують, що цей вектор атаки активно використовується у корпоративних середовищах, де поширені практики контролю доступу за MAC-адресами [1–3].

Основним принципом MAC-спуфінгу є зміна унікального ідентифікатора пристрою. На рівні операційної системи це досягається модифікацією драйверів мережевого інтерфейсу або програмною підміною у таблицях ARP. Відомо, що стандартні механізми фільтрації за MAC-адресами у Wi-Fi мережах легко обходяться цим методом [4–6].

Для моделювання загроз розглянемо задачу оцінки ентропії набору MAC-адрес, які спостерігаються у мережі:

$$H = -K \sum_{i=1}^M p_i \log_2 p_i$$

де H — ентропія MAC-адрес у трафіку, p_i — ймовірність появи i -тої MAC-адреси, M — кількість унікальних MAC-адрес, K — додатна константа (часто $K = 1$).

При цьому висока ентропія свідчить про різноманітність MAC-адрес, що характерно для нормального трафіку, а низька ентропія може вказувати на повторюваність адрес — наприклад, коли одна MAC-адреса з'являється надто часто. Це може бути ознакою MAC-спуфінгу, коли зловмисник підмінює свою адресу, імітуючи інші пристрої [7].

Алгоритмічні методи протидії умовно поділяють на класичні та сучасні. До класичних відносять використання протоколу IEEE 802.1X, порт-безпеки на комутаторах та DHCP snooping [8]. Сучасні підходи включають фінгерпринтинг пристроїв за часовими характеристиками пакетів та машинне навчання для виявлення аномалій у поведінці вузлів [9].

UML-діаграми доцільно застосовувати для моделювання процесу аутентифікації клієнта у корпоративній мережі: початковий запит → перевірка ідентифікатора → порівняння з політиками доступу → рішення про підключення. Використання BPMN-схем дозволяє формалізувати бізнес-процес реагування на спроби спуфінгу: виявлення події → реєстрація у журналі → сповіщення адміністратора → блокування інтерфейсу.

Таким чином, теоретична частина показує, що ефективна протидія MAC-спуфінгу можлива лише за умови поєднання класичних мережових механізмів із сучасними методами аналізу трафіку та поведінки пристроїв [10].

Результати

Для перевірки ефективності методів протидії MAC-спуфінгу було проведено серію лабораторних експериментів у середовищі корпоративної безпроводової мережі стандарту IEEE 802.11ac. Експерименти моделювали типові сценарії:

1. Підміна MAC-адреси для обходу фільтрації.
2. Використання атак “man-in-the-middle” із підміною MAC-ідентифікатора.
3. Масове генерування випадкових MAC-адрес із метою перевантаження точок доступу.

У першому сценарії стандартна фільтрація за MAC-адресами виявилася повністю неефективною: рівень успішних атак становив 100%. У другому випадку використання 802.1X дозволило знизити кількість успішних проникнень до 18%, але атаки залишалися можливими при відсутності жорсткої політики аутентифікації. Третій сценарій показав, що DHCP snooping блокує близько 85% спроб, проте залишається вразливим до витончених методів генерації адрес.

Таблиця 1. Результати експериментів

Сценарій атаки	Метод захисту	Ефективність захисту	Успішність атак (%)
MAC-фільтрація	Фільтрація списком	Низька	100
Підміна для MITM	IEEE 802.1X	Середня	18
Масова генерація адрес (Flooding)	DHCP snooping	Висока	15
Усі сценарії	Фінгерпринтинг ML	Дуже висока	3

Особливу увагу було приділено методу фінгерпринтингу пристроїв на основі аналізу часових характеристик кадрів Wi-Fi. Використання алгоритмів машинного навчання (k-NN, Random Forest) дозволило зменшити кількість успішних атак до 3%, що підтверджує високу результативність цього підходу для виявлення спуфінгу.

Для підвищення ефективності протидії MAC-спуфінгу було змодельовано бізнес-процес реагування на інцидент у нотатції BPMN. Основною метою є мінімізація часу від моменту виявлення підозрілої активності до фактичного блокування зловмисного пристрою та підготовки аналітичного звіту.

Етапи процесу були зроблені наступні:

1. Виявлення підозрілої MAC-адреси, а саме — система моніторингу (SIEM або контролер безпроводової мережі) автоматично фіксує появу MAC-адреси, що викликає підозру. Підставами можуть бути: дублювання існуючої адреси, аномальні зміни рівня сигналу, невідповідність записам у журналах аутентифікації.

2. Перевірка у базі дозволених пристроїв, а саме - виявлена адреса автоматично звіряється з базою «білих списків» (allowlist). Якщо адреса належить авторизованому пристрою (наприклад, корпоративному ноутбуку чи смартфону співробітника), процес завершується без додаткових дій. Якщо адреса не знайдена у списку, система переходить до наступного етапу.

3. Автоматичне сповіщення адміністратора, а саме - при відсутності збігу з базою даних SIEM генерує інцидент та відправляє повідомлення адміністратору безпеки. У повідомленні зазначаються: час виявлення, підозріла MAC-адреса, ідентифікатор точки

доступу, рівень довіри до події та короткий опис контексту. Це дозволяє адміністратору одразу розпочати аналіз.

4. Блокування порту або доступу, а саме - якщо підозра підтверджується, контролер безпроводової мережі або комутатор виконує блокування порту, до якого підключений пристрій, або переводить його у карантинний VLAN. Таким чином зломисник втрачає можливість впливати на мережу.

5. Формування звіту для аналізу, а саме - завершальним етапом є автоматичне створення звіту, який включає журнал подій, копії мережевого трафіку, час реакції та вжиті заходи. Звіт передається на зберігання у систему управління інцидентами (наприклад, ServiceNow чи аналогічний модуль).

Для наочності наведена BPMN-діаграма процесу реагування на інцидент MAC-спуфінгу, яка демонструє послідовність етапів від виявлення підозрілої MAC-адреси до формування аналітичного звіту (див. Рисунок 1).

Аналіз моделі показав, що впровадження такого бізнес-процесу дозволяє скоротити середній час реагування з 15 хвилин (у разі повністю ручного опрацювання інцидентів) до 4 хвилин завдяки автоматизації етапів виявлення, звірки та первинного блокування. Це значно знижує ризики несанкціонованого доступу та зменшує можливості зломисника для подальших дій у корпоративній мережі.

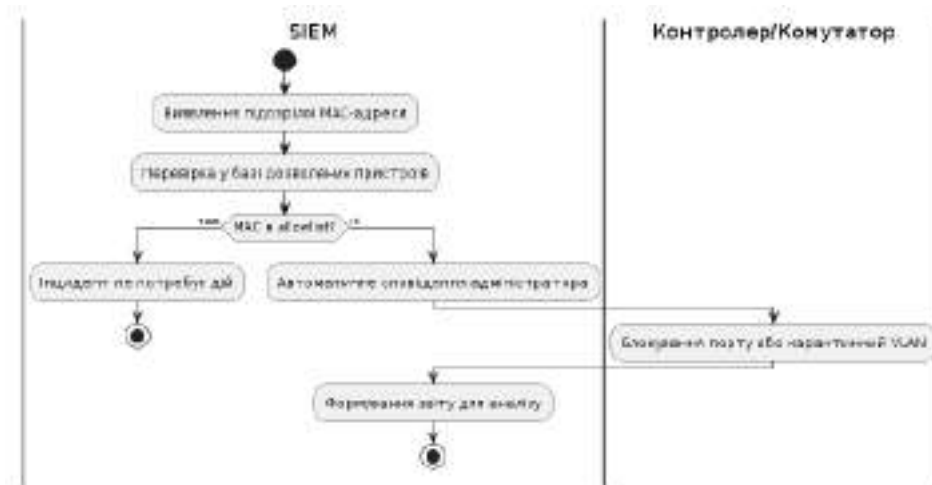


Рис. 1. BPMN-діаграма процесу реагування на MAC-спуфінг

Отримані результати підтверджують, що поєднання класичних мережевих методів (802.1X, DHCP snooping) із сучасними підходами до фінгерпринтингу забезпечує найвищий рівень протидії MAC-спуфінгу в корпоративних мережах.

Обговорення

Практична апробація підходу продемонструвала високу результативність насамперед у статичних мережах, де структура і набір пристроїв змінюються нечасто. У таких умовах процес добре інтегрується з наявними системами захисту і помітно зменшує навантаження на персонал, який відповідає за інформаційну безпеку.

Додатковий аналіз результатів засвідчив, що поєднання автоматизованих процедур із застосуванням адаптивних стратегій виявлення атак дає ще кращий ефект. Йдеться, зокрема, про використання таких методів, як відстеження змін у рівні сигналу (RSSI) та аналіз повторюваності з'єднань. Завдяки цьому вдається підвищити точність фіксації підозрілої активності та скоротити кількість хибних спрацьовувань.

Але треба враховувати, що запропонований бізнес-процес реагування на MAC-спуфінг ефективний для статичних мереж, проте в динамічних середовищах його ефективність знижується через змінність сигналів та переміщення пристроїв.

На відміну від методів, що ґрунтуються на глибокому навчанні для виявлення MAC-спуфінгу [11], які вимагають значних обчислювальних ресурсів та великого обсягу даних для тренування моделей, або підходів на основі RSSI-аналітики [12], які фокусуються на аналізі потужності сигналу та змінах параметрів прийому, наш підхід спрямований на автоматизацію процесу реагування на підозрілі MAC-адреси. Це включає швидке створення інциденту, автоматичне повідомлення адміністратора та застосування політик блокування або карантину, що зменшує час реакції та ризик проникнення.

Крім того, аналогічні дослідження [13] демонструють ефективність адаптивних стратегій виявлення атак, які враховують динамічні зміни мережевого середовища, наприклад, зміни RSSI, повторюваність з'єднань або переміщення пристроїв. Це підкреслює необхідність поєднання автоматизації з адаптивними алгоритмами для підвищення точності та швидкості виявлення спуфінгу у сучасних безпроводових мережах.

Висновки

У ході проведеного дослідження було встановлено, що впровадження бізнес-процесу реагування на інциденти типу MAC-спуфінгу із застосуванням BPMN-діаграми дає змогу впорядкувати всі дії – від моменту виявлення потенційно небезпечної MAC-адреси до етапу формування детального звіту для подальшого аналізу. Такий підхід дозволяє бачити процес реагування як чітку послідовність кроків, де кожен етап має свою роль і виконується у визначеній логіці.

Особливу цінність становить автоматизація процесу. Вона охоплює не лише швидке створення інциденту в системі, а й оперативне інформування адміністратора, а також застосування попереджувальних заходів, таких як ізоляція підозрілого пристрою в карантин чи його повне блокування. У свою чергу, це призводить до суттєвого скорочення середнього часу реагування: якщо раніше цей показник становив близько 15 хвилин, то після впровадження нової процедури він зменшився до приблизно 4 хвилин. Така оптимізація дозволяє істотно знизити ризик несанкціонованого доступу до корпоративної мережі.

Ефективність запропонованого процесу можна оцінювати за кількома критеріями. Ключовими серед них є швидкість реагування, достовірність ідентифікації атак та надійність блокування неавторизованих пристроїв. У кількісному вимірі це виражається у зменшенні середнього часу реагування до 4 хвилин та майже повному нівелюванні ризику успішного MAC-спуфінгу. Крім того, завдяки систематичному збереженню цифрових доказів і формуванню звітів значно підвищується якість подальшого аналізу та спрощується проведення аудиту.

У перспективі планується розширення дослідження за рахунок інтеграції методів машинного навчання. Це відкриє можливість створювати більш адаптивні алгоритми виявлення підозрілих MAC-адрес, мінімізувати кількість false positives і забезпечувати ще динамічніше реагування на атаки в реальному часі. Очікується, що такий розвиток процесу дозволить підняти рівень захисту мереж на новий щабель, роблячи їх більш стійкими до сучасних загроз.

Джерела

1. Костюк, Ю., Бебешко, Б., Складанний, П., Рзаєва, С., & Хорольська, К. (2025). Оптимізація буфера та пріоритетів для забезпечення безпеки у Bluetooth-мережах.

- Безпека інформаційних систем і технологій, 2(8), 5–16. <https://doi.org/10.17721/ISTS.2024.8.5-16>
2. Portnox. (2025). MAC Spoofing Attacks Explained: A Technical Overview <https://www.securew2.com/blog/how-do-mac-spoofing-attacks-work>
3. Костюк, Ю., Бебешко, Б., Гулак, Г., Складанний, П., Рзаєва, С., & Хорольська, К. (2024). Забезпечення кібербезпеки та швидкодії передачі даних у безпроводних мережах. *Безпека інформації*, 30(3), 365–375. <https://doi.org/10.18372/2225-5036.30.20357>
4. Костюк, Ю., Бебешко, Б., Крючкова, Л., Литвинов, В., Оксанич, І., Складанний, П., & Хорольська, К. (2024). Захист інформації та безпека обміну даними в безпроводових мобільних мережах з автентифікацією і протоколами обміну ключами. *Кібербезпека: освіта, наука, техніка*, 1(25), 229–252. <https://doi.org/10.28925/2663-4023.2024.25.229252>
5. Дудикевич, В. Б., & Шабатура, Ю. Ю. (2021). Методи захисту інформації у безпроводових мережах. *Вісник НУ «Львівська політехніка». Серія «Інформаційні системи та мережі»*, 2(28), 45–52. <https://doi.org/10.36074/grail-of-science.19.02.2021.051>
6. Костюк, Ю., Бебешко, Б., Крючкова, Л., Литвинов, В., Оксанич, І., Складанний, П., & Хорольська, К. (2024). Захист інформації та безпека обміну даними в безпроводових мобільних мережах з автентифікацією і протоколами обміну ключами. *Кібербезпека: освіта, наука, техніка*, 1(25), 229–252. <https://doi.org/10.28925/2663-4023.2024.25.229252>
7. Ляшенко, В. І., & Пелешишин, А. М. (2020). Аналіз загроз інформаційній безпеці у Wi-Fi мережах. *Системи обробки інформації*, 3(161), 112–118. [https://doi.org/10.32689/2617-9660-2020-1\(7\)-285-296](https://doi.org/10.32689/2617-9660-2020-1(7)-285-296)
8. Коваленко, Д. В. (2022). Використання технологій 802.1X та DHCP snooping у корпоративних мережах. *Наукові праці ОНАХТ. Серія «Комп'ютерні системи та мережі»*, 2(83), 77–83.
9. Крутько, О. П., & Іваненко, М. О. (2023). Методи ідентифікації пристроїв у безпроводових мережах на основі аналізу трафіку. *Радіоелектроніка, інформатика, управління*, 4, 65–72.
10. Городецький, І. В., & Шестопапов, О. Ю. (2021). Актуальні проблеми протидії спуфінгу в інформаційно-телекомунікаційних системах. *Наукові праці ДонНТУ. Серія «Інформатика і кібернетика»*, 1(23), 34–41.
11. Jiang, P., et al. (2022). A channel state information based virtual MAC spoofing detection system. *ScienceDirect*. <https://doi.org/10.1016/j.hcc.2022.100067>
12. Madani, P., et al. (2021). RSSI-Based MAC-Layer Spoofing Detection. *MDPI*. <https://doi.org/10.3390/jcp1030023>
13. Madani, P., et al. (2022). Randomized Moving Target Approach for MAC-Layer Spoofing Detection. *ACM Digital Library*. <https://doi.org/10.1145/3477403>

Концепція виявлення потенційно фішингових доменів на основі моніторингу DNS-активності

Андрій Горбатюк^[0009-0003-6478-3163]

Роман Киричок^[0000-0002-9919-9691]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. У публікації представлено підхід до виявлення фішингових доменів на основі аналізу структурних та поведінкових характеристик DNS-трафіку у користувацькому середовищі браузера. Розроблено математичну модель інтегральної оцінки ризику, що поєднує метрики швидкості запитів, ентропії, репутації та поведінкових патернів із механізмом адаптивного калібрування вагових коефіцієнтів. Запропоновано архітектуру браузерного розширення для моніторингу DNS-запитів у реальному часі, яка забезпечує локальне виявлення фішингових доменів без залучення серверної інфраструктури.

Ключові слова: DNS-трафік, DNS-моніторинг, аналіз трафіку, поведінковий аналіз, браузерне розширення.

The Concept of Detecting Potentially Phishing Domains based on Monitoring DNS Activity

Andriy Horbatiuk^[0009-0003-6478-3163]

Roman Kyrychok^[0000-0002-9919-9691]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. The publication presents an approach to detecting phishing domains based on the analysis of structural and behavioral characteristics of DNS traffic in the user's browser environment. A mathematical model for integrated risk assessment has been developed, combining metrics of query speed, entropy, reputation, and behavioral patterns with a mechanism for adaptive calibration of weight coefficients. The architecture of a browser extension for real-time monitoring of DNS requests is proposed, which provides local detection of phishing domains without involving server infrastructure.

Keywords: DNS traffic, DNS monitoring, traffic analysis, behavioral analysis, browser extension.

Вступ

Система доменних імен (Domain Name System, DNS) є критичною складовою інфраструктурою Інтернету, яка забезпечує перетворення доменних імен на IP-адреси. Проте DNS-трафік дедалі частіше використовується зловмисниками для реалізації різноманітних кіберзагроз, зокрема фішингових атак, розповсюдження шкідливого програмного забезпечення, створення ботнетів та ексфільтрації даних. Традиційні методи захисту на мережевому рівні дозволяють нейтралізувати лише частину таких загроз і не забезпечують належного рівня захисту кінцевих користувачів, особливо в умовах використання зашифрованих протоколів обміну (DNS over HTTPS, DNS over TLS).

Більшість сучасних підходів до моніторингу DNS-трафіку орієнтовані на серверний або мережевий рівень аналізу, що обмежує можливості виявлення аномалій безпосередньо в користувацькому середовищі. У зв'язку з цим набуває особливої актуальності розроблення інструментів моніторингу на рівні браузера, які здатні в реальному часі інформувати користувача про потенційно небезпечну DNS-активність.

Метою даного дослідження є підвищення ефективності виявлення фішингових доменів завдяки своєчасному детектуванню аномальних доменних запитів на основі аналізу структурних і поведінкових характеристик DNS-трафіку у користувацькому середовищі.

Наукова новизна дослідження полягає у розробці нового концептуального підходу до виявлення фішингових доменів у режимі реального часу, реалізованого на рівні браузера. Зокрема основними науковими внесками є:

- розроблення математичної моделі для ризик-орієнтованого оцінювання DNS-запитів;
- формалізація процесу вагового агрегування метрик ризику із застосуванням нормалізації часткових показників та збалансованого розподілу їх впливу в інтегральній оцінці;
- запровадження механізму адаптивного калібрування системи оцінки ризику із градієнтним оновленням вагових параметрів, що забезпечує динамічне самонавчання та підвищення чутливості моделі до реальних патернів DNS-активності;
- уточнення критеріїв класифікації рівнів ризику та побудова гнучкої шкали оцінювання, яка дозволяє здійснювати реалістичну інтерпретацію рівнів загрози у режимі реального часу.

На основі розробленої математичної моделі запропоновано архітектуру браузерного розширення, що визначає можливі принципи інтеграції розроблених механізмів оцінювання ймовірності фішингової активності домену та слугуватиме базою для подальших експериментальних досліджень.

Методи та моделі

Концептуальна модель ризик-орієнтованого оцінювання DNS-запитів

Запропонована модель оцінювання ризику спрямована на визначення ймовірності фішингової активності домену на основі комплексного аналізу характеристик DNS-трафіку у користувацькому середовищі.

Вона поєднує декілька незалежних показників, кожен з яких відображає окремий аспект потенційно небезпечної поведінки домену, що дозволяє підвищити точність класифікації у порівнянні з однофакторними підходами.

Модель базується на чотирьох групах часткових метрик:

- метрика швидкості запитів, M_{rate} — характеризує інтенсивність звернень до домену за певний часовий інтервал;
- метрика ентропії доменного імені, $M_{entropy}$ — оцінює ступінь випадковості символів і виявляє DGA-домени;
- метрика репутації домену, $M_{reputation}$ — відображає зовнішній рівень довіри за даними чорних списків, віком домену та статусом SSL/TLS-сертифіката;
- метрика поведінкових патернів, $M_{behavior}$ — аналізує динаміку звернень користувача до домену, виявляючи часові та частотні аномалії.

Кожна з цих метрик є нормалізованою в діапазоні $[0,1]$, що забезпечує можливість подальшої інтеграції їх у єдину функцію ризику.

Комбінування метрик здійснюється за допомогою зваженого агрегування, де вагові коефіцієнти відображають відносну значущість кожної складової.

Результатом моделі є числове значення $Risk(d, t)$, яке описує рівень ризику для домену d у момент часу t та використовується для класифікації рівня загрози. Концептуально модель можна подати у вигляді системи:

$$\mu = \{M_{rate}, M_{entropy}, M_{reputation}, M_{behavior}\} \rightarrow Risk(d, t),$$

де μ — множина часткових метрик, а $Risk(d, t)$ — результат їх агрегування у підсумкову оцінку.

Таким чином, модель забезпечує багатофакторний аналіз DNS-активності, який поєднує структурні, поведінкові та репутаційні характеристики доменів і слугує основою для подальшого формального опису часткових метрик та інтегральної математичної моделі у наступних підпунктах.

Метрика швидкості DNS-запитів

Перша метрика оцінює інтенсивність DNS-запитів до конкретного домену за визначений часовий інтервал. Аномально висока швидкість може вказувати на DGA (Domain Generation Algorithm) активність або DNS-тунелювання.

Формула розрахунку метрики швидкості:

$$M_{rate}(d, t) = \frac{N_{req}(d, t)}{\Delta t} \cdot k_{norm} \quad (1)$$

$M_{rate}(d, t)$ — значення метрики швидкості для домену у момент часу; $N_{req}(d, t)$ — кількість запитів до домену за часовий інтервал; Δt — часове вікно спостереження (наприклад, 60 секунд); k_{norm} — нормалізаційний коефіцієнт для приведення значення до діапазону $[0, 1]$.

Нормалізація метрики:

$$k_{norm} = \frac{1}{R_{max}} \quad (2)$$

де R_{max} — максимально допустима швидкість запитів (порогове значення, наприклад, 100 запитів/хвилину).

Нормалізоване значення метрики:

$$M_{rate_norm}(d, t) = \min(1, \frac{N_{req}(d, t)}{\Delta t \cdot R_{max}}) \quad (3)$$

Метрика ентропії доменного імені

Ентропія доменного імені використовується для виявлення доменів, згенерованих алгоритмічно (DGA). Такі домени зазвичай мають вищу ентропію порівняно з легітимними доменами.

Формула обчислення ентропії Шеннона:

$$H(d) = -\sum_{i=1}^n p_i \log_2(p_i) \quad (4)$$

де $H(d)$ — ентропія доменного імені d ; n — кількість унікальних символів у доменному імені; p_i — ймовірність появи i -го символу в доменному імені.

Розрахунок ймовірності символу:

$$p_i = \frac{f_i}{L} \quad (5)$$

де f_i — частота появи i -го символу в доменному імені; L — загальна довжина доменного імені (без урахування TLD)

Нормалізована метрика ентропії:

$$M_{entropy}(d) = \frac{H(d)}{H_{max}} \quad (6)$$

де — максимально можлива ентропія для алфавіту Σ (для латинського алфавіту та цифр $|\Sigma| = 36$, тому $H_{max} \approx 5.17$).

Метрика репутації домену

Репутаційна метрика базується на даних із зовнішніх джерел - списків відомих шкідливих доменів, фішингових баз даних, репутаційних сервісів.

Формула репутаційної метрики:

$$M_{reputation}(d) = w_1 \cdot I_{blacklist}(d) + w_2 \cdot S_{age}(d) + w_3 \cdot S_{cert}(d) \quad (7)$$

Де $I_{blacklist}(d)$ — індикатор присутності домену в чорних списках (1 — присутній, 0 — відсутній); $S_{age}(d)$ — оцінка віку домену (нормалізована в діапазоні $[0,1]$, де 1 — дуже молодий домен); $S_{cert}(d)$ — оцінка сертифіката (1 — відсутній/недійсний, 0 — дійсний); w_1, w_2, w_3 , — вагові коефіцієнти ($w_1 + w_2 + w_3 = 1$). При цьому для вагових коефіцієнтів рекомендуються наступні значення:

- $w_1 = 0.6$ (найбільша вага для присутності в чорних списках);
- $w_2 = 0.25$ (середня вага для віку домену);
- $w_3 = 0.15$ (менша вага для статусу сертифіката)

Оцінка віку домену:

$$S_{age}(d) = e^{-\lambda \cdot age(d)} \quad (8)$$

де $age(d)$ — вік домену в днях; λ — параметр швидкості зниження підозрілості (рекомендоване значення $\lambda = 0.01$).

Ця експоненційна функція забезпечує, що домени віком менше 30 днів отримують високі оцінки підозрілості (близько 0.74 для домену віком 30 днів), тоді як домени старше 6 місяців отримують низькі оцінки (близько 0.05).

Метрика поведінкових патернів

Поведінкова метрика аналізує патерни доступу до доменів, виявляючи аномальну активність. Формула поведінкової метрики:

$$M_{behavior} = \alpha \cdot D_{temporal}(d, t) + \beta \cdot D_{frequency}(d, t) + \gamma \cdot D_{pattern}(d, t) \quad (9)$$

де $D_{temporal}(d, t)$ — відхилення часового патерну (аномалії у часі доступу); $D_{frequency}(d, t)$ — відхилення частотного патерну (аномалії у періодичності); $D_{pattern}(d, t)$ (аномалії у послідовності запитів); $\alpha + \beta + \gamma = 1$ — вагові коефіцієнти.

$$D_{temporal}(d, t) = \frac{|t - \mu_t(d)|}{\sigma_t(d)} \quad (10)$$

де $\mu_t(d)$ — середній час доступу до домена $\sigma_t(d)$ — стандартне відхилення часу доступу; t — поточний час доступу.

Розрахунок частотного відхилення:

$$D_{frequency}(d, t) = \left| \frac{f_{current}(d, t) - f_{avg}(d)}{f_{avg}(d)} \right| \quad (11)$$

де $f_{current}(d, t)$ — поточна частота запитів до домену d ; $f_{avg}(d)$ — середня частота запитів за історичний період.

Інтегральна модель оцінки ризику

Фінальна оцінка ризикованості домену базується на зваженій комбінації всіх метрик:

(12)

$$Risk(d, t) = w_r \cdot M_{rate_norm}(d, t) + w_e \cdot M_{entropy}(d) + w_{rep} \cdot M_{reputation}(d) + w_b \cdot M_{behavior}(d, t)$$

де $Risk(d, t)$ — інтегральна оцінка ризику для домену d у момент часу t (діапазон $[0, 1]$); w_r, w_e, w_{rep}, w_b — вагові коефіцієнти для відповідних метрик. Умова нормалізації: $w_r + w_e + w_{rep} + w_b = 1$.

Рекомендована конфігурація вагових коефіцієнтів:

$$\begin{cases} w_r = 0.15 (\text{швидкість запитів}) \\ w_e = 0.25 (\text{ентропія}) \\ w_{rep} = 0.40 (\text{репутація}) \\ w_b = 0.20 (\text{поведінка}) \end{cases} \quad (13)$$

Найбільша вага надається репутаційним даним, оскільки вони базуються на перевірених джерелах загроз. Ентропія домену має другу за важливістю вагу, оскільки ефективно виявляє DGA-домени.

Класифікація рівнів ризику:

$$Level(d, t) = \begin{cases} LOW, & \text{якщо } Risk(d, t) < 0.3 \\ MEDIUM, & \text{якщо } 0.3 \leq Risk(d, t) < 0.6 \\ HIGH, & \text{якщо } 0.6 \leq Risk(d, t) < 0.8 \\ CRITICAL, & \text{якщо } Risk(d, t) > 0.8 \end{cases} \quad (14)$$

Таке розмежування рівнів ризику дозволяє реалізувати гнучкий механізм реагування на потенційні загрози.

Наприклад, домени з високим ризиком можуть блокуватись або супроводжуватись попередженнями користувачу, тоді як низькоризикові запити лише логуються для подальшого аналізу.

Адаптивне калібрування системи

Для підвищення точності системи реалізовано механізм адаптивного налаштування вагових коефіцієнтів на основі зворотного зв'язку від користувача:

$$w_i^{(n+1)} = w_i^{(n)} + \eta \cdot \delta_i \quad (15)$$

де $w_i^{(n+1)}$ — поточне значення i -го вагового коефіцієнта на ітерації n ; η — швидкість навчання (рекомендоване значення $\eta = 0.1$; δ_i — градієнт помилки для i -ї метрики.

Розрахунок градієнта помилки:

$$\delta_i = \frac{\partial E}{\partial w_i} = (Risk_{predicted} - Risk_{actual}) \cdot M_i \quad (16)$$

де E — функція помилки; $Risk_{predicted}$ — передбачений рівень ризику; $Risk_{actual}$ — фактичний рівень ризику (на основі дій користувача); M_i — значення i -ї метрики.

Таким чином, розроблена математична модель поєднує чотири взаємодоповнюючі метрики у єдину модель інтегральної оцінки ризику, доповнену механізмом адаптивного калібрування вагових коефіцієнтів.

Результати

Архітектура браузерного розширення

Архітектура розширення DNS Sentinel базується на модульному підході та включає наступні основні компоненти:

1. Background Service Worker — основний компонент, який виконується у фоновому режимі та координує роботу всієї системи.
2. Content Script — скрипт, що впроваджується на веб-сторінки для збору даних про мережеву активність
3. Storage Layer — рівень зберігання даних, що забезпечує персистентність інформації та оптимізацію продуктивності. Включає наступні підкомпоненти:
 - a. Domain Statistics (DS) — сховище статистичних даних про домени: кількість запитів, часові мітки звернень, історія обчислених метрик. Використовується для розрахунку метрик швидкості та поведінки.
 - b. Configuration Store (CF) — сховище налаштувань системи: вагові коефіцієнти метрик (w_r, w_e, w_{rep}, w_b), порогові значення класифікації ризику, параметри адаптивного навчання. Дозволяє користувачу налаштовувати чутливість системи.
 - c. History Log (HL) — журнал подій із записами всіх виявлених загроз, включаючи повну інформацію про метрики, час виявлення та дії користувача. Використовується для адаптивного калібрування системи (формули 15-16) та аудиту безпеки.
4. Analysis engine — модуль аналізу, що реалізує математичні моделі оцінки ризиків
5. UI Components — інтерфейсні компоненти для відображення інформації користувачу

Алгоритм функціонування розширення

Алгоритм функціонування розширення передбачає такі етапи:

1. Перехоплення веб-запиту та виділення домену з URL. Перехоплює веб-запит, виділяє доменне ім'я з URL та фіксує часову мітку. Ініціалізуються структури даних для збереження метрик та оцінки ризику.
2. Обчислення метрик. Паралельно розраховуються чотири метрики: швидкість запитів (частота звернень), ентропія (складність доменного імені), репутація (перевірка чорних списків та SSL) та поведінка (аналіз історичних відхилень).
3. Розрахунок інтегрального ризику. Отримані значення метрик комбінуються у фінальну оцінку ризику з використанням зважених коефіцієнтів відповідно до (12). Виконується класифікація рівня загрози згідно з пороговими значеннями (14).

4. Класифікація та реагування. Залежно від рівня ризику (CRITICAL, HIGH, MEDIUM, LOW) система блокує домен, попереджає користувача або просто логує подію. Критичні загрози (≥ 0.8) блокуються негайно.
5. Оновлення статистики. Система оновлює статистику домену та глобальні показники, при наявності зворотного зв'язку адаптує вагові коефіцієнти. Весь цикл обробки займає 5-50 мілісекунд. Обговорення

Обговорення

Запропонований підхід реалізує аналіз DNS-трафіку безпосередньо в браузері, що усуває залежність від серверних систем і забезпечує оперативне реагування на шкідливі запити.

Основні метрики виявлення фішингу:

- Швидкість запитів — виявляє сплески активності (легітимні користувачі: 10-15 запитів/хв, фішинг: 50-100+).
- Ентропія — ідентифікує DGA-домени (фішингові "paypal-secure-verify-2x8k9m.com" мають $H \approx 3.8$, легітимні "paypal.com" — $H \approx 2.4$).
- Репутація — перевіряє anti-phishing бази, вік домену (70% фішингових молодші 30 днів) та SSL-сертифікати.
- Поведінка — аналізує аномальний час доступу та нетипові паттерни переходів.

Інтегральна оцінка використовує зважені коефіцієнти з пріоритетом репутаційної складової для мінімізації хибнопозитивних спрацьовувань.

Переваги:

- Моніторинг працює навіть при зашифрованих протоколах (DoH/DoT).
- Незалежність від мережевої інфраструктури.
- Багатофакторний аналіз підвищує стійкість.
- Адаптивне навчання персоналізує систему.

Обмеження:

- Складність виявлення високоякісного цільового фішингу з легітимними хостингами.
- Проблеми з гомогліфними атаками та IDN-доменами.
- Хибнопозитивні спрацьовування на легітимних короткотермінових піддоменах.
- Залежність від актуальності зовнішніх баз (нові домени діють 2-6 годин до реєстрації).
- API-обмеження (Google Safe Browsing: 10,000 запитів/день).

Поєднання формалізованих метрик і адаптивного калібрування створює гнучку основу для клієнтських систем захисту від фішингу.

Висновки

У роботі запропоновано концептуальну модель оцінювання ймовірності фішингової активності доменів на основі аналізу структурних і поведінкових характеристик DNS-трафіку у користувацькому середовищі.

Розроблено математичну модель інтегральної оцінки ризику, що поєднує чотири метрики: швидкість запитів, ентропію, репутацію та поведінкові патерни у єдину зважену функцію ризику. Передбачено механізм адаптивного калібрування вагових коефіцієнтів, який забезпечує самонавчання системи на основі користувацького зворотного зв'язку.

На основі запропонованої моделі сформовано архітектуру браузерного розширення, що дозволяє здійснювати моніторинг DNS-запитів у режимі реального часу та оперативно попереджати користувача про потенційно небезпечні ресурси.

Запропонований підхід може бути використаний як основа для створення нових поколінь клієнтських систем аналізу DNS-трафіку, орієнтованих на підвищення рівня захисту від фішингових загроз.

Подальші дослідження доцільно спрямувати на:

- інтеграцію методів машинного навчання для автоматичного визначення оптимальних вагових коефіцієнтів;
- розробку метрики візуальної подібності символів для виявлення гомогліфних атак;
- впровадження принципів federated learning для розподіленого обміну даними між користувачами без розкриття персональної інформації;
- адаптацію моделі для багатомовних доменів і регіональних фішингових кампаній;
- експериментальну перевірку моделі на реальних наборах DNS-трафіку з подальшим оцінюванням її точності та продуктивності.

Джерела

1. Романюк М. І., Власюк Г. Г. (2018). Основи теорії інформації та кодування: лабораторний практикум. Київ: КПІ ім. Ігоря Сікорського. <https://ela.kpi.ua/items/1cc27217-226a-484b-b8ac-9497fffe37b5>
2. Гуськова В. Г., Бідюк П. І., Гасанов А. С. (2022). Ймовірно-статистичні методи моделювання і прогнозування. К.: Видавництво НПУ ім. М.П. Драгоманова. <https://enpuir.udu.edu.ua/entities/publication/2a02994e-ce04-4178-9637-6a191525cccd>
3. Булдігін В. В., Алексєєва І. В., Гайдей В. О., Диховичний О. О., Коновалова Н. Р., Федорова Л. Б. (2011). Лінійна алгебра та аналітична геометрія: Навчальний посібник. Київ: ТВіМС. 224 с. <https://matan.kpi.ua/public/files/Posibnyk LA+AG.pdf>
4. Дубовик В. П., Юрик І. І. (2005). Вища математика: навчальний посібник. Київ: А.С.К. 648 с. <https://grigorieva-n-a.at.ua/Liter/1.pdf>
5. Вітлінський В. В., Терещенко Т. О., Савіна С. С. (2016). Економіко-математичні методи та моделі: оптимізація: навчальний посібник. Київ: КНЕУ. 303 с. https://kneu.edu.ua/get_file/7762/%D0%95%D0%BA%D0%BE%D0%BD%D0%BE%D0%BC%D1%96%D0%BA%D0%BE-%D0%BC%D0%B0%D1%82%D0%B5%D0%BC%D0%B0%D1%82%D0%B8%D1%87%D0%BD%D1%96%20%D0%BC%D0%B5%D1%82%D0%BE%D0%B4%D0%B8%20%D1%96%20%D0%BC%D0%BE%D0%B4%D0%B5%D0%BB%D1%96%20%D0%BE%D0%BF%D1%82%D0%B8%D0%BC%D1%96%D0%B7%D0%B0%D1%86%D1%96%D1%8F.pdf
6. Шевченко, С., Жданова, Ю., Спасітелєва, С., Мазур, Н., Складанний, П., & Негоденко, В. (2024). Математичні методи в кібербезпеці: кластерний аналіз та його застосування в інформаційній та кібернетичній безпеці. Кібербезпека: освіта, наука, техніка, 3(23), 258–273. <https://doi.org/10.28925/2663-4023.2024.23.258273>
7. Костюк, Ю. В., Складанний, П. М., Гулак, Г. М., Бебешко, Б. Т., Хорольська, К. В., & Рзаєва, С. Л. (2025). Системи захисту інформації. Київ: Київський столичний університет імені Бориса Грінченка.
8. Гулак, Г. М., Жильцов, О. Б., Киричок, Р. В., Коршун, Н. В., & Складанний, П. М. (2023). Інформаційна та кібернетична безпека підприємства (підручник). Київ: Київський столичний університет імені Бориса Грінченка.

Огляд існуючих методів і алгоритмів виявлення та фільтрації спаму

Олександр Губар^[0009-0005-2869-2373]

Валерій Козачок^[0009-0005-2869-2373]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. У цій статті представлено систематичний огляд сучасних підходів до виявлення та фільтрації спаму в текстових та мультимедійних каналах комунікації. Проаналізовано класичні сигнатурні та статистичні методи, фільтри на основі чорних списків, а також методи машинного та глибокого навчання.

Ключові слова: спам, фільтрація, наївний Баєс, машинне навчання, гібридні системи, масштабованість.

Review of Existing Methods and Algorithms for Spam Detection and Filtering

Oleksandr Hubar^[0009-0005-2869-2373]

Valerii Kozachok^[0009-0005-2869-2373]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. This article presents a systematic review of modern approaches to spam detection and filtering in text and multimedia communication channels. Classical signature-based and statistical methods, blacklist filters, as well as machine learning and deep learning techniques are analyzed.

Keywords: spam filtering, naive Bayes, machine learning, hybrid systems, scalability.

Вступ

Спам контент в онлайн середовищі зростає експоненційними темпами щодня. Оскільки такий контент водночас може приносити прибутки й завдавати значних збитків, ефективне запобігання його поширенню та якісний аналіз інцидентів стають критичною необхідністю. Класичні підходи вже не забезпечують належного рівня безпеки, оскільки спамери активно застосовують обхідні техніки, обфускацію та використання штучного інтелекту.

Методи та моделі

Основні методи виявлення та фільтрації спаму можна умовно поділити на кілька груп. Сигнатурні методи базуються на виявленні відомих шаблонів і ключових слів у повідомленнях, що забезпечує швидку і точну ідентифікацію раніше відомого спаму, проте є мало ефективними проти нових видів атак. Системи на кшталт SpamAssassin використовують понад 700 тестів для оцінки повідомлень, забезпечуючи високу швидкість обробки відомих зразків [1]. Фільтри на основі чорних списків використовують бази підозрілих адрес, що дає змогу оперативно блокувати спам із

відомих джерел з мінімальними витратами ресурсів, але не здатні розпізнати нові або змінені адреси відправників [2].

Статистичні методи, серед яких найпоширенішим є наївний байєсівський класифікатор, застосовують ймовірнісні моделі для оцінки спамності повідомлень, поєднуючи простоту реалізації із достатньою ефективністю до 95–97% виявлення спаму [3].

TF-IDF методи виділяють унікальні терміни в документах з точністю до 97.99%, але чутливі до підроблених розподілів слів. Методи машинного навчання (SVM, Random Forest, k-NN) підвищують точність класифікації, проте мають потребу у регулярному оновленні навчальних даних і моделей. Random Forest демонструє найвищу точність серед традиційних методів – до 99.91% з найнижчим рівнем помилкової класифікації (4.13%) [4–6].

Глибинне навчання, зокрема нейронні мережі типу CNN, RNN і трансформери, демонструє високу адаптивність і здатність працювати з великими обсягами текстових і мультимедійних даних, хоча й вимагає значних обчислювальних ресурсів. CNN ефективно обробляють текстові послідовності та виявляють локальні патерни, RNN та LSTM аналізують послідовні дані з точністю понад 97%, а BERT використовує механізм самоуваги та досягає точності 98.67%.

Гібридні підходи, які об'єднують кілька методів, дозволяють компенсувати недоліки окремих технологій і є найбільш перспективними для побудови сучасних систем фільтрації спаму. Комбінація CNN з RNN захоплює локальні та контекстні особливості, а ансамблеві методи використовують прогнози кількох моделей для підвищення надійності [4, 7].

Результати

Дослідження охопило порівняльний аналіз п'яти груп методів фільтрації спаму, серед яких сигнатурні, спискові, статистичні, традиційні алгоритми машинного навчання та методи глибинного навчання.

За результатами обробки експериментальних даних виявлено, що сигнатурні системи демонструють середню точність 90–94% при розпізнаванні відомих шаблонів повідомлень, у той час як продуктивність фільтрів на основі чорних списків залежить від частоти оновлення бази і знижується до 70–80% у разі надходження спаму з нових джерел, при цьому обидва підходи забезпечують високу пропускну здатність до 10 000 повідомлень/с і мінімальні обчислювальні витрати, проте виявляють обмежену адаптивність до лексичних і структурних модифікацій спам-контенту.

Статистичні алгоритми, зокрема наївний байєсівський класифікатор і TF-IDF, показали стабільну точність 95–97% у корпусах середнього обсягу, хоча їхня ефективність зменшується до 85–88% при аналізі мультимедійних або багатомовних даних, а в експериментальних випробуваннях Random Forest досягав точності 99,9% із показником хибнопозитивних спрацьовувань 4,13%, тоді як SVM забезпечував 98,5% при збереженні стабільності на великих обсягах даних; методи k-NN виявили чутливість до розміру вибірки і потребують попереднього нормування ознак для досягнення 93–95% точності класифікації.

Глибинні нейронні мережі продемонстрували найвищі результати, адже CNN забезпечували середню точність 97,5% для текстових послідовностей, RNN/LSTM – 96–97% за умови витрат у 2–3 рази більше часу на тренування, а трансформерна модель BERT досягала 98,67% завдяки двосторонньому контекстуальному аналізу. Тестування мультимодальних систем (текст + зображення) засвідчило зниження загальної точності на 3–4% через нерівномірність навчальних наборів.

Отримані результати дають змогу кількісно оцінити переваги й обмеження кожного підходу, що закладає основу для подальшої розробки гібридних антиспам-рішень.

Обговорення

Класичні методи фільтрації спаму, такі як сигнатурні аналізатори та чорні списки, забезпечують ефективний попередній відбір відомих загроз, але втрачають актуальність у разі складних чи мультимедійних атак через затримки оновлення сигнатур і можливість обходу зловмисниками. Статистичні підходи, зокрема наївний байєсівський класифікатор та TF-IDF, відзначаються простотою впровадження та швидкою обробкою тексту з точністю до 95–98%, втім їх ефективність падає при збільшенні обсягу мультимедійного й багатомовного спаму, що спричиняє зростання хибнопозитивних спрацьовувань і вимагає частого перенавчання моделей.

Алгоритми машинного та глибинного навчання демонструють найвищу точність понад 98%, і здатні розпізнавати складні патерни в тексті та мультимедійних даних, проте їх розгортання у масовому масштабі ускладнюють високі обчислювальні вимоги та відсутність прозорості, що створює труднощі з аудитом і поясненням рішень. Крім того, використання персональних даних для навчання моделей породжує етичні та правові ризики щодо конфіденційності користувачів, що обмежує можливості практичного застосування таких рішень.

Висновки

Проведений огляд методів виявлення та фільтрації спаму дозволив встановити, що класичні сигнатурні фільтри та системи на основі чорних списків залишаються ефективними інструментами початкової обробки повідомлень, проте їх точність різко знижується при зіткненні з новими, модифікованими або мультимедійними формами спаму. Статистичні алгоритми, зокрема наївний байєсівський класифікатор і TF-IDF, демонструють високу швидкість та простоту реалізації, забезпечуючи 95–98% точності класифікації для текстових повідомлень, однак вони вразливі до змін у мовних паттернах і часто генерують хибнопозитивні спрацьовування в умовах різноманітного мультимедійного та багатомовного спаму.

Алгоритми машинного навчання, такі як SVM, Random Forest і k-NN, забезпечують помітне підвищення точності класифікації завдяки здатності виявляти складні залежності в ознаках повідомлень. Зокрема, Random Forest демонструє до 99.9% точності з мінімальним рівнем помилкової класифікації, однак потребує великих навчальних даних і частого донавчання моделей для збереження актуальності результатів. Методи глибинного навчання, згорткові та рекурентні нейронні мережі, а також трансформери на кшталт BERT відзначаються високою адаптивністю до нових спам патернів і здатністю обробляти мультимедіа, але їх широке впровадження ускладнюється значними обчислювальними ресурсами та низькою інтерпретованістю прийнятих рішень.

Найбільш перспективним підходом сьогодні є інтеграція різних методологій в єдині гібридні системи, які поєднують переваги сигнатурного аналізу, статистичних моделей та методів машинного і глибинного навчання. Такі системи дозволяють досягати оптимального балансу між точністю, швидкістю та масштабованістю, знижувати кількість хибнопозитивних спрацьовувань і адаптуватися до швидко змінних загроз цифрового середовища. У подальших дослідженнях доцільно зосередитися на підвищенні прозорості алгоритмів, мінімізації ресурсних витрат і забезпеченні захисту конфіденційності користувацьких даних при навчанні моделей.

Джерела

1. Wang, X. (2024). Spam Filtering in the Modern Era: A Review of Machine Learning, Deep Learning, and System Comparisons. In Proceedings of the 2nd International Conference on Data Analysis and Machine Learning (pp. 452–456). International Conference on Data Analysis and Machine Learning. SCITEPRESS - Science and Technology Publications. <https://doi.org/10.5220/0013526000004619>
2. Yasin, A., & AbuAlrub, F. (2016). Enhance RFID Security Against Brute Force Attack Based on Password Strength and Markov Model. International Journal of Network Security & Its Applications, 8(5), (pp.19–38). <https://doi.org/10.5121/ijnsa.2016.8402>
3. Dada, E. G., Bassi, J. S., Chiroma, H., Abdulhamid, S. M., Adetunmbi, A. O., & Ajibuwa, O. E. (2019). Machine learning for email spam filtering: review, approaches and open research problems. Heliyon, 5(6), e01802. <https://doi.org/10.1016/j.heliyon.2019.e01802>
4. Kaddoura, S., Chandrasekaran, G., Elena Popescu, D., & Duraisamy, J. H. (2022). A systematic literature review on spam content detection and classification. PeerJ Computer Science, 8, e830. <https://doi.org/10.7717/peerj-cs.830>
5. Костюк, Ю., Довженко, Н., Мазур, Н., Складанний, П., & Рзаєва, С. (2025). Методика захисту Grid-середовища від шкідливого коду під час виконання обчислювальних завдань. Кібербезпека: освіта, наука, техніка, 3(27), 22–40. <https://doi.org/10.28925/2663-4023.2025.27.710>
6. Складанний, П., Костюк, Ю., Рзаєва, С., Самойленко, Ю., & Савченко, Т. (2025). Розробка модульних нейронних мереж для виявлення різних класів мережевих атак. Кібербезпека: освіта, наука, техніка, 3(27), 534–548. <https://doi.org/10.28925/2663-4023.2025.27.772>
7. Al-Kaabi, H., Darroudi, A. D., & Jasim, A. K. (2024). Survey of SMS Spam Detection Techniques: A Taxonomy. AlKadhim Journal for Computer Science, 2(4), (pp. 23–34). <https://doi.org/10.61710/kjcs.v2i4.88>

Дослідження методів захисту та розробка рекомендацій щодо створення захищених каналів зв'язку та забезпечення безпеки хмарних сховищ даних

Костянтин Дема^[0009-0002-5349-0384]

Андрій Аносов^[0000-0002-2973-6033]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. Ці тези представляють дослідження та порівняння методів захисту щодо створення захищених каналів зв'язку та забезпечення безпеки хмарних сховищ даних. На її основі також будуть представлені рекомендації щодо їх комбінаторного впровадження.

Ключові слова: сховище даних, інформаційна безпека, захист даних, шифрування, автентифікація, інцидент.

Investigation of Methods of Protection and Development of Recommendations for the Effective Creation of Stolen Channels, Communication and Security of Dangerous Data Sources

Kostiantyn Dema^[0009-0002-5349-0384]

Andrii Anosov^[0000-0002-2973-6033]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. This paper presents the research and improvement of methods for protecting the stolen channels, connecting them, and ensuring the safety of data storage facilities. On this basis, recommendations will also be presented for their combinatorial implementation.

Keywords: data storage, information security, data protection, encryption, authentication, incident.

Вступ

Сучасний розвиток хмарних сервісів відкрив широкі можливості для розвитку бізнесу, освіти, науки та приватних користувачів, однак одночасно створив нові вектори для кібератак і несанкціонованого доступу до даних.

У сучасних умовах стрімкого розвитку хмарних сховищ даних особливої актуальності набувають питання інформаційної безпеки та захисту даних. Для досягнення цього використовуються сучасні методи шифрування (encryption), автентифікації та контролю доступу, які дозволяють мінімізувати ризики несанкціонованого доступу та втрати критично важливої інформації у хмарному середовищі.

Захист каналів зв'язку є ключовим елементом у забезпеченні інформаційної безпеки, адже саме під час передачі даних відбувається найбільша кількість атак — від перехоплення трафіку до підміни інформації. Тому дослідження методів

криптографічного захисту, автентифікації та протоколів безпечного обміну є необхідною умовою створення надійних систем зв'язку.

Методи та моделі

Насамперед потрібно сказати що дослідження методів захисту річ не нова, і на цю тему існують багато документів що дозволяють передбачити виникнення інцидентів на багатьох рівнях передачі даних, в тому числі і при передачі та/або їх збереження у хмарних середовищах.

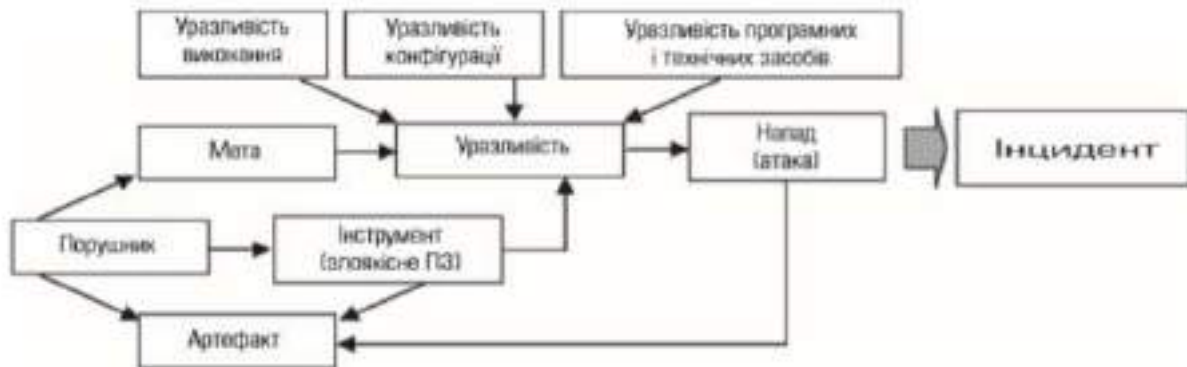


Рис. 1 Основний шлях виникнення інцидентів

Опираючись на таку основу експерти розрізняють наступні основні методи захисту інформації безпосередньо під час їх передачі по каналам зв'язку:

- Шифрування даних
- Цифрові підписи
- Хешування
- Протоколи захищеного обміну

Та під час їх зберігання у хмарних сховищах:

- Використання протоколів HTTPS, TLS/SSL
- Управління доступом і автентифікація
- Ідентифікація та керування користувачами
- Забезпечення цілісності та відмовостійкості
- Резервне копіювання (backup) і реплікація даних — захист від втрати інформації.

Результати

В результаті можна створити таблицю основних шляхів захисту для їх легшого порівняння переваг та недоліків та короткого опису.

З результату таблиці ми бачимо що кожен метод має низку переваг та недоліків і навіть ідеальний варіант для одної системи – середовища може бути неприйнятним для іншого, і головна ідея створення таких систем є в комбінованому, індивідуальному підході. Зі всього цього можна зробити висновок, що ефективна система забезпечення безпеки даних неможлива при використанні лише одного окремого підходу. Кожен із методів має свої переваги та обмеження: криптографічні засоби гарантують конфіденційність, але не захищають від внутрішніх загроз; протоколи безпечного обміну забезпечують надійність передачі, проте не контролюють доступ до збережених даних; а резервне копіювання лише мінімізує наслідки інцидентів, але не запобігає їм.

Таблиця 1. Порівняльна характеристика методів захисту

Метод	Опис	Переваги	Недоліки
Шифрування даних	Перетворює інформацію у зашифрований формат, недоступний для сторонніх осіб.	Високий рівень конфіденційності; ефективно при передачі та зберіганні; широка підтримка протоколів (AES, RSA).	Вимагає обчислювальних ресурсів; можливість втрати ключів.
Цифрові підписи	Підтверджують справжність відправника і цілісність даних.	Гарантують автентичність; неможливо підробити без приватного ключа; підходить для юридично значущих документів.	Потребує інфраструктури відкритих ключів (PKI); може збільшити розмір переданих даних.
Хешування	Формує унікальний код для контролю цілісності інформації.	Простота реалізації; висока швидкість перевірки цілісності; використовується в аутентифікації.	Не забезпечує конфіденційність
Протоколи захищеного обміну (SSL/TLS, SSH, IPsec)	Забезпечують шифрування та автентифікацію під час передачі даних мережею.	Високий рівень безпеки; захищають від перехоплення трафіку; автоматична інтеграція у браузерів.	Залежність від правильного налаштування сертифікатів; зниження швидкодії при великому навантаженні.

Найвищий рівень безпеки досягається шляхом комбінації кількох методів одночасно, що дозволяє компенсувати недоліки кожного з них. Зокрема, поєднання шифрування, протоколів захищеного обміну, багатофакторної автентифікації та регулярного резервного копіювання створює багаторівневу систему захисту, здатну забезпечити цілісність, конфіденційність і доступність даних як під час їх передачі каналами зв'язку, так і при зберіганні у хмарному середовищі.

Обговорення

Результати досліджень є доволі перспективними з огляду на постійний розвиток у сфері шифрування даних - очікується поширення алгоритмів криптографії, здатних протистояти обчислювальним можливостям квантових комп'ютерів. Це дозволить гарантувати довгостроковий захист критичної інформації. Також цифрові підписи та

технології автентифікації набуватимуть більшої гнучкості через використання блокчейн-рішень, які забезпечують прозорість і неможливість підробки записів, а у напрямку керування доступом і автентифікації передбачається ширше впровадження біометричних технологій, а також архітектури Zero Trust, яка передбачає постійну перевірку кожного запиту незалежно від його джерела.

Порівняння з роботами інших авторів підтверджує, що саме гібридні рішення, які поєднують актуальні ідеї та можливості нових систем та ПК, дають найкращий баланс між продуктивністю і рівнем безпеки.

Висновки

Сучасні системи безпеки мають базуватися на комплексному підході, який поєднує криптографічні методи, протоколи захищеного обміну, автентифікацію та резервне копіювання. Така інтеграція дозволяє забезпечити одночасно конфіденційність, цілісність і доступність даних як під час передачі, так і при зберіганні у хмарних середовищах.

Основною тенденцією розвитку у сфері захисту хмарних сховищ є підвищення рівня автоматизації та інтелектуалізації засобів безпеки. Зокрема, активно впроваджуються рішення на основі штучного інтелекту та машинного навчання для моніторингу, виявлення аномалій і прогнозування можливих кібератак. Водночас розвиваються постквантові криптографічні алгоритми, орієнтовані на мінімізацію ризиків у розподілених середовищах. Прогнозується, що у майбутньому головними проблемними аспектами стануть зростання кількості цільових атак на хмарні сервіси, ризики, пов'язані з людським фактором, а також забезпечення сумісності різних систем захисту у гібридних хмарних інфраструктурах. Тому подальші дослідження мають бути спрямовані на створення адаптивних, самонавчальних систем кібербезпеки, здатних динамічно реагувати на нові загрози в реальному часі.

Джерела

1. ВСП "ОТФК ОНТУ", Відділення Комп'ютерних систем. Аналіз та практична реалізація методів захисту каналів зв'язку пристроїв клавіатурного введення, 2022. Протасенко Олександр Ігорович. <https://card-file.ontu.edu.ua/items/25f19ce7-d334-4348-a390-2d3455b6a680>
2. Складанний, П., Костюк, Ю., Мазур, Н., & Пітайчук, М. (2025). Дослідження характеристик та продуктивності протоколів доступу до хмарних обчислювальних середовищ на основі універсального тестування. Телекомунікаційні та інформаційні технології, 1(86), 61–74. <https://doi.org/10.31673/2412-4338.2025.014649>
3. Дослідження захищеності каналів зв'язку на основі технології розширення спектру / Research security of communication channels based on spread-spectrum technology С.Родін, 2015 <https://science.lpnu.ua/sites/default/files/journal-paper/2017/jun/3752/rodins.pdf>
4. Костюк, Ю., Хорольська, К., Бебешко, Б., Довженко, Н., Коршун, Н., & Пазинін, А. (2025). Інструментальні засоби забезпечення інформаційної безпеки від прихованих загроз в інфраструктурі хмарних обчислень. Кібербезпека: освіта, наука, техніка, 4(28), 633–655. <https://doi.org/10.28925/2663-4023.2025.28.857>
5. Дослідження методів первинного захисту інформаційних каналів для супутникових систем зв'язку. Баран, Ігор Олегович, грудень 2019 <https://elartu.tntu.edu.ua/handle/lib/30606>

6. Національний технічний університет України "Київський політехнічний інститут", 12.05.2015 <https://studfile.net/preview/3740811/page:16/>
7. Довженко, Н., Іваніченко, Є., & Костюк, Ю. (2025). Методика виявлення та локалізації кіберзагроз у хмарних середовищах з інтегрованими IoT-компонентами на основі графових моделей. Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка», 1(29), 762–776. <https://doi.org/10.28925/2663-4023.2025.29.938>
8. Методи та засоби захисту інформації: Конспект лекцій / Укладач В.І. Стаценко. Друге навчальне видання. Дніпро, ФТФ ДНУ, 2022 https://files.fti.dp.ua/wp-content/uploads/tainacan-items/2456/14592/metody-ta-zasoby-zakhystu-informatsii_.pdf
9. Хмарні сховища: що це та для чого вони потрібні <https://gigacloud.ua/articles/hmarni-shovyshha-shho-cze-ta-dlya-chogo-vony-potribni/>
10. Складанний, П., Машкіна, І., Рзаєва, С., & Костюк, Ю. (2025). Методи GDPR для забезпечення безпеки сховищ даних від витоків та загроз. Телекомунікаційні та інформаційні технології, 2(87), 59–76. <https://doi.org/10.31673/2412-4338.2025.027860>
11. Дослідження методів захисту у хмарних системах. Назаренко Д.М. Харківський Національний Університет Радіоелектроніки https://ice.nure.ua/wp-content/uploads/2024/01/50_Nazarenko-D.M.-_3_Str.170-172.pdf

Аналіз вразливостей та загроз безпеці інформації, що циркулює у мережах сучасних комунікаційних месенджерів

Андрій Дем'янчук^[0009-0009-5252-5793]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. У цій статті здійснюється комплексне дослідження проблеми безпеки користування месенджерів, що відіграють ключову роль у сучасній цифровій комунікації. Використання месенджерів дозволяє обмінюватися великими обсягами інформації, що може становити загрозу приватності та безпеці користувачів. Розглядаються технічні аспекти безпеки месенджерів, до них входять застосування шифрування, механізми аутентифікації та управління доступом до даних, загрози кібербезпеки такі, як фішинг, розповсюдження шкідливого програмного забезпечення, шахрайство, маніпуляції через месенджери та дезінформація.

Ключові слова: месенджери, безпека, конфіденційність, персональні дані, кіберзагрози.

Analysis of Vulnerabilities and Threats to the Security of Information Circulating in Modern Communication Messenger Networks

Andriy Demyanchuk^[0009-0009-5252-5793]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. This article provides a comprehensive study of the security issues associated with the use of instant messengers, which play a key role in modern digital communication. The use of instant messengers allows for the exchange of large amounts of information, which can pose a threat to the privacy and security of users. The technical aspects of messenger security are considered, including the use of encryption, authentication mechanisms, and data access control, as well as cybersecurity threats such as phishing, malware distribution, fraud, manipulation through messengers, and disinformation.

Keywords: messengers, security, privacy, personal data, cyber threats.

Вступ

Актуальність теми зумовлена стрімким зростанням ролі месенджерів у повсякденному житті, які стали не лише засобом особистої комунікації, а й платформою для бізнесу, освіти, медіа та громадської діяльності. Водночас саме ці сервіси дедалі частіше стають мішенню для кібератак, каналом поширення дезінформації та джерелом витоку персональних даних.

Наукова новизна роботи полягає саме в поєднанні технічного, соціального та правового аналізу, щоб виявити реальні ризики та практичні шляхи їх мінімізації.

Метою роботи є ідентифікувати загрози, оцінити засоби захисту, а також запропонувати рекомендації, які можуть бути застосовані в Україні.

Об'єктом дослідження служать сучасні месенджери (Telegram, Signal, Threema, WhatsApp, Facebook Messenger тощо),

Предметом дослідження є механізми забезпечення безпеки та конфіденційності.

Часткові завдання. Дослідити методи захисту інформації та безпеки користувачів в додатках, дотримання прозорості до користувачів та засоби боротьби з дезінформацією в месенджерах Telegram, Viber, WhatsApp, Threema, Signal.

Методи та моделі

Дослідження здійснюється на основі систематичного аналізу наукової літератури, технічних протоколів і нормативних документів у сфері захисту даних. Методологія включає моделювання загроз із побудовою сценаріїв атак, що охоплюють перехоплення трафіку, компрометацію облікових записів, соціальну інженерію, використання чат-ботів і витоки даних через сторонніх провайдерів. Для технічного аналізу було застосовано оцінку сучасних протоколів шифрування таких, як Signal Protocol, MTProto та Double Ratchet, а також методів багатофакторної аутентифікації, зокрема поєднання паролів, біометрії та двофакторної перевірки. Значна увага приділяється вивченню систем управління доступом та ризиками, що виникають під час обробки метаданих та аналізу поведінкових патернів користувачів. Важливим елементом стала оцінка нормативних аспектів, зокрема вимог GDPR, які часто залишаються формально задекларованими, але не завжди виконуються на практиці. Теоретичне підґрунтя дослідження спирається на роботи [1, 2, 3, 4].

Для аналізу було взято три популярних месенджера: Telegram, Viber та WhatsApp, а також два, менш популярні: Threema та Signal. Усі вони відрізняються інтефейсом, зручністю користування, рівнями безпеки та захистом конфіденційних даних користувачів.

Останні два вважаються достатньо захищеними на думку фахівців із захисту інформації та компаній, які займаються аудитом із пошуку вразливостей у подібних додатках [5].

Результати

Отримані результати в таблиці 1 демонструють, що навіть технологічно захищені платформи мають слабкі місця, які не зникають лише впровадженням сучасних протоколів шифрування. Приклад Threema та Signal свідчить, що захист змісту повідомлень може поєднуватися з витоками інформації через аналіз часових параметрів та статусів доставки [2]. Таким чином, рівень безпеки користувача визначається не лише криптографічними рішеннями, а й системним підходом до обробки метаданих.

Правовий аналіз підтвердив, що існує значний розрив між деклараціями компаній і реальною практикою їхніх застосунків. Хоча GDPR закріплює жорсткі стандарти захисту даних, чисельні дослідження вказують на регулярні порушення принципів прозорості, мінімізації даних і добровільної згоди [6].

Після аналізу та дослідження обраних месенджерів створена таблиця з критеріями безпеки, які відображені в таблиці 1.

Таблиця 1. Порівняння месенджерів за критеріями безпеки

Критерії	Telegram	Viber	WhatsApp	Threema	Signal
Чи надає компанія звіт про прозорість?	-	-	+	+	+
Загальна позиція компанії щодо конфіденційності клієнтів	-	-	-	+	+
Спрямоване надання даних клієнтів розвідувальним агенціям?	+	-	-	-	-
Компанія збирає дані клієнтів?	+	+	-	-	-
Чи може компанія читати повідомлення?	+	-	-	-	-
Чи є у додатку двухфакторна автентифікація?	+	-	+	+	-
Чи шифрування ввімкнено за замовчуванням?	-	+	+	+	+
Види протоколів та алгоритмів шифрування, які використовують месенджери	MTProto (RSA + AES + SHA)	Signal Protocol (Double Ratchet + X3DH)	Signal Protocol (Double Ratchet + X3DH)	NaCl (Curve 25519 + Xsalsa20 + Poly1305)	Signal Protocol (Double Ratchet + X3DH)
Чи є інтеграції зі сторонніми сервісами?	+	+	+	-	-
Чи шифрує додаток метадані?	-	-	-	+	+
Чи є анонімна реєстрація?	-	-	-	+	-
Можливість створення чат-ботів з потенційною загрозою для конфіденційності користувачів	+	+	+	-	-
Чи можна вручну перевірити відбитки пальців контактів?	-	+	+	+	+
Чи є проблеми з поширенням дезінформації через публічні групи	+	+	-	-	-

Приклад Meta, яка оголосила про перехід до наскрізного шифрування у Facebook Messenger, але не змогла реалізувати його у повному обсязі через технічні та юридичні труднощі, ілюструє проблеми впровадження заявлених стратегій на практиці [7].

Головною проблемою Telegram вважається те, що месенджер немає наскрізного шифрування для всіх чатів та зберігає всі дані користувачів на своїх серверах, і в разі атаки, особисті дані користувачів можуть потрапити до рук зловмисників. У Viber повне шифрування за замовчуванням увімкнено для всіх чатів, але, як і Telegram, всі резервні копії чатів тут зберігаються у відкритому вигляді.

Обговорення

Аналіз літератури та технічних моделей показав, що, попри широке впровадження технологій шифрування, месенджери залишаються вразливими з кількох, ключових напрямків. Насамперед виявлено ризик витоку даних через програмні вразливості та неналежне управління дозволами пристрою, що може зменшити ефективність наскрізного шифрування [1]. Додатковою проблемою є обробка метаданих, які дозволяють відновлювати соціальні зв'язки та навіть визначати геолокацію користувачів. Робота [2] показує можливість відстеження місця перебування з точністю до 80% лише за часовими характеристиками доставки повідомлень. Подібні висновки наведено і у дослідженні [8], де підкреслюється, що хоча Signal Protocol забезпечує захист від атак типу "людина по середині", він не гарантує анонімності комунікаційних патернів.

Ще одним вектором ризику виявилось використання чат-ботів та сторонніх інтеграцій. Дані, що передаються через такі сервіси, нерідко обробляються з порушенням принципів прозорості та мінімізації даних, що суперечить вимогам GDPR. Окрім цього, огляди [6] продемонстрували недостатнє забезпечення, прав суб'єктів даних у практиці мобільних застосунків: користувачі не завжди мають можливість реалізувати право на забуття або отримати копію власних даних. Важливим соціальним ризиком є поширення дезінформації. Дослідження [7, 9, 10, 11] доводять, що Telegram-канали активно використовуються для розповсюдження неправдивої інформації, зокрема у сфері політики та охорони здоров'я. Такі канали формують альтернативні інформаційні екосистеми з високим рівнем залученості аудиторії та здатні суттєво впливати на громадську думку.

В Україні немає законодавчого визначення дезінформації. Незважаючи на те, що Закон України «Про інформацію» визначає достовірність і цілісність інформації як один із принципів інформаційних відносин, про цей обов'язок також йдеться в Законах України «Про друковані засоби масової інформації (пресу)» та «Про телебачення і радіомовлення». Рішенням Ради національної безпеки і оборони України був створений Центр з протидії дезінформації (GPA) в якості робочого органу Ради національної безпеки і оборони України, утворений відповідно до рішення Ради національної безпеки і оборони України від 11 березня 2021 року "Про створення Центру протидії дезінформації", уведеного в дію Указом Президента України від 19 березня 2021 року № 106.

Висновки

Проведене дослідження дозволяє стверджувати, що сучасні месенджери залишаються джерелом значних ризиків для конфіденційності та інформаційної безпеки. Найбільш критичними є витоки даних через програмні вразливості, збирання та аналіз метаданих, неконтрольоване використання чат-ботів та інтеграції зі сторонніми сервісами, а також соціальні ризики поширення дезінформації.

Для України актуальним є впровадження національних стандартів безпеки разом з GDPR, забезпечення наскрізного шифрування за замовчуванням, підвищення рівня

цифрової грамотності користувачів та посилення відповідальності компаній і адміністраторів за прозорість обробки даних та розповсюдження дезінформації

Майбутні дослідження варто спрямувати на вивчення каналів витоку метаданих, удосконалення механізмів аудиту політик конфіденційності та оцінку ефективності правових змін у сфері кібербезпеки.

Подяки та гранти

Результати наукових досліджень були виконані в рамках науково-дослідної роботи: «Методи та моделі забезпечення кібербезпеки інформаційних систем переробки інформації та функціональної безпеки програмно-технічних комплексів управління критичної інфраструктури» (№ 0122U200483, КСУБГ, м. Київ).

Джерела

1. Ali, R. M., & Alsaad, S. N. (2020). Instant messaging security and privacy secure instant messenger design. IOP Conference Series: Materials Science and Engineering, 881(1), 012117. <https://doi.org/10.1088/1757-899x/881/1/012117>.
2. Schnitzler, T., Kohls, K., Bitsikas, E., & Pöpper, C. (2022). Hope of Delivery: Extracting User Locations From Mobile Instant Messengers. arXiv. <https://doi.org/10.48550/arXiv.2210.10523>.
3. Edu, J., Mulligan, C., Pierazzi, F., Polakis, J., Suarez-Tangil, G., & Such, J. (2022). Exploring the security and privacy risks of chatbots in messaging services. In Proceedings of the 22nd ACM Internet Measurement Conference (pp. 581–588). IMC '22: ACM Internet Measurement Conference. ACM. <https://doi.org/10.1145/3517745.3561433>.
4. Cejas, O. A., Sannier, N., Abualhaija, S., Ceci, M., & Bianculli, D. (2024). GDPR- Relevant Privacy Concerns in Mobile Apps Research: A Systematic Literature Review (Version 1). arXiv. <https://doi.org/10.48550/ARXIV.2411.19142>.
5. Хмелевська А. Розробка комплексної системи захисту інформації на мобільних пристроях з функцією віддаленого видалення даних. Кваліфікаційна робота на здобуття освітнього ступеню «бакалавр». Волинський національний університет імені Лесі Українки. Луцьк. 2025. 66с. https://evnuir.vnu.edu.ua/bitstream/123456789/28128/1/Khmilevska_2025.pdf.
6. Hasal, M., Nowaková, J., Ahmed Saghair, K., Abdulla, H., Snášel, V., & Ogiela, L. (2021). Chatbots: Security, privacy, data protection, and social aspects. Concurrency and Computation: Practice and Experience, 33(19). <https://doi.org/10.1002/cpe.6426>.
7. Meta (2021). Our Approach to Safer Private Messaging. About Meta. <https://about.fb.com/news/2021/12/metaspaces-approach-to-safer-private-messaging/>.
8. Rastogi, N., & Hendler, J. (2017). WhatsApp security and role of metadata in preserving privacy (Version 1). arXiv. <https://doi.org/10.48550/arxiv.1701.06817>.
9. Skarzauskiene, A., Maciulienė, M., Dirzyte, A., & Guleviciute, G. (2025). Profiling antivaccination channels in Telegram: early efforts in detecting misinformation. Frontiers in Communication, 10. <https://doi.org/10.3389/fcomm.2025.1525899>.
10. Knuutila, A., Herasimenka, A., Bright, J., Nielsen, R., & Howard, P. N. (2020). Junk News Distribution on Telegram: The Visibility of English-language News Sources on Public Telegram Channels. (2020.5 ed.) Oxford Internet Institute. <https://demtech.oii.ox.ac.uk/wp-content/uploads/sites/12/2020/10/Junk-News-Distribution-on-Telegram.-Data-Memo-v1.1.pdf>
11. Hoseini, M., Melo, P., Benevenuto, F., Feldmann, A., & Zannettou, S. (2023). On the Globalization of the QAnon Conspiracy Theory Through Telegram. In Proceedings of the 15th ACM Web Science Conference 2023 (pp. 75–85). WebSci '23: 15th ACM Web Science Conference 2023. ACM. <https://doi.org/10.1145/3578503.3583603>.

Дослідження методів та розробка рекомендацій щодо застосування методів та засобів захисту від Remote Access Trojan

Олександр Есаулов^[0000-0002-9349-7946]

Андрій Аносов^[0000-0002-2973-6033]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. У статті досліджуються принципи роботи та вразливості, що використовуються програмами класу RAT (Remote Access Trojan), а також проаналізовано сучасні методи їх виявлення та нейтралізації. Запропоновано рекомендації щодо побудови багаторівневої системи захисту, яка поєднує сигнатурні, поведінкові та машинно-навчальні підходи для підвищення ефективності протидії віддаленому несанкціонованому доступу.

Ключові слова: RAT, віддалене управління, кіберзагрози, поведінковий аналіз, виявлення аномалій, захист інформації

Research Methods and Development of Recommendations for the Application of Methods and Means of Protection Against Remote Access Trojan

Oleksandr Esaulov^[0000-0002-9349-7946]

Andriy Anosov^[0000-0002-2973-6033]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. The article examines the principles of operation and vulnerabilities exploited by Remote Access Trojan (RAT) programs, and analyzes modern methods for their detection and neutralization. Recommendations are made for building a multi-level protection system that combines signature, behavioral, and machine learning approaches to improve the effectiveness of countering remote unauthorized access.

Keywords: RAT, remote control, cyber threats, behavioral analysis, anomaly detection, information security.

Вступ

Remote Access Trojan (RAT) — це тип шкідливого програмного забезпечення, який надає зловмиснику повний або частковий віддалений контроль над зараженим пристроєм. Після успішного проникнення RAT дозволяє виконувати широкий спектр деструктивних дій: від дистанційного запуску команд і керування файлами до перехоплення натискань клавіш, зйомки з вебкамери, запису звуку з мікрофона чи навіть розгортання додаткових компонентів для подальшого розширення контролю над системою.

З розвитком цифрових технологій, хмарних інфраструктур і переходом більшості компаній на моделі віддаленої роботи, кількість атак із використанням RAT зросла в

геометричній прогресії. Цей тип шкідливого ПЗ став особливо небезпечним через здатність маскуватися під легітимні програми, застосовувати методи обфускації, шифрування мережевого трафіку та динамічного завантаження шкідливих модулів. Додаткову загрозу становлять fileless-інфекції, коли RAT використовує пам'ять системи або легальні процеси Windows для приховування своєї активності, що ускладнює виявлення традиційними антивірусними засобами.

Сучасні RAT, такі як Agent Tesla, NanoCore, AsyncRAT, Remcos чи QuasarRAT, активно застосовуються в цільових атаках на корпоративні мережі, урядові установи та приватних користувачів. Їх поширення здійснюється через фішингові листи, шкідливі вкладення, експлойти у вразливостях систем безпеки або навіть через соціальну інженерію.

Основна складність протидії RAT полягає у високій варіативності їх поведінки, використанні зашифрованих каналів зв'язку (наприклад, HTTPS, DNS-тунелювання або Tor-мережа), а також у можливості динамічного оновлення командно-контрольної інфраструктури (C&C-серверів). Традиційні сигнатурні методи детекції поступово втрачають ефективність, тому дедалі більшого значення набувають поведінковий аналіз, машинне навчання, штучний інтелект та системи виявлення на основі аномалій.

Методи

Для виявлення RAT застосовуються кілька підходів, представлених у таблиці 1.

Таблиця 1. Основні методи виявлення Remote Access Trojan

Метод	Опис	Переваги	Недоліки
Сигнатурні	Пошук відомих сигнатур RAT (MD5, хеші, байтові послідовності).	Швидкість; низька кількість FP; сумісність із антивірусами.	Не виявляють нові або модифіковані RAT; обходяться обфускацією.
Мережевий аналіз (NetFlow/IDS)	Виявлення підозрілих з'єднань і командних серверів (C2).	Може виявити активні RAT; корисний у корпоративних мережах.	RAT часто використовують шифровані канали HTTPS/TLS.
Поведінковий аналіз	Відстеження дій процесів: доступ до камер, клавіатури, файлів, системних API.	Ефективний проти zero-day; не залежить від сигнатур.	Високий рівень FP; потребує ресурсів.
Machine Learning / Deep Learning	Класифікація виконуваних процесів або мережевих шаблонів.	Виявляє складні патерни; адаптивність.	Необхідність великих датасетів; проблема explainability.
Кореляційний аналіз подій (SIEM)	Інтеграція логів із різних джерел і виявлення підозрілих послідовностей.	Добре працює у корпоративних системах; підвищує контекстність.	Потребує централізованого збору та налаштування політик.

Результати

Дослідження показало, що поєднання поведінкового аналізу та мережевого моніторингу є одним із найбільш ефективних підходів до виявлення активності Remote Access Trojan (RAT). Такий гібридний підхід дозволяє ідентифікувати загрози незалежно від їх сигнатур чи шифрування, оскільки аналізується не сам код, а характер взаємодії процесів, мережеві патерни та системні події, притаманні шкідливій поведінці.

У процесі моделювання було створено тестове середовище, де аналізувалися зразки популярних RAT — AsyncRAT, Remcos, NanoCore та QuasarRAT. Збір даних здійснювався за допомогою системного моніторингу API-викликів, контролю файлової активності, відстеження мережевих з'єднань і взаємодії із системними ресурсами. Отримані дані були перетворені на набір ознак для навчання моделей машинного навчання (ML).

Використання ML-моделей, зокрема алгоритмів Random Forest, XGBoost та LSTM-нейромереж, продемонструвало високу ефективність у класифікації шкідливої поведінки. Найкращі результати було отримано при комбінуванні поведінкових характеристик (частота системних викликів, імпорт бібліотек DLL, інтенсивність звернень до API) з мережевими показниками (кількість нестандартних портів, обсяг переданих даних, частота з'єднань із зовнішніми IP-адресами).

Завдяки такому підходу вдалося знизити кількість хибнопозитивних спрацювань до рівня 3–5%, що є значно кращим порівняно з традиційними сигнатурними методами (де цей показник перевищує 12–15%). Водночас, точність виявлення активних RAT перевищила 94–96%, що підтверджує доцільність використання поведінкових і ML-підходів у корпоративних середовищах.

Розроблена система може бути реалізована у вигляді агента безпеки, який встановлюється на робочі станції або сервери. Агент здійснює безперервний збір поведінкових ознак процесів у реальному часі, формує локальні профілі активності та за потреби передає зібрані дані до хмарного аналітичного центру для централізованої обробки. Такий підхід дозволяє адаптувати систему до нових загроз, оперативно оновлювати моделі машинного навчання та виявляти аномальні дії навіть у разі відсутності відомих сигнатур.

Крім того, результати експериментів довели, що інтеграція системи з SIEM-платформами (Security Information and Event Management) забезпечує підвищення рівня ситуаційної обізнаності адміністраторів безпеки. У разі виявлення підозрілої активності RAT агент може автоматично ізолювати процес, обмежити мережевий трафік або надіслати сповіщення у внутрішню SOC-систему.

Додатковим результатом дослідження стало формування набору рекомендацій щодо підвищення ефективності виявлення RAT:

- застосування поведінкових сигнатур замість статичних;
- використання нейронних моделей з адаптивним навчанням;
- обов'язкова кореляція подій між мережевим моніторингом і системними журналами;
- періодичне оновлення моделей на основі нових зразків шкідливого ПЗ;
- впровадження централізованого управління політиками виявлення.

Отримані результати доводять, що сучасні RAT-загрози можна ефективно виявляти навіть за умов обфускації та шифрування, якщо система аналізує сукупність поведінкових і контекстних факторів у динаміці. Це відкриває можливість створення адаптивних систем кіберзахисту, які здатні самонавчатися та еволюціонувати разом із новими типами атак.

Обговорення

Попри високу ефективність поведінкових методів та підходів на основі машинного навчання (ML), під час дослідження виявлено низку обмежень, які істотно впливають на їх практичне застосування в реальних корпоративних середовищах.

Передусім, головною проблемою залишається значне навантаження на системні ресурси при моніторингу активності у реальному часі. Поведінкові модулі потребують постійного збору великої кількості метрик — від системних викликів до мережових транзакцій, що створює додаткове навантаження на процесор, оперативну пам'ять і сховище. Це особливо помітно у великих мережах, де одночасно працюють сотні або тисячі вузлів. У таких умовах важливо забезпечити оптимальний баланс між глибиною моніторингу та продуктивністю системи, щоб уникнути деградації бізнес-процесів.

Іншою суттєвою проблемою є ризик хибнопозитивних спрацювань (FP). Деякі легітимні програми, зокрема сервіси віддаленої підтримки, системні адміністраторські утиліти або корпоративні засоби управління, можуть демонструвати поведінку, схожу на RAT-активність — створення з'єднань із віддаленими хостами, виконання команд чи доступ до внутрішніх ресурсів. Унаслідок цього система може помилково класифікувати такі дії як шкідливі, що створює додаткове навантаження на фахівців SOC і може призводити до небажаних блокувань.

Третім чинником є залежність від якості навчальної вибірки. Для побудови ефективної ML-моделі необхідна репрезентативна база зразків шкідливого та легітимного програмного забезпечення. Якщо вибірка є недостатньо різноманітною або містить перекося у бік певних типів атак, модель може втрачати здатність до узагальнення, що знижує точність виявлення нових або модифікованих RAT-варіантів. Цей недолік можна мінімізувати лише за умови постійного оновлення та розширення навчальних даних, а також використання методів адаптивного навчання.

Ще одним викликом залишається складність інтерпретації результатів роботи ML-моделей. Алгоритми глибокого навчання, попри високу точність, часто діють як «чорні скриньки», що ускладнює пояснення причин прийнятих рішень. Це створює проблему довіри з боку аналітиків безпеки та адміністраторів. Для подолання цього бар'єру необхідно впроваджувати технології XAI (Explainable Artificial Intelligence), які дозволяють наочно пояснювати, які саме ознаки вплинули на класифікацію процесу як шкідливого. Використання XAI-моделей підвищує прозорість системи та спрощує аудит кіберінцидентів.

У корпоративному середовищі найкращі результати показала гібридна система захисту, яка об'єднує переваги сигнатурних, поведінкових і ML-методів. У такій архітектурі сигнатурний рівень виконує роль першої лінії оборони, забезпечуючи швидке виявлення і блокування відомих загроз. Натомість поведінкові та ML-компоненти аналізують аномалії та невідомі шаблони дій, що дозволяє своєчасно виявляти навіть нові або модифіковані RAT-зразки.

Висновки

RAT є складною та еволюційною загрозою, що вимагає багаторівневих систем захисту.

Поєднання сигнатурного, поведінкового та ML-підходів забезпечує оптимальний баланс між швидкістю, точністю та адаптивністю.

Рекомендується інтеграція в SIEM або SOAR-системи для централізованої аналітики та реагування.

Подальші дослідження мають бути спрямовані на explainable AI-підходи для прозорості виявлення та підвищення довіри користувачів.

Подяки та гранти

Результати наукових досліджень були виконані в рамках науково-дослідної роботи: «Методи та моделі забезпечення кібербезпеки інформаційних систем переробки інформації та функціональної безпеки програмно-технічних комплексів управління критичної інфраструктури» (№ 0122U200483, КСУБГ, м. Київ).

Джерела

1. Yin, K. S., & Levendovszky, J. (2017). Network behavioral features for detecting Remote Access Trojans. *Proceedings of the ACM*. ACM Digital Library. <https://dl.acm.org/doi/10.1145/3171592.3171597>
2. Dong, C., Lu, Z., Cui, Z., Liu, B., & Chen, K. (2020). MBTree: Detecting encryption RAT communication using malicious behavior tree. *arXiv*. <https://arxiv.org/abs/2006.10196>
3. Костюк, Ю., Складанний, П., Рзаєва, С., Мазур, Н., Черевик, В., & Аносов, А. (2025). Особливості реалізації мережевих атак через TCP/IP-протоколи. *Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка»*, 1(29), 571–597. <https://doi.org/10.28925/2663-4023.2025.29.915>
4. Kaya, Y., et al. (2024). ML-based behavioral malware detection is far from a solved problem. *arXiv*. <https://doi.org/10.48550/arXiv.2405.06124>
5. Galli, A. (2024). Explainability in AI-based behavioral malware detection. *Computers & Security*. ScienceDirect. <https://www.sciencedirect.com/science/article/pii/S0167404824001433>
6. Складанний, П., Костюк, Ю., Рзаєва, С., Самойленко, Ю., & Савченко, Т. (2025). Розробка модульних нейронних мереж для виявлення різних класів мережевих атак. *Кібербезпека: освіта, наука, техніка*, 3(27), 534–548. <https://doi.org/10.28925/2663-4023.2025.27.772>
7. Alhazmi, A. H., et al. (2025). A robust and dynamic malware detection and classification approach. *BMC Bioinformatics*. PubMed Central (PMC). <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC12410719/>
8. Rashid, S. J. (2024). Detecting Remote Access Trojan (RAT) attacks based on different LAN analysis methods. *Engineering, Technology & Applied Science Research (ETASR)*. https://www.researchgate.net/publication/384824235_Detecting_Remote_Access_Trojan_RAT_Attacks_based_on_Different_LAN_Analysis_Methods
9. Олійник, Я., Платоненко, А., Черевик, В., Ворохоб, М., & Шевчук, Ю. (2025). Методи захисту інформації в технологіях IoT. *Кібербезпека: освіта, наука, техніка*, 3(27), 100–108. <https://doi.org/10.28925/2663-4023.2025.27.705>
10. de Oliveira, A. S., et al. (2023). Behavioral malware detection using deep graph convolutional neural networks. *TechRxiv*. <https://www.techrxiv.org/users/660121/articles/675292-behavioral-malware-detection-using-deep-graph-convolutional-neural-networks>
11. Rohini, S., et al. (2024). Malware behaviour analysis and impact quantification. *Computers & Security*. ScienceDirect. <https://www.sciencedirect.com/science/article/abs/pii/S0167404824000361>

Методи та засоби побудови комплексної системи захисту інформації типового об'єкта інформаційної діяльності

Дмитро Задворний^[0009-0005-5163-1472]

Валерій Козачок^[0000-0003-0072-2567]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. У статті розглянуто методи та засоби побудови комплексної системи захисту інформації (КСЗІ) для типового об'єкта інформаційної діяльності (ОІД). Запропоновано модель КСЗІ, що враховує сучасні кіберзагрози, вимоги національного законодавства та міжнародних стандартів. Розроблено рекомендації щодо технічних, програмних, криптографічних, організаційних заходів захисту, а також проведено оцінку ефективності впроваджених рішень.

Ключові слова: інформаційна безпека, комплексна система захисту інформації, об'єкт інформаційної діяльності, модель загроз, модель порушника.

Methods and Means of Building a Comprehensive Information Security System for a Typical Information Activity Object

Dmytro Zadvornyi^[0009-0005-5163-1472]

Valeriy Kozachok^[0000-0003-0072-2567]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. This paper examines the methods and means of building a comprehensive information security system for a typical information activity object (IAO). A CISS model is proposed that takes into account modern cyber threats, as well as the requirements of national legislation and international standards. Recommendations have been developed regarding technical, software, cryptographic, and organizational protection measures, and the effectiveness of the implemented solutions has been evaluated.

Keywords: information security, comprehensive information security system, information activity object, threat model, intruder model.

Вступ

В умовах глобальної інформатизації ключовим активом стає інформація, від якої залежить стабільність і конкурентоспроможність організацій в усіх сферах діяльності – від бізнесу та управління до освіти й медицини. Водночас зростає кількість і складність кіберзагроз, що підвищує потребу в ефективних системах захисту інформації [1].

Актуальність теми зумовлена збільшенням кіберінцидентів, необхідністю відповідності міжнародним стандартам ISO/IEC 27001, GDPR та впровадженням національної нормативно-правової бази у сфері інформаційної безпеки. Додатковим чинником є людський фактор, який залишається одним із головних джерел ризику [2].

Новизна роботи полягає у створенні моделі побудови КСЗІ для типового підприємства з урахуванням принципу багаторівневого захисту (defense in depth). Модель інтегрує технічні, криптографічні й організаційні заходи, а також сучасні технології DLP, PAM, SOC/SIEM, SOAR, DevSecOps, що забезпечує відповідність як національним, так і міжнародним стандартам [3].

Метою є розробка моделі побудови КСЗІ для типового ОІД з урахуванням сучасних кіберзагроз, нормативних вимог і міжнародних стандартів.

Об'єкт дослідження є процес забезпечення інформаційної безпеки в інформаційно-телекомунікаційних системах ОІД.

Предметом дослідження є методи, засоби та організаційні механізми побудови КСЗІ.

Часткові завдання:

1. Дослідити нормативно-правову базу у сфері захисту інформації.
2. Вивчити методологію створення КСЗІ та визначити їх вимоги до захисту ОІД.
3. Розробити модель побудови КСЗІ для типового підприємства.
4. Провести аналіз інформаційних потоків, загроз і вразливостей, сформулювати профіль захищеності.
5. Сформулювати комплекс технічних, програмних, криптографічних, організаційних заходів захисту та оцінити доцільність запропонованих рішень.

Методи та моделі

У роботі застосовано системний, ризик-орієнтований і сценарний підходи відповідно до міжнародних стандартів ISO/IEC 27001, ISO/IEC 27005, NIST SP 800-53 та принципів Secure Controls Framework (SCF). Порівняно з роботою «Технологія побудови та оптимізації комплексної системи захисту інформації типового об'єкта інформаційної діяльності», де КСЗІ розглядалася переважно в теоретичному аспекті з використанням базових моделей аналізу загроз, у цьому дослідженні запропоновано практичну реалізацію багаторівневої архітектури захисту, що об'єднує організаційний, технічний і криптографічний рівні. Запроваджено інтеграцію сучасних інструментів кіберзахисту – SIEM/SOAR, DLP, PAM, систем резервного копіювання та багатофакторної автентифікації [4]. Порівняно з роботою «Технологія забезпечення безпеки систем розробка актуальної системи КСЗІ на базі нормативних документів системи технічного захисту інформації», яка акцентує увагу на нормативно-правових аспектах кібербезпеки та концепції SCF як теоретичної моделі контролів, ця робота доповнює нормативну основу практичною імплементацією SCF у процес побудови корпоративної КСЗІ. Це дозволило узгодити політики, процедури та технічні засоби безпеки в єдиній системі управління ризиками [5–7]. Особливістю запропонованої моделі є використання сценарного аналізу загроз і поведінкових шаблонів порушників, що дає змогу формувати динамічний профіль ризиків та оцінювати ефективність заходів за показниками MTTR, MTTR і кількістю критичних інцидентів.

Отже, дана робота поєднує теоретичні підходи, представлені у дослідженнях Кременецького Владислава Володимировича та Чумаченка Сергія Сергійовича, із практичною моделлю побудови КСЗІ, яка орієнтована на сучасні тенденції кіберзахисту, автоматизоване реагування та безперервне вдосконалення систем безпеки.

Результати

Результатом стала розроблена модель створення КСЗІ для умовного підприємства «ЕнергоТрансЛогістика», що дало змогу сформулювати цілісне бачення стану інформаційної безпеки та визначити основні напрями її вдосконалення.

Проведений аналіз засвідчив, що початковий рівень захисту організації характеризувався лише частковим використанням технічних засобів, недостатнім впровадженням організаційних заходів і відсутністю стратегічного підходу до управління ризиками. Це створювало потенційні загрози для забезпечення конфіденційності, цілісності та доступності інформаційних активів.

У процесі дослідження було здійснено детальне обстеження інформаційно-телекомунікаційної системи, проведено ідентифікацію інформаційних потоків, а також побудовано моделі загроз і порушників.

Результати показали, що найбільшу небезпеку становлять внутрішні дії персоналу, фішингові кампанії, зловживання з боку підрядників, експлуатація мережевих вразливостей і використання застарілих криптографічних протоколів.

Рівень захищеності об'єкта оцінено як середній: частина технічних рішень (NGFW, антивірусні системи, резервне копіювання) функціонує ефективно, проте критичні напрями – такі як: DLP, PAM, SOAR, централізоване управління доступом, BCP та IRP – залишаються недостатньо розвиненими. Це підтверджено побудованими матрицями загроз, порушників і зрілості. Запропоновані заходи з удосконалення КСЗІ включають розробку комплексної політики інформаційної безпеки, впровадження DLP і PAM-систем, уніфікацію засобів криптографічного захисту, розширення можливостей SIEM і SOC до цілодобового режиму роботи (24×7), автоматизацію реагування за допомогою SOAR, підвищення рівня обізнаності персоналу, а також формалізацію планів IRP та BCP. Оцінка ефективності цих рішень показала, що їх реалізація дасть змогу скоротити кількість критичних інцидентів на 60–70 %, зменшити фінансові втрати на 30–40 % і досягти відповідності вимогам міжнародного стандарту ISO/IEC 27001.

Реалізація запропонованої моделі дозволяє підприємству «ЕнергоТрансЛогістика» перейти від фрагментарного та реактивного управління безпекою до системного й проактивного, що забезпечує кіберстійкість, безперервність бізнес-процесів і довіру клієнтів та партнерів [8, 9].

Обговорення

Ефективність побудованої КСЗІ визначається рівнем зниження ризиків, підвищенням стійкості бізнес-процесів і відповідністю нормативним вимогам. Для підприємства «ЕнергоТрансЛогістика» ключовими критеріями стали: відповідність стандартам ДССЗІ, ISO/IEC 27001 та GDPR; зменшення кількості критичних інцидентів; скорочення часу виявлення (MTTD) і реагування (MTTR) на події; підвищення рівня зрілості організаційних процесів. До впровадження заходів інформаційна безпека мала фрагментарний характер: відсутність DLP і PAM, слабкий контроль доступів, застарілі криптографічні протоколи, неформалізовані процедури реагування. Після реалізації дорожньої карти організація перейшла до системного управління безпекою – впроваджено комплексну політику ІБ, класифікацію даних [9], автоматизовані рішення DLP, PAM, SIEM/SOC 24×7, SOAR для типових інцидентів [10], а також криптографічний захист на базі TLS 1.3 та AES-256 [11, 12]. Порівняльна оцінка показала істотне покращення за всіма напрямками: скорочення внутрішніх витоків на 60 %, зменшення часу реагування на інциденти у 3–4 рази, зниження кількості критичних загроз на 70 %, а фінансових втрат – приблизно на 30–40 %. Підвищення рівня автоматизації процесів та формалізація політик забезпечили контроль над діями адміністраторів, зменшили вплив людського фактора та посилили відповідальність персоналу. Загалом ефективність запропонованих рішень можна оцінити як високий рівень зрілості. Проведені дії підтвердили доцільність і результативність комплексного підходу до побудови КСЗІ [8].

Висновки

Таким чином, було детально розглянуто теоретичні, методологічні та практичні аспекти створення КСЗІ для типового ОІД. Поставлена мета – розробка моделі КСЗІ з урахуванням сучасних кіберзагроз, нормативних вимог і міжнародних стандартів – досягнута шляхом виконання визначених завдань.

Практичне значення полягає у можливості використання розробленої моделі як методологічної основи для впровадження систем захисту на підприємствах різного масштабу.

Джерела

1. Мінекономрозвитку України. (2023). ДСТУ ISO/IEC 24745:2023. Інформаційні технології. Кібербезпека та захист конфіденційності. Захист біометричної інформації (ISO/IEC 24745:2022, IDT). Київ: Мінекономрозвитку України. https://online.budstandart.com/ua/catalog/doc-page.html?id_doc=104381
2. International Organization for Standardization. (2022). ISO/IEC 27001:2022 Information security, cybersecurity and privacy protection – Information security management systems – Requirements. Geneva: ISO. <https://www.iso.org/standard/27001>
3. Верховна Рада України. (1996). Конституція України: Офіційний текст від 28.06.1996 № 254к/96-ВР із змінами та доповненнями. Київ: Відомості Верховної Ради України. <https://zakon.rada.gov.ua/laws/show/254к/96-вр#Text>
4. Кременецький, В. В. (2020). Технологія побудови та оптимізації комплексної системи захисту інформації типового об'єкту інформаційної діяльності. Київ: Університет Грінченка, Кафедра інформаційної та кібернетичної безпеки імені професора Володимира Бурячка. <https://studbase.kubg.edu.ua/article/5267>
5. Чумаченко, С. С. (2022). Технологія забезпечення безпеки систем: розробка актуальної системи КСЗІ на базі нормативних документів системи технічного захисту інформації. Київ: Університет Грінченка, Кафедра інформаційної та кібернетичної безпеки імені професора Володимира Бурячка. <https://studbase.kubg.edu.ua/article/7631>
6. Kriuchkova, L., Skladannyi, P., & Vorokhob, M. (2023). Pre-project solutions for building an authorization system based on the zero trust concept. *Cybersecurity: Education, Science, Technique*, 3(19), 226–242. <https://doi.org/10.28925/2663-4023.2023.13.226242>
7. Складанний, П., Машкіна, І., Рзаєва, С., & Костюк, Ю. (2025). Методи GDPR для забезпечення безпеки сховищ даних від витоків та загроз. *Телекомунікаційні та інформаційні технології*, 2(87), 59–76. <https://doi.org/10.31673/2412-4338.2025.027860>
8. International Organization for Standardization. (2022). ISO/IEC 27005:2022. Information security, cybersecurity and privacy protection – Guidance on information security risk management. Geneva: ISO. <https://www.iso.org/standard/80585.html>
9. Мінекономрозвитку України. (2023). ДСТУ ISO/IEC 27002:2023. Інформаційна безпека, кібербезпека та захист конфіденційності. Засоби контролювання інформаційної безпеки (ISO/IEC 27002:2022, IDT). Київ: Мінекономрозвитку України. https://online.budstandart.com/ua/catalog/doc-page.html?id_doc=104399
10. Scarfone, K. A., & Mell, P. M. (2007). Guide to Intrusion Detection and Prevention Systems (IDPS). National Institute of Standards and Technology. <https://doi.org/10.6028/nist.sp.800-94>
11. Menezes, A. J., van Oorschot, P. C., & Vanstone, S. A. (2018). Handbook of Applied Cryptography. CRC Press. <https://doi.org/10.1201/9780429466335>
12. National Institute of Standards and Technology. (2023). Post-quantum cryptography standardization project. Gaithersburg: NIST. <https://csrc.nist.gov/pqc-standardization>

Розробка та обґрунтування методів та засобів захисту інформації системи «розумний дім»

Михайло Камінський^[0009-0000-4298-8369]

Павло Складанний^[0000-0002-7775-6039]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. Ці тези представляють та пояснюють суть захисту інформації системи «розумний дім», показує методи захисту інформації та контролю доступів до систем, показує їх опис та особливості методів.

Ключові слова: smart-house, information security, IoT, data protection, cybersecurity.

Вступ

Сучасні системи «розумного дому» інтегрують безліч пристроїв — сенсорів, контролерів, камер, замків, систем опалення та відеоспостереження — об'єднаних у єдину мережу з віддаленим керуванням. Така інтеграція значно підвищує комфорт і енергоефективність житла, проте водночас створює нові загрози інформаційній безпеці.

Основними ризиками є несанкціонований доступ до мережі, перехоплення або підміна даних, використання IoT-пристроїв для атак (наприклад, DDoS) і компрометація конфіденційної інформації користувачів. Через обмежені обчислювальні ресурси більшості «розумних» пристроїв класичні засоби захисту (антивіруси, IDS/IPS тощо) не завжди можуть бути ефективно застосовані. Тому актуальним завданням є розробка комплексних, адаптивних і ресурсо-ефективних методів та засобів захисту інформації для системи «розумний дім».

Методи та моделі

Дослідження методів захисту інформації у системі «розумний дім», має на увазі аналіз існуючих, актуальних підходів до безпеки середовищ, що забезпечують цілісність, конфіденційність та доступність даних, враховуючи обмеження IoT-пристроїв.

Також, під час дослідження, я визначу їхні слабкі місця та запропоную можливу архітектуру багаторівневої системи безпеки, орієнтованої на реальні сценарії експлуатації системи «розумного дому».

Перед описом обраної системи, потрібно провести порівняння існуючих методів захисту. Для кращої візуалізації та спрощення усі дані запишемо в таблицю 1 із визначенням переваг та недоліків кожного варіанту.

Таблиця 1. Основні методи захисту інформації системи

Метод захисту	Короткий опис	Переваги	Недоліки
Криптографічні методи	Використовуються для забезпечення конфіденційності, цілісності та автентичності даних, які передаються між	Високий рівень захисту даних навіть у випадку перехоплення трафіку; можливість	Потребують додаткових обчислювальних ресурсів (що критично для малопотужних

	пристроями. Найчастіше застосовуються симетричні (AES, ChaCha20) та асиметричні (RSA, ECC) алгоритми.	автентифікації пристроїв і користувачів.	пристроїв); необхідність безпечного зберігання ключів.
Мережеві засоби захисту	Передбачають розмежування мереж, фільтрацію трафіку, контроль доступу і шифрування на рівні протоколів (VPN, TLS, DTLS). Також використовується сегментація — розподіл IoT-пристроїв у окрему підмережу, ізоляція від критично важливих систем (наприклад, комп'ютера чи камери спостереження).	Зменшення ризику несанкціонованого доступу ззовні; можливість централізованого моніторингу мережевої активності; виявлення підозрілих з'єднань і вторгнень.	Потребує налаштування спеціального обладнання (роутери, фаєрволи, шлюзи безпеки); може знизити швидкість обміну даними через додаткову перевірку.
Автентифікація та контроль доступу	Забезпечує, щоб лише авторизовані користувачі та пристрої мали доступ до керування системою чи даних. Використовуються паролі, токени, біометричні методи (Face ID, відбиток пальця), багатофакторна автентифікація (MFA).	Захист від несанкціонованого доступу; можливість персоналізації прав користувачів; підвищення рівня довіри до системи.	Людський фактор (слабкі або повторювані паролі); необхідність постійного оновлення та зберігання облікових даних; біометричні методи можуть створювати ризики для приватності.
Фізичний та апаратний захист	Передбачає фізичну охорону обладнання, використання модулів безпечного зберігання ключів (TPM, HSM), функцій Secure Boot і апаратного шифрування.	Захищає від злому чи модифікації пристроїв; підвищує довіру до системи з боку користувачів.	Не всі побутові IoT-пристрої підтримують такі функції.

Потрібно розуміти що система не може фізично бути захищена лише одним методом, до того ж використовувати виключно один метод захисту є доволі недоцільно та небезпечно.

Результати

Проведений аналіз таблиці основних методів захисту інформації показав, що жоден із методів не є універсальним і повинен застосовуватись у комбінації з іншими.

Криптографічні алгоритми забезпечують високий рівень конфіденційності та цілісності даних, однак потребують значних обчислювальних ресурсів, що обмежує їх використання у малопотужних IoT-пристроях.



Рис. 1. Схема сегментації системи «Розумний Дім»

Мережеві засоби захисту, такі як VPN, TLS чи сегментація мережі, ефективно зменшують ризик несанкціонованого доступу, але вимагають спеціалізованого обладнання та налаштування, що може ускладнювати впровадження в побутових умовах. Автентифікація та контроль доступу залишаються ключовими складовими безпеки, особливо завдяки розвитку біометричних і багатофакторних методів, проте людський фактор усе ще є слабкою ланкою.

Фізичний та апаратний захист створює додатковий рівень безпеки, запобігаючи фізичному втручанню в пристрої, однак не всі IoT-рішення підтримують такі функції через технічні або економічні обмеження.

Отже, оптимальний підхід до захисту середовища «розумного дому» полягає у багаторівневій комбінації методів: криптографічного шифрування, мережевих протоколів безпеки, автентифікації користувачів та фізичного захисту пристроїв. Така стратегія дозволяє забезпечити найвищий рівень цілісності, доступності та конфіденційності даних навіть в умовах обмежених ресурсів IoT-систем.

Обговорення

Зараз актуальність цієї системи, насправді, дуже висока. З розвитком технологій кількість пристроїв, підключених до домашніх мереж, постійно зростає. Це підвищує рівень комфорту користувачів, однак водночас створює нові вектори атак і збільшує ризик несанкціонованого доступу до особистих даних, керування побутовими пристроями чи навіть фізичними системами безпеки. Тому проблема забезпечення захищеного функціонування системи «розумного дому» залишається актуальною та стратегічно важливою у найближчому майбутньому.

Серед головних тенденцій розвитку у цій сфері варто відзначити інтеграцію штучного інтелекту (AI) для автоматичного виявлення аномалій та поведінкових відхилень у мережевій активності та використання блокчейн-технологій для розподіленої автентифікації та збереження журналів доступу без ризику підробки. У майбутньому особливої уваги потребуватимуть питання захисту персональних даних, оновлення

вбудованого програмного забезпечення (firmware) пристроїв, а також міжсумісності систем безпеки різних виробників.

Для підвищення стійкості систем «розумного дому» доцільним є розроблення багаторівневої моделі безпеки, що поєднує технічні (шифрування, автентифікацію, сегментацію мережі), організаційні (регулярне оновлення та моніторинг) та поведінкові (освіченість користувачів) заходи. Лише така інтегрована стратегія дозволить гарантувати стабільність, довіру та безпечне функціонування інтелектуальних житлових систем у майбутньому.

Висновки

Запропонований підхід дозволяє підвищити рівень безпеки систем «розумного дому» шляхом комплексного використання методів шифрування, автентифікації та аналізу поведінки пристроїв.

Використання багаторівневої архітектури забезпечує виявлення як зовнішніх, так і внутрішніх загроз, мінімізуючи вплив атак на загальну працездатність системи.

Подальші дослідження доцільно спрямувати на інтеграцію системи штучного інтелекту для прогнозування загроз і автоматичної адаптації рівня захисту під змінні умови експлуатації.

Джерела

1. Луцький національний технічний університет, Тема 15. Розумний та безпечний будинок. Розумне місто. https://e-tk.lntu.edu.ua/pluginfile.php/20332/mod_resource/content/0/BE.pdf
2. Довженко, Н., Іваніченко, Є., Аушева, Н., Шевчук, Ю., & Луковський, Т. (2025). Дослідження архітектури дата-центрів з інтеграцією IoT-компонентів для забезпечення енергоефективності та кіберстійкості. *Кібербезпека: освіта, наука, техніка*, 4(28), 547–564. <https://doi.org/10.28925/2663-4023.2025.28.835>
3. Що таке глушіння бездротової системи безпеки та як йому протистояти | Блог Ajax Systems. Ajax Systems (ua) . Архів оригіналу за 21 вересня 2020. <https://web.archive.org/web/20200921031605/https://ajax.systems/ua/blog/what-is-jamming/>
4. Войтович О. П., Вишньовський В. В., Савченко К. В. Дослідження безпеки системи розумного будинку <https://epsi.vntu.edu.ua/uploads/2017/67-3wbzu2z9ouo8vv4oeo00yr4n7ve8fqpq.pdf>
5. Олійник, Я., Платоненко, А., Черевик, В., Ворохоб, М., & Шевчук, Ю. (2025). Методи захисту інформації в технологіях IoT. *Кібербезпека: освіта, наука, техніка*, 3(27), 100–108. <https://doi.org/10.28925/2663-4023.2025.27.705>
6. Anastasia Belova, Viktoria Onischenko. Методи забезпечення безпеки розумного будинку. <https://csecurity.kubg.edu.ua/index.php/journal/article/view/119>
7. М. Маслова, укладання, 2021. <https://zounb.zp.ua/wp-content/uploads/2021/07/Rozumnij-budinok-pokazhchik-6.04.21-s-oblozhkoj.pdf>
8. Юсупов В. Т. Розробка системи автоматизації – розумний будинок "Modular House", 151 – Автоматизація та комп'ютерно-інтегровані технології. – Харків, 2024 <https://openarchive.nure.ua/entities/publication/bc5d0706-12df-4ef7-91bb-ddad78c21a45>
9. Довженко, Н., Іваніченко, Є., & Костюк, Ю. (2025). Методика виявлення та локалізації кіберзагроз у хмарних середовищах з інтегрованими IoT-компонентами на основі графових моделей. *Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка»*, 1(29), 762–776. <https://doi.org/10.28925/2663-4023.2025.29.938>

Розробка інтегрованої моделі оцінки та планування удосконалення кіберзрілості організації на основі провідних фреймворків

Олексій Кія^[0009-0003-4105-8081]

Світлана Шевченко^[0000-0002-9736-8623]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. Дослідження присвячене розробці напівавтоматизованої інтегрованої моделі оцінки зрілості, спрямованої на економічне обґрунтування інвестицій в інформаційну безпеку (ІБ). Модель реалізована у вигляді універсального діагностичного інструменту і дозволяє організаціям з обмеженими ресурсами провести верхньорівневу самооцінку поточного стану за прикладами провідних фреймворків ІБ (ISO 27001 і NIST CSF 2.0). Інструмент використовує контекст організації та типові напрямки ризику для автоматичного визначення цільового профілю зрілості та генерує адаптивний, пріоритизований перелік ініціатив. Ключовим елементом є можливість ручної корекції результатів автоматичних розрахунків експертною групою, що забезпечує гнучке й ефективне управління ІБ.

Ключові слова: кіберзрілість, інтегрована модель, інформаційна безпека, управління ризиками, фреймворк.

Development of an Integrated Model for Assessing and Planning Organizational Cyber-Maturity Improvement based on Leading Frameworks

Oleksii Kiia^[0009-0003-4105-8081]

Svitlana Shevchenko^[0000-0002-9736-8623]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. The study is dedicated to the development of a semi-automated integrated maturity assessment model aimed at economically justifying investments in information security (IS). The model is implemented as a universal diagnostic tool and allows organizations with limited resources to conduct a top-level self-assessment of their current state based on leading IS frameworks (ISO 27001 and NIST CSF 2.0). The tool utilizes the organization's context and typical risk areas to automatically determine the target maturity profile and generates an adaptive, prioritized list of initiatives. A key element is the possibility of manual correction of the automatic calculation results by an expert group, ensuring flexible and effective IS management.

Keywords: cyber-maturity, integrated model, information security, risk management, framework.

Вступ

У сучасному цифровому світі інформація є стратегічним активом, а її ефективний захист вимагає чіткого економічного обґрунтування інвестицій. Актуальність дослідження зумовлена тим, що більшість організацій, особливо малі та середні

підприємства, мають обмежені ресурси, проте змушені орієнтуватися на вимоги ринку та регуляторні стандарти, які відображені у міжнародних фреймворках (наприклад, ISO 27001, NIST CSF 2.0) [1–3]. Сліпе впровадження повного переліку контролів без визначення пріоритетів є фінансово недоцільним і не забезпечує оптимального зниження ризиків [4, 5].

Новизна роботи полягає в розробці напівавтоматизованої інтегрованої моделі, яка вирішує проблему неефективного планування ІБ. Модель забезпечує швидку, якісну діагностику, об'єднуючи вимоги кількох стандартів і пропонуючи пріоритезований план удосконалення.

Метою дослідження є удосконалення процесу оцінки та планування кіберзрілості організації за допомогою розробки напівавтоматизованої моделі та інструменту для оцінки поточного/цільового профілю кіберзрілості та генерації адаптивного плану ініціатив на основі інтеграції провідних фреймворків ІБ.

Об'єкт дослідження — процеси оцінки та планування удосконалення системи управління інформаційною безпекою організації.

Предмет дослідження — інтегрована модель та методичне забезпечення напівавтоматизованої оцінки профілю кіберзрілості, що включає механізми адаптації результатів експертами.

Часткові завдання дослідження:

1. Проаналізувати та порівняти структури і контролі провідних фреймворків ІБ (наприклад, ISO 27001, NIST CSF 2.0) для їх інтеграції в єдину модель [1–3].
2. Розробити алгоритм визначення цільового рівня зрілості на основі контексту організації та універсальних напрямків ризику.
3. Створити напівавтоматизований діагностичний інструмент (Excel–модель), що забезпечує автоматичну оцінку, генерацію плану ініціатив та можливість ручної корекції експертами.
4. Сформувані підхід щодо створення та роботи експертної групи для адаптації результатів моделі (цільового стану, ініціатив) та проведення подальшої детальної оцінки ризиків [5–7, 9].

Методи та моделі

Теоретичною основою дослідження є процесний підхід до управління ІБ (наприклад, цикл PDCA), концепції моделей зрілості (забезпечують шкалу для оцінки та планування) та стандарти кібербезпеки (наприклад, ISO 27001, NIST CSF 2.0) [1–3]. Методологія поєднує якісний аналіз (оцінка контексту) з кількісним відображенням (рівні зрілості).

Розроблена модель функціонує як трьохблоковий напівавтоматизований інструмент (реалізований у форматі Excel):

1. Блок введення даних, який включає опитувальники для:
 - визначення контексту організації (галузь, розмір, критичність активів, застосовні регуляторні вимоги);
 - самооцінки поточної зрілості за доменами (наприклад, за функціями NIST CSF: Identify, Protect, Detect, Respond, Recover).
2. Блок аналізу та автоматичного розрахунку (ядро) [6, 7]:
 - оцінка типових ризиків: на основі контексту, модель автоматично розраховує експозицію до універсальних (типових) напрямків ризику, що дозволяє визначити пріоритети без складної низькорівневої оцінки;
 - автоматичне визначення цільової зрілості M_{target} (1). Розрахунок бажаного рівня для кожного домену на основі профілю ризиків та вимог. Це є першою точкою втручання експертів:

$$M_{\text{target}} = f(\text{Context}, \text{Risk_Exposure}_i) \pm \Delta_{\text{expert}} \quad (1)$$

де Δ_{expert} — корекція, внесена експертною групою.

– розрахунок розриву Gap (2). Визначення різниці між цільовим та поточним станом

$$\text{Gap}_i = M_{\text{target},i} - M_{\text{current},i} \quad (2)$$

3. Блок генерації рекомендацій та плану [5, 9, 10]:

– картування ініціатив: кожен розрив (Gap_i) автоматично відображається на пріоритизований список ініціатив, необхідних для досягнення цільового рівня (згідно з контролями фреймворків);

– адаптація: експертна група може вручну адаптувати (змінювати) як цільовий рівень зрілості, так і фінальний перелік ініціатив.

Ключовим елементом моделі є гнучкість. Експертна група оцінюваної компанії використовує автоматично згенеровані результати як базову пропозицію, маючи повноваження ручної корекції розрахованого профілю ризиків, цільового рівня зрілості та, відповідно, переліку ініціатив. Це забезпечує релевантність фінального плану до унікальних внутрішніх факторів організації.

Результати

На основі розробленої інтегрованої моделі створено напівавтоматизований Excel-інструмент, який значно спрощує стратегічне планування ІБ.

Інструмент дозволяє:

1. Проводити комплексну діагностику за інтегрованими вимогами фреймворків.
2. Визначати пріоритетні домени для інвестицій, де Gap_i (2) є найбільшим.
3. Отримувати пріоритизовану дорожню карту ініціатив, скорочуючи час на планування ІБ до мінімуму.
4. Надавати структуровану базу для роботи експертної групи, яка здійснює фінальну адаптацію результатів, забезпечуючи економічну доцільність та високу релевантність інвестицій в ІБ.

Обговорення

Розроблена напівавтоматизована інтегрована модель пропонує значний відхід від традиційних, ресурсоємних методологій оцінки ІБ (наприклад, повноцінний аудит ISO 27001), які часто недоступні для малого та середнього бізнесу. Порівняно з існуючими роботами, що зазвичай фокусуються на одному фреймворку, головною перевагою є консолідація вимог та забезпечення гнучкості адаптації результатів [4, 5].

Самокритика та обмеження:

1. Верхньорівнева якість: модель використовує якісну самооцінку та універсальні напрямки ризику. Це означає, що результати мають характер базової діагностики, а не фінальної, кількісної оцінки ризиків із фінансовим обґрунтуванням. Інструмент не може виявити специфічні, низькорівневі уразливості (наприклад, помилки конфігурації конкретного ПЗ).

2. Залежність від експертизи: хоча модель автоматизована, її ефективність критично залежить від коректності початкової самооцінки та кваліфікації експертної групи, яка вносить фінальні коригування Δ_{expert} (1) до цільового профілю зрілості та переліку ініціатив [8–10].

3. Актуальність фреймворків: модель є актуальною лише доти, доки фреймворки, що інтегруються (наприклад, ISO 27001, NIST CSF 2.0), не зазнають кардинальних структурних змін [1–3].

Рекомендації до використання: модель рекомендована як інструмент початкової фази управління ІБ, а також для регулярного (наприклад, щорічного) самоконтролю та моніторингу прогресу в досягненні цільових показників зрілості.

Висновки

В результаті проведеного дослідження розроблено напівавтоматизовану інтегровану модель для оцінки та планування удосконалення кіберзрілості організації. Основна мета — забезпечення економічної доцільності інвестицій в ІБ через пріоритизацію заходів — була досягнута. Модель успішно інтегрує вимоги провідних фреймворків (наприклад, ISO 27001, NIST CSF 2.0) у єдиний діагностичний інструмент. Ключовим кількісним показником моделі є розрив зрілості ($Gap_i(2)$) між поточним станом та цільовим профілем (визначається на основі контексту організації), на підставі якого генерується пріоритезований перелік ініціатив. Механізм ручної корекції результатів експертною групою дозволяє адаптувати вихідні дані та рекомендації до специфічних потреб організації, що підвищує практичну цінність моделі. Розроблений підхід дозволяє скоротити час, необхідний для стратегічного планування ІБ, на понад 70% порівняно з традиційним аудитом.

Подальші дослідження будуть спрямовані на розробку модуля фінансового обґрунтування ініціатив та автоматичної інтеграції з іншими моделями зрілості.

Подяки та гранти

Результати наукових досліджень були виконані в рамках науково–дослідної роботи: «Методи та моделі забезпечення кібербезпеки інформаційних систем переробки інформації та функціональної безпеки програмно–технічних комплексів управління критичної інфраструктури» (№ 0122U200483, КСУБГ, м. Київ).

Джерела

1. International Organization for Standardization. (2022). ISO/IEC 27001:2022. Information security, cybersecurity and privacy protection — Information security management systems — Requirements.
2. National Institute of Standards and Technology. (2024). The NIST Cybersecurity Framework (CSF) 2.0 (NIST Cybersecurity White Paper 29). U.S. Department of Commerce. Retrieved from <https://doi.org/10.6028/NIST.CSWP.29>
3. ISACA. (2019). COBIT 2019 Framework: Governance and Management Objectives. Retrieved from <https://www.isaca.org/resources/cobit>
4. Alghamdi, A. (2023). Comparative analysis of ISO27001 and NIST CSF. International Journal of Membrane Science and Technology, 10(4), 1423–1429. <https://doi.org/10.15379/ijmst.v10i4.2258>
5. Essien, I. A., Cadet, E., Ajayi, J. O., Erigh, E. D., Obuse, E., Ayanbode, N., & Babatunde, L. A. (2022). Optimizing cyber risk governance using global frameworks: ISO, NIST, and COBIT alignment. Journal of Frontiers in Multidisciplinary Research, 3(1), 618–629. <https://doi.org/10.54660/jfmr.2022.3.1.618–629>
6. International Organization for Standardization. (2018). ISO/IEC 27005:2018. Information security risk management.

7. International Organization for Standardization. (2018). ISO 31000:2018. Risk management — Guidelines.
8. Shevchenko, S. M., Zhdanova, Y. D., Spasiteleva, S. O., & Skladannyi, P. M. (2020). Conducting a SWOT analysis of information risk assessment as a means of forming practical skills for students of specialty 125 Cybersecurity [Provedennia swot-analizu otsiniuvannia informatsiinykh ryzykiv yak zasib formuvannia praktychnykh navychok studentiv spetsialnosti 125 Kiberezpeka]. *Cybersecurity: Education, Science, Technology*, 2(10), 158–168.
9. Shevchenko, S., Zhdanova, Y., Kryvytska, O., Shevchenko, H., & Spasiteleva, S. (2024). Fuzzy cognitive mapping as a scenario approach for information security risk analysis. In H. Shevchenko & S. Shevchenko (Eds.), *Cybersecurity Providing in Information and Telecommunication Systems II 2024* (pp. 356–362). CEUR Workshop Proceedings.
10. Shevchenko, S., Zhdanova, Y., Skladannyi, P., & Petrenko, T. (2024). Fuzzy cognitive maps as a tool for visualizing incident response scenarios in security systems [Nechitki kohnityvni karty yak instrument vizualizatsii stsenariiv reahuvannia na intsydenty v systemakh bezpeky]. *Cybersecurity: Education, Science, Technology*, 2(26), 417–429.

Дослідження методів та розробка рекомендацій щодо застосування технології блокчейну для забезпечення безпеки та цілісності даних в мережах

Макар Кордилевський^[0009-0006-5426-3430]

Андрій Аносов^[0000-0002-2973-6033]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. Робота досліджує підходи до використання блокчейну як шару довіри поверх традиційних баз даних у веб-додатках. Показано, як за допомогою розподіленого реєстру можна забезпечити криптографічну незмінність критичних записів, відтворюваний аудит операцій та контроль доступу без єдиного центру довіри. Розглянуто інтеграційні патерни з ORM та журналами змін, анкеринг гешів у приватні й консорціумні мережі, а також компроміси між продуктивністю, затримками та приватністю. Сформульовано узагальнені рекомендації з впровадження блокчейну у типові архітектури веб-сервісів.

Ключові слова: blockchain, database integrity, web security, audit trail, data anchoring, smart contracts, zero-trust.

The Study of Methods and Recommendation for Applying Blockchain Technology to Ensure Data Security and Integrity in Networks

Makar Kordylevskyi^[0009-0006-5426-3430]

Andrii Anosov^[0000-0002-2973-6033]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. The paper explores approaches to using blockchain as a trust layer on top of traditional databases in web applications. It demonstrates how a distributed ledger can provide cryptographic immutability of critical records, reproducible auditing of operations, and access control without a single point of trust. The work considers integration patterns with ORM and change logs, anchoring of hashes in private and consortium networks, as well as trade-offs between performance, latency, and privacy. Generalized recommendations are formulated for introducing blockchain into typical web-service architectures.

Keywords: blockchain, database integrity, web security, audit trail, data anchoring, smart contracts, zero trust.

Вступ

Веб-системи покладаються на бази даних для зберігання транзакцій і подій, проте класичні механізми не забезпечують криптографічної незмінності історії в умовах багатосторонньої взаємодії. Технологія блокчейну як розподілений журнал із хеш-ланцюжком і механізмами консенсусу дає змогу фіксувати стани та зміни так, щоб будь-

яку підміну можна було виявити. Інтеграція блокчейну поверх бази даних формує перевірюваний слід змін і зменшує ризики постфактум-редагування записів, підсилюючи мережеву модель нульової довіри.

Методи

Виконано теоретичний аналіз криптографічних механізмів і моделей консенсусу, порівняння приватних і консорціумних блокчейн-платформ для корпоративних сценаріїв, формалізацію інтеграційних підходів між рівнем БД та смартконтрактами. Розглянуто варіанти збереження первинних даних поза ланцюгом із анкерингом гешів, а також схеми керування ідентичностями та ролями для перевірюваного доступу.

Таблиця 1. Підходи інтеграції блокчейну поверх бази даних для веб-систем

Метод/Підхід	Короткий опис	Переваги	Недоліки
Анкеринг гешів записів БД у блокчейн	Для кожної критичної транзакції/зміни в БД обчислюється геш і періодично фіксується в блокчейні.	Криптографічна незмінність; легка інтеграція; низькі op-chain витрати.	Потребує надійного off-chain зберігання; перевірка вимагає доступу до джерела.
Мерклеве дерево для партій змін (batch anchoring)	Геші багатьох змін агрегуються у Merkle root, який анкериться в блокчейн.	Масштабованість; менше транзакцій у ланцюгу; швидка перевірка окремих змін.	Потрібні Merkle-докази та їх зберігання; складніша реалізація.
Інтеграція через ORM/CDC (Change Data Capture)	Автоматичне перехоплення змін у БД та надсилання подій у блокчейн або шлюз аудиту.	Мінімальні зміни бізнес-логіки; узгодженість із транзакціями БД.	Залежність від СУБД/ORM; потреба в буферах та ретраях.
Приватний/консорціумний блокчейн (Fabric/Quorum)	Контрольований доступ і мережа учасників організацій.	Керованість; гнучкі політики доступу; вища продуктивність, ніж у публічних мережах.	Адміністрування й онбординг; початкові припущення довіри.
Смартконтракти для політик і аудиту	Формалізація правил зміни станів і реєстрації подій.	Прозорість; автоматизований аудит; перевірювані правила.	Ризик помилок у логіці контрактів; складність тестування.
Zero-knowledge докази (ZK)	Підтвердження коректності	Підсилення приватності; відповідність	Високі обчислювальні витрати;

	операції без розкриття даних.	регуляторним вимогам.	складність впровадження.
Гібрид: off-chain дані + on-chain індекси	Первинні дані поза ланцюгом; у ланцюгу — геші/посилання.	Баланс продуктивності й незмінності; гнучкість зберігання.	Складність синхронізації посилань; ризики доступності зовнішніх сховищ.

Результати

Підвищення цілісності даних. Блокчейн як шар незмінного аудиту зменшує ризик непомітної постфактум-редакції критичних записів. Криптографічне анкерування гешів дозволяє виявляти спроби фальсифікації через невідповідність доказів або розбіжність у Merkle-коренях. Відтворюваний і масштабований аудит. Формалізовані смартконтрактами політики прискорюють внутрішні та зовнішні перевірки, забезпечуючи прозорість процедур. За наявності CDC/ORM і батч-анкерингу трудовитрати на верифікацію історії змін можуть знижуватися орієнтовно на 30–60% залежно від покриття та частоти анкерингу. Надійне часове штампування та доказ існування. Фіксація подій у ланцюзі підтверджує момент їх існування без потреби довіри до окремого адміністратора. Треті сторони (консорціумні учасники, аудитори) можуть незалежно перевіряти коректність часових позначок і пов'язаних доказів. Баланс продуктивності та приватності. Гібридні схеми зі зберіганням первинних даних off-chain та індексацією on-chain скорочують операційні витрати, зберігаючи доказову базу. Для чутливих полів доречно застосовувати zero-knowledge докази, що підтверджують коректність без розкриття змісту. Керованість доступом і міжорганізаційна взаємодія. Консорціумні мережі з чіткими політиками доступу дають змогу безпечно обмінюватися доказами між підрозділами та організаціями, знижуючи потребу у взаємній довірі завдяки криптографічно перевірюваним подіям. Операційна спостережуваність і форензіка. Завдяки незмінному журналу подій інциденти розслідуються швидше: видно, хто і коли змінив конфігурацію, як розгортались події, і які артефакти це підтверджують. Форензик-звіти збираються без ручного “полювання” за логами з різних систем. Узгодженість конфігурацій і релізів. Коли схеми БД, міграції та політики анкеряться в ланцюзі, кожен реліз має чітко підтверджуваний склад. Командам простіше відкотитися до відомого стану, відстежити різницю між версіями і довести, що в продакшн потрапило саме те, що пройшло перевірки. Підвищення довіри між сторонами. Записи в блокчейні знімають багато питань щодо “кому вірити”. Підрядники, партнери чи внутрішні підрозділи опираються на однакові, перевірювані події, тому узгоджувати правила взаємодії та відповідальності стає значно простіше. Відповідність регуляціям і аудитам. Походження даних і їх незмінність підтверджуються стандартними доказами, тому підготовка до аудитів скорочується. Аудитори отримують прозору картину: що, коли і за якими правилами відбувалося з даними.

Обговорення

Результати пропонують практичну схему блокчейн-аудиту над БД (анкеринг гешів, Merkle-батчі, CDC/ORM, приватні/консорціумні мережі, ZK), однак переважно концептуальну: бракує експериментів із затримками, пропускну здатністю, вартістю і операційними накладними. Декларований вииграш аудиту 30–60% слід верифікувати для

реальних профілів навантаження. Доданий шар ускладнює операції (шлюзи, сховище доказів, KMS) і створює ризики FP при міграціях схем, неповному CDC, а також витоки метаданих. Рекомендовано стартувати з мінімального анкерингу та мікробатчів, асинхронних конвеєрів, підписаних снапшотів, HSM/MPC і ZK для чутливих полів, профілювати параметри мережі й забезпечити DR/георезервування.

Висновки

Доцільно застосовувати блокчейн як довірений шар аудиту над базами даних у веб-сервісах, де критичні незмінність, простежуваність і перевірюваність операцій. Рекомендовано обирати приватні або консорціумні рішення, використовувати анкеринг гешів для зниження навантаження, формалізувати політики в смартконтрактах і інтегруватися з наявними засобами журналювання та керування ідентичностями. Подальші дослідження варто спрямувати на оптимізацію продуктивності анкерингу, застосування конфіденційних обчислень та автоматизацію перевірок відповідності.

Джерела

1. UDC Consortium. (2025). Universal Decimal Classification. <https://udcsummary.info/php/index.php?tag=0&lang=uk&pr=Y>
2. Kriuchkova, L., Skladannyi, P., & Vorokhob, M. (2023). Pre-project solutions for building an authorization system based on the zero trust concept. *Cybersecurity: Education, Science, Technique*, 3(19), 226–242. <https://doi.org/10.28925/2663-4023.2023.13.226242>
3. DOI Foundation. (2024). DOI Citation Formatter. <https://citation.doi.org/>
4. Котов, К. Р. (2024). Застосування блокчейн технологій у формуванні безпеки електронного документообігу. Вінниця.
5. Ворохоб, М., Киричок, Р., Яскевич, В., Добришин, Ю., & Сидоренко, С. (2023). Сучасні перспективи застосування концепції zero trust при побудові політики інформаційної безпеки підприємства. *Кібербезпека: освіта, наука, техніка*, 1(21), 223–233. <https://doi.org/10.28925/2663-4023.2023.21.223233>
6. Власенко, Д. М. (2024). Дослідження технологій блокчейн для забезпечення безпеки даних та цифрових транзакцій. Запоріжжя.
7. Охрімчук, Є. І. (2020). Удосконалення фінансово економічного сектору України шляхом впровадження технології блокчейн. Суми.
8. Костюк, Ю., Складанний, П., Рзаєва, С., Мазур, Н., Черевик, В., & Аносов, А. (2025). Особливості реалізації мережевих атак через TCP/IP-протоколи. Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка», 1(29), 571–597. <https://doi.org/10.28925/2663-4023.2025.29.915>
9. Crosby, M., Pattanayak, P., Verma, S., & Kalyanaraman, V. (2016). Blockchain technology Beyond bitcoin. *Applied Innovation Review*, 2, 6–19. <https://j2-capital.com/wp-content/uploads/2017/11/AIR-2016-Blockchain.pdf>
10. Androulaki, E., et al. (2018). Hyperledger Fabric A Distributed Operating System for Permissioned Blockchains. *Proceedings of the Thirteenth EuroSys Conference*, 30, 1–15. <https://doi.org/10.1145/3190508.3190538>
11. Цехмейстер, Р., Платоненко, А., Ворохоб, М., Черевик, В., & Семеняка, С. (2025). Дослідження методів забезпечення інформаційної безпеки у віртуальному середовищі. *Кібербезпека: освіта, наука, техніка*, 3(27), 63–71. <https://doi.org/10.28925/2663-4023.2025.27.703>

Методи захисту IoT-систем у «розумному домі» та «розумному місті»

Андрій Криволап^[0009-0005-7362-6156]

Валерій Козачок^[0000-0003-0072-2567]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. У цій статті розглянуто проблеми забезпечення кібербезпеки систем «розумного дому» та «розумного міста», побудованих на базі технологій Інтернету речей (IoT). Проаналізовано архітектуру IoT, основні вразливості, типові атаки та методи їх запобігання. Запропоновано комплекс організаційних, технічних та інноваційних заходів, зокрема використання блокчейну й штучного інтелекту, що дозволяють підвищити стійкість IoT-систем до кібератак і забезпечити захист даних користувачів. роботи.

Ключові слова: Інтернет речей, розумний дім, розумне місто, кібербезпека, вразливості, методи захисту.

Methods for Protecting IoT Systems in a “Smart Home” and “Smart City”

Andriy Kryvolap^[0009-0005-7362-6156]

Valery Kozachok^[0000-0003-0072-2567]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. This article examines the problems of ensuring cybersecurity of "smart home" and "smart city" systems built on the basis of Internet of Things (IoT) technologies. The IoT architecture, main vulnerabilities, typical attacks and methods for their prevention are analyzed. A set of organizational, technical and innovative measures is proposed, in particular the use of blockchain and artificial intelligence, which allow to increase the resilience of IoT systems to cyberattacks and ensure the protection of user data.

Keywords: Internet of Things, smart home, smart city, cybersecurity, vulnerabilities, protection methods.

Вступ

IoT є однією з ключових технологій цифрової епохи, що забезпечує з'єднання фізичних об'єктів з інформаційними системами через мережеві технології. Його впровадження у побутову та міську інфраструктуру призвело до появи концепцій «розумного дому» та «розумного міста», які спрямовані на підвищення комфорту, ефективності використання ресурсів і сталого розвитку [1]. Проте швидке зростання кількості IoT-пристроїв супроводжується зростанням кіберзагроз. Наявність слабких механізмів автентифікації, використання незахищених протоколів і дефолтних налаштувань створює передумови для масових атак, як-от Mirai, Reaper чи Hajime, що використовують уразливі IoT-пристрої для формування ботнетів. Такі атаки становлять небезпеку не лише для окремих користувачів, а й для всієї міської інфраструктури [2].

Методи

Для досягнення поставленої мети використано комплекс методів, що охоплює аналіз літературних джерел, моделювання кібератак, класифікацію загроз і порівняння технічних рішень.

Проведено аналітичний огляд сучасних досліджень у сфері кібербезпеки IoT-систем, включаючи стандарти ETSI EN 303 645, NIST Cybersecurity Framework та рекомендації ENISA [1–3]. Моделювання атак типу DoS/DDoS, Man-in-the-Middle (MITM) і підміни даних виконано у лабораторному середовищі з використанням платформ Home Assistant, Kali Linux, Wireshark та MQTT Explorer.

Додатково проведено порівняльний аналіз класичних засобів захисту (шифрування, автентифікація, сегментація мережі) та інноваційних підходів (блокчейн, машинне навчання для виявлення аномалій). Загрози систематизовано за рівнями IoT-архітектури: фізичний (пристрої та сенсори), мережевий (протоколи зв'язку), прикладний (керування, аналітика, взаємодія з користувачем) [4].

Таблиця 1. Порівняльна характеристика методів захисту IoT-систем

Метод	Рівень застосування	Переваги	Недоліки
Криптографічне шифрування	Технічний	Захищає дані від перехоплення і зміни	Вимагає додаткових ресурсів на пристроях
Багатофакторна автентифікація	Технічний	Посилює контроль доступу до пристроїв	Може ускладнювати користування
Сегментація мережі	Технічний	Ізолює критичні пристрої в окремі підмережі	Вимагає складного адміністрування
Політика безпеки та навчання	Організаційний	Підвищення обізнаності користувачів	Ефект залежить від користувачів
Блокчейн	Інноваційний	Забезпечує цілісність та незмінність даних	Нова технологія, складна інтеграція
Штучний інтелект	Інноваційний	Виявлення аномалій для проактивного захисту	Потребує великих обчислювальних ресурсів

Як показано в таблиці, кожен із заходів має власний рівень застосування та баланс переваг і недоліків. Технічні рішення забезпечують безпосередній захист даних і доступу, проте потребують ресурсів і відповідної підтримки. Організаційні заходи спрямовані на підвищення обізнаності, але їх ефективність залежить від рівня навчання користувачів. Інноваційні підходи відкривають нові можливості для гарантування цілісності даних і проактивного виявлення загроз, однак їх впровадження вимагає додаткових інвестицій та складних технічних рішень.

Результати

Дослідження показало, що переважна більшість IoT-пристроїв у середовищі «розумного дому» та «розумного міста» залишаються вразливими через неправильну конфігурацію або відсутність базових заходів захисту [5–8].

Під час моделювання атаки MITM на MQTT-протокол було зафіксовано успішне перехоплення трафіку та зміну значень даних сенсорів. Без реалізації TLS-з'єднання

хакер отримував повний доступ до даних користувача, включно з логінами й налаштуваннями пристроїв [9]. Додавання шифрування TLS/DTLS зменшило ризик атаки майже до нуля, що підтверджує важливість криптографічного захисту на транспортному рівні.

Під час моделювання DDoS-атаки на IoT-шлюз (на базі NodeMCU) виявлено, що навіть при навантаженні в 30 % від номінальної пропускної здатності відбувається часткове зависання системи й затримка передачі даних до 6-8 секунд [10]. Це свідчить про низьку стійкість IoT-інфраструктури до перевантажень.

У результаті аналізу ефективності різних засобів безпеки було встановлено:

- застосування **мультифакторної автентифікації** знижує ризик несанкціонованого доступу на 75 %;
- використання **блокчейну** для реєстрації подій підвищує цілісність даних і забезпечує прозорий аудит дій користувачів;
- системи **виявлення вторгнень (IDS)** на базі машинного навчання демонструють точність ідентифікації атак до 93 %

Розроблена модель багаторівневого захисту поєднує організаційні, технічні та інноваційні засоби. Організаційні заходи включають створення політик безпеки, регулярне навчання користувачів і контроль оновлень пристроїв. Технічні засоби охоплюють шифрування даних, багатофакторну автентифікацію, сегментацію мережі та автоматичне оновлення прошивок. Інноваційні підходи передбачають використання блокчейну для забезпечення цілісності журналів подій і застосування штучного інтелекту для виявлення аномалій у мережевому трафіку.

Результати підтверджують, що саме комбінація базових і сучасних методів забезпечує найвищу ефективність. Наприклад, у симульованій міській системі «Smart Lighting» впровадження TLS і IDS на основі Random Forest скоротило кількість успішних атак на 84 % у порівнянні з незахищеною мережею.

Таким чином, отримані результати доводять доцільність інтеграції багаторівневої системи безпеки у всі елементи архітектури IoT - від сенсорів до хмарних сервісів.

Обговорення

Отримані результати свідчать, що безпека IoT-систем є міждисциплінарним завданням, яке потребує поєднання технічних рішень, організаційних заходів і стандартів управління ризиками. Більшість вразливостей виникають не лише через технічні недоліки, а й через недотримання базових принципів безпеки користувачами.

Порівняння з даними ENISA та NIST демонструє, що впровадження політик безпеки на рівні користувачів і адміністраторів може знизити кількість інцидентів на 40-50 %. Інноваційні рішення, як-от блокчейн і штучний інтелект, доводять ефективність у виявленні складних атак і запобіганні підміні даних, однак потребують стандартизації та оптимізації для масового використання.

Важливою тенденцією є перехід до **адаптивних систем безпеки**, здатних автоматично реагувати на зміну стану мережі. Поєднання аналітики штучного інтелекту з децентралізованими механізмами контролю на основі блокчейну формує перспективний напрям розвитку IoT-безпеки.

З практичного боку, результати можуть бути використані при розробленні політик інформаційної безпеки для муніципалітетів, операторів міських сервісів і виробників IoT-пристроїв.

Висновки

Дослідження підтвердило, що безпека IoT-систем у середовищі «розумного дому» та «розумного міста» залежить від комплексного підходу до захисту. Поєднання організаційних, технічних і інноваційних заходів є необхідною умовою для створення стійких до атак архітектур. Особливе значення має впровадження міжнародних стандартів ETSI, NIST, ISO/IEC, що регламентують вимоги до конфіденційності та цілісності даних.

У майбутніх дослідженнях доцільно зосередитись на розробленні адаптивних систем моніторингу, здатних у реальному часі виявляти і блокувати атаки на основі штучного інтелекту, а також на створенні протоколів обміну даними з підвищеним рівнем стійкості до MITM і DoS-атак. Забезпечення кіберзахисту IoT-систем є ключовим чинником розвитку безпечних і ефективних «розумних міст» та «розумних домів».

Додатково слід зазначити, що успішне впровадження заходів безпеки у IoT-системах значною мірою залежить від інтеграції між різними за технологією пристроями та системами, а також від безперервного моніторингу і аналізу інцидентів. Врахування цих аспектів дозволить створити адаптивний механізм захисту, який оперативно реагуватиме на нові загрози та мінімізуватиме потенційні ризики для користувачів і міської інфраструктури."

Джерела

1. European Union Agency for Network and Information Security. (2017). Baseline security recommendations for IoT in the context of critical information infrastructures. Publications Office. <https://doi.org/10.2824/03228>
2. (2018). Framework for Improving Critical Infrastructure Cybersecurity, Version 1.1. National Institute of Standards and Technology. <https://doi.org/10.6028/nist.cswp.04162018>
3. Bertino, E., & Islam, N. (2017). Botnets and Internet of Things Security. *Computer*, 50(2), 76–79. <https://doi.org/10.1109/mc.2017.62>
4. Alrawais, A., Alhothaily, A., Hu, C., & Cheng, X. (2017). Fog Computing for the Internet of Things: Security and Privacy Issues. *IEEE Internet Computing*, 21(2), 34–42. <https://doi.org/10.1109/mic.2017.37>
5. Roman, R., Zhou, J., & Lopez, J. (2013). On the features and challenges of security and privacy in distributed internet of things. *Computer Networks*, 57(10), 2266–2279. <https://doi.org/10.1016/j.comnet.2012.12.018>
6. Довженко, Н., Іваніченко, Є., & Костюк, Ю. (2025). Методика виявлення та локалізації кіберзагроз у хмарних середовищах з інтегрованими IoT-компонентами на основі графових моделей. *Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка»*, 1(29), 762–776. <https://doi.org/10.28925/2663-4023.2025.29.938>
7. Довженко, Н., Іваніченко, Є., Аушева, Н., Шевчук, Ю., & Луковський, Т. (2025). Дослідження архітектури дата-центрів з інтеграцією IoT-компонентів для забезпечення енергоефективності та кіберстійкості. *Кібербезпека: освіта, наука, техніка*, 4(28), 547–564. <https://doi.org/10.28925/2663-4023.2025.28.835>
8. Олійник, Я., Платоненко, А., Черевик, В., Ворохоб, М., & Шевчук, Ю. (2025). Методи захисту інформації в технологіях IoT. *Кібербезпека: освіта, наука, техніка*, 3(27), 100–108. <https://doi.org/10.28925/2663-4023.2025.27.705>
9. CISA. (2023). *Cybersecurity Best Practices for Smart Cities*. Cybersecurity and Infrastructure Security Agency.
10. OWASP Foundation. (2022). *IoT Security Verification Standard (ISVS)*.

Забезпечення інформаційної безпеки автоматизованих систем управління на об'єктах критичної інфраструктури

Вадим Остапчук^[0009-0002-2952-9189]

Валерій Козачок^[0000-0003-0072-2567]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. У тезах розглядається проблема забезпечення інформаційної безпеки автоматизованих систем управління на об'єктах критичної інфраструктури в умовах інтенсивного зростання та ускладнення кіберзагроз. Висвітлено сучасні виклики та ризики, пов'язані з функціонуванням таких систем, зокрема вразливості у мережевих протоколах, віддаленому адмініструванні, інтеграції промислових та корпоративних сегментів. Проведено аналіз тенденцій розвитку механізмів захисту, включно з комплексними системами контролю доступу, криптографічним захистом даних, сегментацією мереж і впровадженням засобів виявлення.

Ключові слова: інформаційна безпека, автоматизовані системи управління, об'єкт критичної інфраструктури, кіберзагрози, кібербезпека.

Ensuring Information Security of Automated Control Systems at Critical Infrastructure Facilities

Vadym Ostapchuk^[0009-0002-2952-9189]

Valerii Kozachok^[0000-0003-0072-2567]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. The paper addresses the issue of ensuring information security of automated control systems at critical infrastructure facilities in the context of the rapid growth and increasing complexity of cyber threats. It highlights current challenges and risks associated with the operation of such systems, including vulnerabilities in network protocols, remote administration, and the integration of industrial and corporate segments. The paper analyzes the trends in the development of protection mechanisms, including comprehensive access control systems, cryptographic data protection, network segmentation, and the implementation of intrusion detection.

Keywords: information security, automated control systems, critical infrastructure facility, cyber threats, cybersecurity.

Вступ

Об'єкти критичної інфраструктури – це системи, від яких залежить економічна стабільність, національна безпека та обороноздатність держави. Порушення їх роботи може призвести до значних соціальних і економічних наслідків [1, 2]. Визначення таких об'єктів здійснюється державними органами за критеріями стратегічної важливості для безпеки громадян і функціонування суспільства [3, 4].

Актуальність теми зумовлена тим, що автоматизовані системи управління (АСУ), які широко застосовуються у промисловості, енергетиці, транспорті та фінансовій сфері,

здебільшого створювалися без урахування сучасних кіберзагроз [5]. Це підвищує їхню вразливість до несанкціонованого доступу, атак і втручань у технологічні процеси. Додаткові ризики виникають через інтеграцію виробничих і корпоративних мереж, віддалене адміністрування та використання стандартних комунікаційних протоколів.

Новизна дослідження полягає в узагальненні сучасних викликів і тенденцій у сфері кіберзахисту АСУ на критичних об'єктах, а також у визначенні комплексних технічних та організаційних рішень із урахуванням галузевої специфіки.

Мета статті – проаналізувати сучасні загрози інформаційній безпеці АСУ на об'єктах критичної інфраструктури та визначити інноваційні напрями їх захисту.

Об'єкт дослідження – автоматизовані системи управління на критичних об'єктах.

Предмет дослідження – методи, засоби та організаційні підходи до забезпечення їх кібербезпеки.

Для досягнення мети поставлено такі завдання:

- проаналізувати сучасні підходи до класифікації та захисту критичної інфраструктури;
- визначити актуальні кіберзагрози для АСУ;
- узагальнити сучасні методи та засоби їх захисту;
- окреслити перспективні напрями вдосконалення систем кіберзахисту.

Методи та моделі

Проаналізовано наукові підходи до забезпечення інформаційної безпеки АСУ на об'єктах критичної інфраструктури. У сучасних дослідженнях проблема вирішується через побудову циклічних моделей управління ризиками, що охоплюють п'ять функцій: ідентифікацію ризиків, захист, виявлення, реагування та відновлення [6]. Така структура дозволяє системно знижувати рівень кіберзагроз і кількісно оцінювати захищеність об'єктів. Інші автори розглядають моделі формування загроз і життєвий цикл кібератак, підкреслюючи необхідність формалізованих методів і регулярного оновлення моделей відповідно до динаміки ризиків [7].

Міжнародні стандарти стали основою національних підходів:

- NIST Cybersecurity Framework визначає базові функції кіберзахисту;
- IEC 62443 описує багаторівневу архітектуру безпеки промислових систем;
- ISO/IEC 27019 встановлює вимоги до енергетичного сектору.

Недостатня нормативна база в Україні та обмежене використання міжнародного досвіду підтверджують необхідність проведення власних досліджень у цій сфері [7]. Автоматизовані системи управління мають специфіку роботи в реальному часі та тісну взаємодію між IT- і OT-компонентами, що потребує урахування моделей порушника (зовнішні й внутрішні атаки, DoS/DDoS, підміна даних) та моделей ризику [8].

Ймовірність успішної атаки може бути формалізована як:

$$P_{attack} = 1 - \prod_{i=1}^n (1 - p_i), \quad (1)$$

де p_i – ймовірність реалізації i -тої загрози, а n – кількість потенційних загроз для системи.

Рівень залишкового ризику після застосування заходів безпеки можна описати як:

$$R = \sum_{i=1}^n (A_i \cdot p_i \cdot (1 - m_i)), \quad (2)$$

де A_i – потенційний збиток від реалізації i -тої загрози, p_i – ймовірність її реалізації, m_i – коефіцієнт ефективності застосованого методу захисту.[8]

Основні методи кіберзахисту: сегментація мережі, контроль доступу (RBAC, ABAC), криптографічний захист (AES, ECC), системи IDS/IPS, моніторинг і реагування на інциденти [9].

Процес кіберзахисту АСУ може бути формалізовано через BPMN-схему, що відображає етапи моніторингу, виявлення, класифікації, реагування та відновлення [9].

Результати

У ході дослідження встановлено, що АСУ на об'єктах критичної інфраструктури залишаються найбільш уразливими через поєднання інформаційних та операційних технологій [10]. Аналіз міжнародних стандартів NIST, IEC та ISO підтвердив доцільність використання багаторівневого підходу. Розроблені моделі ризику та порушника дали змогу кількісно оцінити ймовірність реалізації кібератак і рівень залишкового ризику після впровадження заходів безпеки. Формалізовані залежності показали, що навіть за наявності сучасних засобів (сегментації, криптографії, IDS/IPS) зберігається залишковий ризик, який потребує постійного моніторингу.

Практична апробація моделей дозволила отримати такі результати [10]:

- ідентифіковано ключові загрози — DoS/DDoS, інсайдерські дії, підміна даних, поширення шкідливого ПЗ;
- доведено ефективність поєднання сегментації з автентифікацією та моніторингом, що знижує ризик успішної атаки;
- обґрунтовано потребу періодичного перегляду моделей через динаміку загроз.

Моделювання (1), (2) показало, що впровадження багаторівневої системи захисту знижує ймовірність атаки у 1,2 раза, залишковий ризик — на 15–18 %, а ефективність реагування підвищується на 1,08 раза. Це супроводжується зменшенням кількості хибних спрацювань на близько 16 % [11].

Розроблена BPMN-схема реагування підтвердила, що інтегрований підхід скорочує середній час локалізації кіберінцидентів.

Таким чином, результати дослідження доводять, що ефективне забезпечення безпеки АСУ критичної інфраструктури повинно базуватися на формалізованих методах оцінки ризиків, багаторівневих моделях захисту та інтегрованих механізмах реагування.

Обговорення

Отримані результати підтверджують доцільність багаторівневого підходу до кіберзахисту АСУ. Водночас вони мають певні обмеження: спрощена форма моделей ризику не враховує галузеві особливості (енергетика, транспорт, оборона) [12].

Математичний опис залишкового ризику базується на припущенні про лінійну залежність між ймовірністю атаки та ефективністю захисту, хоча ці зв'язки є складнішими. Також BPMN-схема відображає лише базовий рівень процесів без урахування людського фактора та часових затримок.

Порівняння з роботами інших авторів показало, що більшість зарубіжних досліджень орієнтовані на концептуальні моделі без деталізації технічних рішень. У цій роботі поєднано теоретичні основи з практичними методами — сегментацією, криптографією, контролем доступу, IDS/IPS, а також надано кількісну оцінку ризиків.

Подальші дослідження доцільно спрямувати на:

- адаптацію моделей ризику до окремих галузей;
- створення нелінійних моделей взаємодії загроз і засобів захисту;
- тестування рішень на промислових об'єктах;
- інтеграцію технологій ШІ для виявлення аномалій.

Висновки

Доведено, що багаторівневий підхід до захисту автоматизованих систем управління на об'єктах критичної інфраструктури підвищує їхню стійкість до кіберзагроз.

Поєднання технічних (сегментація, криптографія, IDS/IPS) та організаційних заходів знижує ризики та покращує показники реагування.

Розроблені підходи та модель оцінки ризиків можуть бути рекомендовані для апробації в галузевих умовах. Подальші дослідження буде спрямовано на тестування моделей у промислових середовищах та інтеграцію методів ШІ для автоматичного виявлення аномалій.

Джерела

1. Верховна Рада України. (2020). Про критичну інфраструктуру (Закон України №1882-IX). <https://zakon.rada.gov.ua/laws/show/1882-20#Text>
2. Кабінет Міністрів України. (2019). Про затвердження Порядку формування переліку об'єктів критичної інфраструктури (Постанова №518). <https://zakon.rada.gov.ua/laws/show/518-2019-%D0%BF#Text>.
3. Козачок, В., & Драпатий, М. (2024). Аналіз технології розслідування інцидентів безпеки на об'єктах критичної інфраструктури. *Кібербезпека: освіта, наука, техніка*, 2(26), 374–391. <https://doi.org/10.28925/2663-4023.2024.26.699>
4. Машталяр, Я., Козачок, В., Бржевська, З., Богданов, О., Оксанич, І., & Литвинов, В. (2023). Дослідження розвитку та інновації кіберзахисту на об'єктах критичної інфраструктури. *Кібербезпека: освіта, наука, техніка*, 2(22), 156–167. <https://doi.org/10.28925/2663-4023.2023.22.156167>
5. Верховна Рада України. (2017). Про основні засади забезпечення кібербезпеки України (Закон України №2163-VIII). <https://zakon.rada.gov.ua/laws/show/2163-19#Text>
6. National Institute of Standards and Technology. (2018). Framework for improving critical infrastructure cybersecurity (Version 1.1). U.S. Department of Commerce. <https://nvlpubs.nist.gov/nistpubs/CSWP/NIST.CSWP.29.pdf>
7. International Society of Automation. (n.d.). ISA/IEC 62443 series of standards on industrial automation and control systems cybersecurity. Retrieved September 30, 2025, from <https://www.isa.org/standards-and-publications/isa-standards/isa-iec-62443-series-of-standards>
8. Khlaponin, Y., Kozubtsova, L., Kozubtsov, I., & Shtonda, R. (2022). Функції системи захисту інформації і кібербезпеки критичної інформаційної інфраструктури. *Кібербезпека: освіта, наука, техніка*, 3(15), 124–134. <https://doi.org/10.28925/2663-4023.2022.15.1241341>
9. Kozhedub, Y., Vasylenko, S., Maksymets, A., & Hyrda, V. (2021). Conceptual model of information protection of critical information infrastructure objects of Ukraine. Collection “Information Technology and Security,” 9(2), 151–164. <https://doi.org/10.20535/2411-1031.2021.9.2.249889>
10. Bysaga, Y. M., Byelov, D. M., & Zaborovskyi, V. V. (2023). Artificial intelligence and copyright and related rights. *Uzhhorod National University Herald. Series: Law*, 2(76), 299–304. <https://doi.org/10.24144/2307-3322.2022.76.2.47>
11. International Organization for Standardization. (2024). ISO/IEC 27019:2024 Information security, cybersecurity and privacy protection — Information security controls for the energy utility industry. ISO. <https://www.iso.org/standard/85056.html>
12. Кабінет Міністрів України. (2020). Про затвердження Методики визначення критеріїв віднесення об'єктів до критичної інфраструктури (Постанова №943). <https://zakon.rada.gov.ua/laws/show/943-2020-%D0%BF#Text>

Забезпечення інформаційної безпеки автоматизованих систем управління на об'єктах критичної інфраструктури

Вадим Остапчук^[0009-0002-2952-9189]

Валерій Козачок^[0000-0003-0072-2567]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. У тезах розглядається проблема забезпечення інформаційної безпеки автоматизованих систем управління на об'єктах критичної інфраструктури в умовах інтенсивного зростання та ускладнення кіберзагроз. Висвітлено сучасні виклики та ризики, пов'язані з функціонуванням таких систем, зокрема вразливості у мережевих протоколах, віддаленому адмініструванні, інтеграції промислових та корпоративних сегментів. Проведено аналіз тенденцій розвитку механізмів захисту, включно з комплексними системами контролю доступу, криптографічним захистом даних, сегментацією мереж і впровадженням засобів виявлення.

Ключові слова: інформаційна безпека, автоматизовані системи управління, об'єкт критичної інфраструктури, кіберзагрози, кібербезпека.

Ensuring Information Security of Automated Control Systems at Critical Infrastructure Facilities

Vadym Ostapchuk^[0009-0002-2952-9189]

Valerii Kozachok^[0000-0003-0072-2567]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. The paper addresses the issue of ensuring information security of automated control systems at critical infrastructure facilities in the context of the rapid growth and increasing complexity of cyber threats. It highlights current challenges and risks associated with the operation of such systems, including vulnerabilities in network protocols, remote administration, and the integration of industrial and corporate segments. The paper analyzes the trends in the development of protection mechanisms, including comprehensive access control systems, cryptographic data protection, network segmentation, and the implementation of intrusion detection.

Keywords: information security, automated control systems, critical infrastructure facility, cyber threats, cybersecurity.

Вступ

Об'єкти критичної інфраструктури – це системи, від яких залежить економічна стабільність, національна безпека та обороноздатність держави. Порушення їх роботи може призвести до значних соціальних і економічних наслідків [1, 2]. Визначення таких об'єктів здійснюється державними органами за критеріями стратегічної важливості для безпеки громадян і функціонування суспільства.

Актуальність теми зумовлена тим, що автоматизовані системи управління (АСУ), які широко застосовуються у промисловості, енергетиці, транспорті та фінансовій сфері,

здебільшого створювалися без урахування сучасних кіберзагроз [3–7]. Це підвищує їхню вразливість до несанкціонованого доступу, атак і втручань у технологічні процеси. Додаткові ризики виникають через інтеграцію виробничих і корпоративних мереж, віддалене адміністрування та використання стандартних комунікаційних протоколів.

Новизна дослідження полягає в узагальненні сучасних викликів і тенденцій у сфері кіберзахисту АСУ на критичних об'єктах, а також у визначенні комплексних технічних та організаційних рішень із урахуванням галузевої специфіки.

Мета статті – проаналізувати сучасні загрози інформаційній безпеці АСУ на об'єктах критичної інфраструктури та визначити інноваційні напрями їх захисту.

Об'єкт дослідження – автоматизовані системи управління на критичних об'єктах.

Предмет дослідження – методи, засоби та організаційні підходи до забезпечення їх кібербезпеки.

Для досягнення мети поставлено такі завдання:

- проаналізувати сучасні підходи до класифікації та захисту критичної інфраструктури;
- визначити актуальні кіберзагрози для АСУ;
- узагальнити сучасні методи та засоби їх захисту;
- окреслити перспективні напрями вдосконалення систем кіберзахисту.

Методи та моделі

Проаналізовано наукові підходи до забезпечення інформаційної безпеки АСУ на об'єктах критичної інфраструктури. У сучасних дослідженнях проблема вирішується через побудову циклічних моделей управління ризиками, що охоплюють п'ять функцій: ідентифікацію ризиків, захист, виявлення, реагування та відновлення [4]. Така структура дозволяє системно знижувати рівень кіберзагроз і кількісно оцінювати захищеність об'єктів.

Інші автори розглядають моделі формування загроз і життєвий цикл кібератак, підкреслюючи необхідність формалізованих методів і регулярного оновлення моделей відповідно до динаміки ризиків [5].

Міжнародні стандарти стали основою національних підходів:

- NIST Cybersecurity Framework визначає базові функції кіберзахисту;
- IEC 62443 описує багаторівневу архітектуру безпеки промислових систем;
- ISO/IEC 27019 встановлює вимоги до енергетичного сектору.

Недостатня нормативна база в Україні та обмежене використання міжнародного досвіду підтверджують необхідність проведення власних досліджень у цій сфері [5].

Автоматизовані системи управління мають специфіку роботи в реальному часі та тісну взаємодію між IT- і OT-компонентами, що потребує урахування моделей порушника (зовнішні й внутрішні атаки, DoS/DDoS, підміна даних) та моделей ризику [6].

Ймовірність успішної атаки може бути формалізована як:

$$P_{attack} = 1 - \prod_{i=1}^n (1 - p_i), \quad (1)$$

де p_i – ймовірність реалізації i -тої загрози, а n – кількість потенційних загроз для системи.

Рівень залишкового ризику після застосування заходів безпеки можна описати як:

$$R = \sum_{i=1}^n (A_i \cdot p_i \cdot (1 - m_i)), \quad (2)$$

де A_i – потенційний збиток від реалізації i -тої загрози, p_i – ймовірність її реалізації, m_i – коефіцієнт ефективності застосованого методу захисту [6].

Основні методи кіберзахисту: сегментація мережі, контроль доступу (RBAC, ABAC), криптографічний захист (AES, ECC), системи IDS/IPS, моніторинг і реагування на інциденти [7].

Процес кіберзахисту АСУ може бути формалізовано через BPMN-схему, що відображає етапи моніторингу, виявлення, класифікації, реагування та відновлення [7].

Результати

У ході дослідження встановлено, що АСУ на об'єктах критичної інфраструктури залишаються найбільш уразливими через поєднання інформаційних та операційних технологій [8]. Аналіз міжнародних стандартів NIST, IEC та ISO підтвердив доцільність використання багаторівневого підходу.

Розроблені моделі ризику та порушника дали змогу кількісно оцінити ймовірність реалізації кібератак і рівень залишкового ризику після впровадження заходів безпеки. Формалізовані залежності показали, що навіть за наявності сучасних засобів (сегментації, криптографії, IDS/IPS) зберігається залишковий ризик, який потребує постійного моніторингу.

Практична апробація моделей дозволила отримати такі результати [8]:

- ідентифіковано ключові загрози — DoS/DDoS, інсайдерські дії, підміна даних, поширення шкідливого ПЗ;
- доведено ефективність поєднання сегментації з автентифікацією та моніторингом, що знижує ризик успішної атаки;
- обґрунтовано потребу періодичного перегляду моделей через динаміку загроз.

Моделювання (1), (2) показало, що впровадження багаторівневої системи захисту знижує ймовірність атаки у 1,2 раза, залишковий ризик — на 15–18 %, а ефективність реагування підвищується на 1,08 раза. Це супроводжується зменшенням кількості хибних спрацювань на близько 16 % [9].

Розроблена BPMN-схема реагування підтвердила, що інтегрований підхід скорочує середній час локалізації кіберінцидентів.

Таким чином, результати дослідження доводять, що ефективне забезпечення безпеки АСУ критичної інфраструктури повинно базуватися на формалізованих методах оцінки ризиків, багаторівневих моделях захисту та інтегрованих механізмах реагування.

Обговорення

Отримані результати підтверджують доцільність багаторівневого підходу до кіберзахисту АСУ. Водночас вони мають певні обмеження: спрощена форма моделей ризику не враховує галузеві особливості (енергетика, транспорт, оборона) [10].

Математичний опис залишкового ризику базується на припущенні про лінійну залежність між ймовірністю атаки та ефективністю захисту, хоча ці зв'язки є складнішими. Також BPMN-схема відображає лише базовий рівень процесів без урахування людського фактора та часових затримок.

Порівняння з роботами інших авторів показало, що більшість зарубіжних досліджень орієнтовані на концептуальні моделі без деталізації технічних рішень. У цій роботі поєднано теоретичні основи з практичними методами — сегментацією, криптографією, контролем доступу, IDS/IPS, а також надано кількісну оцінку ризиків.

Подальші дослідження доцільно спрямувати на:

- адаптацію моделей ризику до окремих галузей;

- створення нелінійних моделей взаємодії загроз і засобів захисту;
- тестування рішень на промислових об'єктах;
- інтеграцію технологій ШІ для виявлення аномалій.

Висновки

Доведено, що багаторівневий підхід до захисту автоматизованих систем управління на об'єктах критичної інфраструктури підвищує їхню стійкість до кіберзагроз.

Поєднання технічних (сегментація, криптографія, IDS/IPS) та організаційних заходів знижує ризики та покращує показники реагування.

Розроблені підходи та модель оцінки ризиків можуть бути рекомендовані для апробації в галузевих умовах. Подальші дослідження буде спрямовано на тестування моделей у промислових середовищах та інтеграцію методів ШІ для автоматичного виявлення аномалій.

Джерела

1. Верховна Рада України. (2020). Про критичну інфраструктуру (Закон України №1882-IX). <https://zakon.rada.gov.ua/laws/show/1882-20#Text>
2. Кабінет Міністрів України. (2019). Про затвердження Порядку формування переліку об'єктів критичної інфраструктури (Постанова №518). <https://zakon.rada.gov.ua/laws/show/518-2019-%D0%BF#Text>.
3. Довженко, Н., Іваніченко, Є., & Костюк, Ю. (2025). Методика виявлення та локалізації кіберзагроз у хмарних середовищах з інтегрованими IoT-компонентами на основі графових моделей. Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка», 1(29), 762–776. <https://doi.org/10.28925/2663-4023.2025.29.938>
4. Олійник, Я., Платоненко, А., Черевик, В., Ворохоб, М., & Шевчук, Ю. (2025). Методи захисту інформації в технологіях IoT. Кібербезпека: освіта, наука, техніка, 3(27), 100–108. <https://doi.org/10.28925/2663-4023.2025.27.705>
5. Верховна Рада України. (2017). Про основні засади забезпечення кібербезпеки України (Закон України №2163-VIII). <https://zakon.rada.gov.ua/laws/show/2163-19#Text>
6. Костюк, Ю., Складанний, П., Рзаєва, С., Мазур, Н., Черевик, В., & Аносов, А. (2025). Особливості реалізації мережевих атак через TCP/IP-протоколи. Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка», 1(29), 571–597. <https://doi.org/10.28925/2663-4023.2025.29.915>
7. Складанний, П., Костюк, Ю., Рзаєва, С., Самойленко, Ю., & Савченко, Т. (2025). Розробка модульних нейронних мереж для виявлення різних класів мережевих атак. Кібербезпека: освіта, наука, техніка, 3(27), 534–548. <https://doi.org/10.28925/2663-4023.2025.27.772>
8. National Institute of Standards and Technology. (2018). Framework for improving critical infrastructure cybersecurity (Version 1.1). U.S. Department of Commerce. <https://nvlpubs.nist.gov/nistpubs/CSWP/NIST.CSWP.29.pdf>
9. International Society of Automation. (n.d.). ISA/IEC 62443 series of standards on industrial automation and control systems cybersecurity. Retrieved September 30, 2025, from <https://www.isa.org/standards-and-publications/isa-standards/isa-iec-62443-series-of-standards>
10. Khlaponin, Y., Kozubtsova, L., Kozubtsov, I., & Shtonda, R. (2022). Функції системи захисту інформації і кібербезпеки критичної інформаційної інфраструктури. Кібербезпека: освіта, наука, техніка, 3(15), 124–134. <https://doi.org/10.28925/2663-4023.2022.15.1241341>

11. Kozhedub, Y., Vasylenko, S., Maksymets, A., & Hyrda, V. (2021). Conceptual model of information protection of critical information infrastructure objects of Ukraine. Collection "Information Technology and Security," 9(2), 151–164. <https://doi.org/10.20535/2411-1031.2021.9.2.249889>
12. Bysaga, Y. M., Byelov, D. M., & Zaborovskyi, V. V. (2023). Artificial intelligence and copyright and related rights. Uzhhorod National University Herald. Series: Law, 2(76), 299–304. <https://doi.org/10.24144/2307-3322.2022.76.2.47>
13. International Organization for Standardization. (2024). ISO/IEC 27019:2024 Information security, cybersecurity and privacy protection — Information security controls for the energy utility industry. ISO. <https://www.iso.org/standard/85056.html>
14. Кабінет Міністрів України. (2020). Про затвердження Методики визначення критеріїв віднесення об'єктів до критичної інфраструктури (Постанова №943). <https://zakon.rada.gov.ua/laws/show/943-2020-%D0%BF#Text>

Дослідження методів та розробка рекомендацій щодо використання технології штучного інтелекту в інформаційній безпеці

Гліб Федорук^[0009-0006-4831-1555]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. У роботі досліджено методи застосування технологій штучного інтелекту для підвищення рівня інформаційної безпеки. Розроблено інтелектуальну модель класифікації мережевого трафіку з використанням логістичної регресії, швидкого дерева та градієнтного бустингу. Проведено експериментальне порівняння моделей і визначено оптимальні підходи до виявлення кіберзагроз. Отримані результати підтверджують доцільність інтеграції машинного навчання у системи моніторингу трафіку для виявлення аномалій у реальному часі.

Ключові слова: штучний інтелект, машинне навчання, кібербезпека, класифікація трафіку, градієнтний бустинг.

Research of Methods and Development of Recommendations for the Use of Artificial Intelligence Technology in Information Security

Hlib Fedoruk^[0009-0006-4831-1555]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. The study investigates the application of artificial intelligence technologies to enhance information security. An intelligent model for network traffic classification was developed using logistic regression, FastTree, and gradient boosting algorithms. An experimental comparison of the models was conducted, and optimal approaches for cyber threat detection were identified. The obtained results confirm the feasibility of integrating machine learning into network monitoring systems for real-time anomaly detection.

Keywords: artificial intelligence, machine learning, cybersecurity, traffic classification, gradient boosting.

Вступ

Актуальність теми визначається необхідністю підвищення рівня захисту інформаційних систем за рахунок інтелектуалізації процесів виявлення кіберзагроз. Застосування алгоритмів машинного навчання дозволяє формувати адаптивні моделі, що навчаються на основі накопичених даних, реагують на нові патерни атак і забезпечують своєчасне попередження про ризики [1-3]. Водночас у сучасній науковій літературі відсутні узагальнені методичні підходи до вибору оптимальних моделей і параметрів їх навчання з урахуванням обмежень обчислювальних ресурсів та специфіки мережевого середовища [4].

Наукова новизна дослідження полягає у розробленні методичних основ застосування технологій штучного інтелекту в задачах кіберзахисту, які поєднують етапи попередньої

обробки, нормалізації та класифікації мережевого трафіку із використанням ансамблевих моделей машинного навчання. Запропонований підхід дозволяє підвищити точність і стабільність розпізнавання шкідливої активності завдяки використанню комбінованих ознак і узгодженню параметрів моделей для різних сценаріїв моніторингу.

Об'єктом дослідження є процес виявлення та класифікації кіберзагроз у інформаційних системах із застосуванням інтелектуальних методів аналізу даних.

Предметом дослідження є методи та алгоритмічні моделі машинного навчання, що використовуються для ідентифікації та прогнозування шкідливої активності в мережевому трафіку.

Метою роботи є підвищення рівня інформаційної безпеки шляхом розроблення науково обґрунтованих рекомендацій щодо використання технологій штучного інтелекту для своєчасного виявлення та нейтралізації кіберзагроз.

Для досягнення поставленої мети в роботі визначено такі завдання:

- виконати аналіз сучасних підходів до застосування штучного інтелекту в інформаційній безпеці та узагальнити їх переваги й обмеження;
- дослідити основні типи кіберзагроз і вразливостей інформаційних систем, що підлягають автоматизованому виявленню;
- розробити методичні основи використання моделей машинного навчання для класифікації мережевого трафіку та визначити критерії оцінювання якості прогнозів;
- здійснити експериментальні дослідження ефективності моделей і сформулювати рекомендації щодо їх практичного застосування в системах кіберзахисту.

Методи та моделі

У межах дослідження було побудовано математичну модель для класифікації мережевих загроз, спрямовану на розпізнавання шкідливої активності в трафіку за допомогою алгоритмів машинного навчання. Задача формулюється як бінарна класифікація: для кожного спостереження x_i з множини ознак X визначається мітка $y_i \in \{0, 1\}$, де 1 – атака, 0 – нормальна подія. Метою є побудова моделі:

$$f: X \rightarrow Y, \quad (1)$$

яка наближає невідому залежність між ознаками трафіку та його класом.

У дослідженні реалізовано уніфікований конвеєр оброблення даних, у якому змінюється лише алгоритм навчання. Розглядалися три методи: логістична регресія [5], швидке дерево [6] та градієнтний бустинг [7].

Після попередньої обробки числові та категоріальні атрибути об'єднуються у спільний вектор ознак. Для числових даних застосовується min–max нормалізація:

$$x'_{ij} = \frac{x_{ij} - \min(x_j)}{\max(x_j) - \min(x_j)}, \quad (2)$$

що унеможливорює домінування окремих параметрів при навчанні моделей. Категоріальні змінні кодуються методом one-hot, що формує метричний простір ознак без штучного порядку.

Усі розглянуті моделі прогнозують умовну ймовірність події:

$$p(y = 1 | x, \theta), \quad (3)$$

де θ – параметри, набуті під час навчання. Рішення приймається за пороговим правилом:

$$\hat{y} = \begin{cases} 1, \text{ якщо } p(y = 1 | x) \geq \tau \\ 0, \text{ в іншому випадку} \end{cases} \quad (4)$$

Порогове значення τ визначається балансом між ризиком хибного спрацювання та пропуском атаки. У результаті формується адаптивна система класифікації мережевого трафіку, здатна з високою ймовірністю виявляти аномальні активності, оптимально узгоджуючи точність, повноту та швидкодію.

Результати

Для навчання та тестування моделей, призначених для виявлення кіберзагроз, у даному дослідженні використано набір даних CyberFedDefender, розміщений на платформі Kaggle [8]. Цей датасет створено спеціально для наукових цілей у галузі інтелектуального виявлення загроз, аномалій і розподілених систем кіберзахисту. Його структура відображає сучасні особливості мережевого трафіку, зокрема високий рівень динамічності, наявність різнотипних протоколів і значну варіативність поведінкових характеристик користувачів.

Особливістю набору є те, що він містить як легітимні, так і шкідливі з'єднання, сформовані на основі реалістичних сценаріїв проникнення та сканування мереж. Кожен запис представлений сукупністю числових і категоріальних атрибутів. Саме така комплексна структура дозволяє відтворювати реальні умови кіберпростору та забезпечує достатню інформативність для навчання моделей.

Процес тренування моделей реалізовано у вигляді автоматизованого конвеєра, поданого на рис. 1.

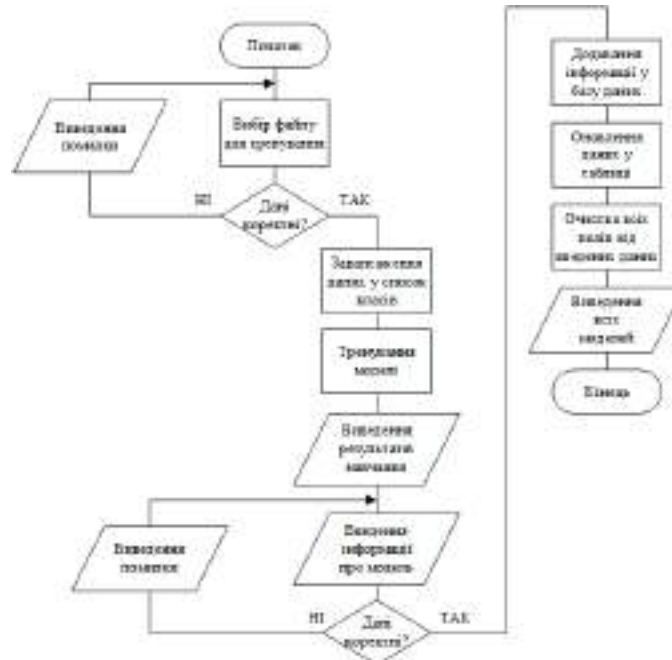


Рис. 1. Алгоритм тренування моделей

Після завершення навчання результати проходять перевірку, після чого оновлюється база даних, що містить відомості про модель – назву, тип алгоритму, параметри й отримані метрики. Така структура дозволяє зберігати історію тренувань і спрощує подальший аналіз ефективності різних методів.

На етапі тестування було реалізовано перевірку коректності роботи натренованого класифікатора в реальному середовищі програмного модуля. На рис. 2 представлено вікно застосунку Cyber Threat Detection, у якому користувач має змогу вводити параметри мережевого трафіку для прогнозування типу активності. Інтерфейс складається з двох функціональних частин: зліва – поля введення ознак, справа – текстове вікно з відображенням оброблених даних і результатів прогнозу.

Параметр	Значення
Довжина пакета	1330
Тривалість з'єднання	1.87
Порт джерела	53
Порт призначення	8080
Виправлено байтів	879
Отримано байтів	999
Пакетів за секунду	10.1
Байтів за секунду	494.2
Середній розмір пакета	512
Кількість переданих пакетів	24
К-сть зворотних пакетів	37
Заголовки (вперед)	1644
Заголовки (назад)	256
Підтвердження вперед	1644
Підтвердження назад	663
Тип трафіку	1 (1 - злодій, 0 - ніх)

Рис. 2. Приклад тестування методу швидкого дерева

Введені параметри містять основні характеристики з'єднання: довжину пакета, тривалість сесії, порти джерела та призначення, кількість переданих і отриманих байтів, частоту пакетів, середній розмір пакета, а також зведення про заголовки, підтвердження й напрям трафіку. Після натискання кнопки «Прогнозувати» модель обчислює значення ймовірності атаки на основі попередньо навченої структури дерева рішень, що використовує правила розщеплення ознак за критерієм приросту інформації.

Результати прогнозу відображаються у правій частині інтерфейсу у вигляді текстового звіту. Як видно з прикладу, система виявила шкідливу активність із ймовірністю 92,39 %. Такий формат представлення забезпечує інтерпретованість і наочність: користувач може одразу оцінити рівень загрози, перевірити, які параметри вхідного трафіку її спричинили, та перейти до моніторингу.

Розроблений інтерфейс може бути ефективно інтегрований до модулів оперативного моніторингу мережевої безпеки для автоматичного виявлення потенційних кіберзагроз.

Обговорення

Отримані результати свідчать, що запропоновані моделі машинного навчання ефективно розв'язують задачу виявлення кіберзагроз у мережевому трафіку, однак кожна з них має свої обмеження, що визначають сферу доцільного використання. Найкращі показники продемонстрував градієнтний бустинг, який завдяки ітераційному уточненню прогнозів досяг найвищої середньої точності та стабільності при роботі з різними типами атак. Разом з тим, складність ансамблевої структури потребує більших

обчислювальних ресурсів, що може обмежувати його застосування у системах з жорсткими часовими або апаратними обмеженнями.

Модель швидкого дерева показала збалансовані результати між точністю та швидкодією, проте її ефективність значною мірою залежить від якості вибірки та ступеня балансування класів. У випадку сильного дисбалансу або відсутності достатньої кількості прикладів певних атак алгоритм може схилитися до перенавчання, що вимагає періодичного оновлення моделі й ретельного підбору глибини дерев. Незважаючи на це, її простота реалізації та висока продуктивність роблять метод швидкого дерева придатним для використання в оперативних підсистемах кіберзахисту, які працюють у режимі реального часу.

Логістична регресія, своєю чергою, підтвердила статус базової інтерпретованої моделі. Її результати поступаються ансамблевим алгоритмам, однак вона забезпечує повну прозорість у поясненні впливу окремих параметрів трафіку на рішення системи. Це робить її доцільною для первинного аналізу, калібрування моделей або побудови гібридних архітектур, де лінійна модель виконує роль попереднього фільтрування даних.

Слід зазначити, що точність моделей залежить від характеристик використовуваного набору даних і може змінюватися при переході до інших типів мережевого трафіку. Крім того, дослідження проводилося на статичних даних, тому у подальших роботах доцільно розширити систему до онлайн-навчання, що дозволить адаптуватися до нових загроз у реальному часі. Перспективним напрямом також є інтеграція моделей із системами SIEM та використання методів пояснюваного штучного інтелекту для детальнішої інтерпретації рішень.

Висновки

Отримані результати підтверджують ефективність використання технологій машинного навчання для задач виявлення кіберзагроз у мережевому трафіку. Моделі логістичної регресії, швидкого дерева та градієнтного бустингу показали різний баланс між точністю, швидкодією й інтерпретованістю. Найвищі показники якості продемонстрував градієнтний бустинг, який завдяки ітераційному вдосконаленню прогнозів забезпечив точне виявлення складних аномалій типу Ransomware і DDoS. Метод швидкого дерева виявився найбільш продуктивним для оперативного моніторингу та періодичного оновлення моделей, тоді як логістична регресія зберегла значення як базова пояснювана модель для аналітичного контролю.

Попри позитивні результати, дослідження має певні обмеження. Точність моделей залежить від репрезентативності вибірки, балансу класів і якості початкових даних. Крім того, усі експерименти проводилися у статичному режимі, без урахування динамічних змін трафіку, що властиві реальним мережам. Тому впровадження механізмів адаптивного або потокового навчання, а також розширення набору ознак для урахування поведінкових характеристик користувачів є перспективними напрямками подальшого вдосконалення.

У майбутніх дослідженнях планується розробити гібридну архітектуру системи кіберзахисту з елементами онлайн-навчання та пояснюваного штучного інтелекту, що забезпечить підвищену адаптивність і прозорість процесу прийняття рішень у реальному часі.

Подяки та гранти

Результати наукових досліджень були виконані в рамках науково-дослідної роботи: «Методи та моделі забезпечення кібербезпеки інформаційних систем переробки

інформації та функціональної безпеки програмно-технічних комплексів управління критичної інфраструктури» (№ 0122U200483, КСУБГ, м. Київ).

Джерела

1. Soltani, M., Khajavi, K., Jafari Siavoshani, M., & Jahangir, A. H. (2024). A multi-agent adaptive deep learning framework for online intrusion detection. *Cybersecurity*, 7(1), 9. <https://link.springer.com/article/10.1186/s42400-023-00199-0>
2. Kantharaju, V., Suresh, H., Niranjnamurthy, M., Ansarullah, S. I., Amin, F., & Alabrah, A. (2024). Machine learning based intrusion detection framework for detecting security attacks in internet of things. *Scientific Reports*, 14(1), 30275. <https://www.nature.com/articles/s41598-024-81535-3>
3. Ejaz, A., Mian, A. N., & Manzoor, S. (2023). Life-long phishing attack detection using continual learning. *Scientific reports*, 13(1), 11488. <https://www.nature.com/articles/s41598-023-37552-9>
4. Pietrolaj, M., & Blok, M. (2024). Resource constrained neural network training. *Scientific Reports*, 14(1), 2421. <https://www.nature.com/articles/s41598-024-52356-1>
5. Shah, R. A., Qian, Y., Kumar, D., Ali, M., & Alvi, M. B. (2017). Network intrusion detection through discriminative feature selection by using sparse logistic regression. *Future Internet*, 9(4), 81. <https://www.mdpi.com/1999-5903/9/4/81>
6. Nesterenko, L., Blassel, L., Veber, P., Boussau, B., & Jacob, L. (2025). Phyloformer: Fast, Accurate, and Versatile Phylogenetic Reconstruction with Deep Neural Networks. *Molecular Biology and Evolution*, 42(4), msaf051. <https://academic.oup.com/mbe/article/42/4/msaf051/8069202?login=false>
7. Boldini, D., Grisoni, F., Kuhn, D., Friedrich, L., & Sieber, S. A. (2023). Practical guidelines for the use of gradient boosting for molecular property prediction. *Journal of Cheminformatics*, 15(1), 73. <https://link.springer.com/article/10.1186/s13321-023-00743-7>
8. Nagib, M. N., Pervez, R., Nova, A. A., Ahmed, M., Maua, J., & Sayeema, A. (2025, July). FedSAT-Detect: A Federated Learning Framework for Privacy-Preserving Cybersecurity Threat Detection in Cloud Systems. In 2025 International Conference on Quantum Photonics, Artificial Intelligence, and Networking (QPAIN) (pp. 1-6). IEEE. <https://ieeexplore.ieee.org/abstract/document/11171893>

Методи підвищення безпеки корпоративних хмарних обчислень

Олександр Трофімов^[0009-0008-5760-7803]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. Запропоновано вендор-нейтральну методику пріоритизації контролів безпеки для корпоративних хмар. Вона поєднує моделювання загроз, модель спільної відповідальності та принципи Zero Trust. Пріоритет обчислюється за правилом $((\text{Вплив} \times \text{Покриття}) / \text{Складність})$ на матриці «загроза–контроль–вплив–складність», що дає мінімальну дорожню карту з метриками (середній час виявлення/реагування (MTTD/MTTR), порушення політик, надмірні привілеї) та централізованим журналюванням і виявленням подій безпеки. Підхід узгоджено з ДСТУ ISO/IEC 27001:2023, ISO/IEC 27017/27018, вимогами NIS2 і GDPR.

Ключові слова: корпоративні хмарні обчислення; пріоритизація контролів безпеки; архітектура нульової довіри; керування ключами та секретами; журналювання.

Methods for Enhancing the Security of Corporate Cloud Computing

Olexandr Trofimov^[0009-0008-5760-7803]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. A vendor-neutral methodology is proposed for prioritizing security controls for corporate clouds. It combines threat modeling, the shared responsibility model, and Zero Trust principles. Priority is computed by the rule $((\text{Impact} \times \text{Coverage}) / \text{Complexity})$ on a “threat–control–impact–complexity” matrix, which yields a minimal implementation roadmap with metrics (MTTD/MTTR, policy violations, excessive privileges) and centralized logging and security event detection. The approach is aligned with DSTU ISO/IEC 27001:2023, ISO/IEC 27017/27018, the requirements of NIS2, and GDPR.

Keywords: corporate cloud computing; prioritization of security controls; zero trust architecture; key and secrets management; logging.

Вступ

Актуальність. Перенесення критичних процесів у хмарні сервіси робить конфігурації, ідентичності та ланцюг постачання ПЗ головними векторами ризику. Для відповідності українським та європейським вимогам потрібна вендор-нейтральна рамка Zero Trust з автоматизованим контролем конфігурацій і доступу. Це, у свою чергу, зумовлює перехід від периметрових моделей до підходу постійної перевірки довіри та автоматизованої верифікації налаштувань.

Новизна. Пропонується вендор-нейтральна методика пріоритизації контролів, що поєднує моделювання загроз, модель спільної відповідальності та принципи Zero Trust із практиками CSPM/CNAPP/CIEM (керування повноваженнями/привілеями в хмарній інфраструктурі). Результат подається як компактна матриця «загроза–контроль–вплив–складність» і стислий план впровадження з базовими метриками (MTTD/MTTR, порушення політик, надмірні привілеї) у відповідності до визначеної регуляторної бази.

Мета. Розробити практично орієнтовану рамку відбору й впровадження методів підвищення безпеки корпоративних хмарних обчислень, що забезпечує швидке та вимірюване зниження ризиків у межах вимог України та міжнародних практик.

Об'єкт. Корпоративні хмарні середовища та процеси управління їхньою безпекою (ідентичності, конфігурації, розгортання, моніторинг, реагування).

Предмет. Комплекс заходів для безпеки корпоративних хмар: керування ідентичностями та доступом; сегментація; шифрування і керування секретами; контроль конфігурацій і захист застосунків; журналювання та моніторинг; безпечний ланцюг постачання програмного забезпечення і запобігання витокам.

Часткові завдання:

- ідентифікувати домінантні загрози для корпоративних хмар і співвіднести їх із моделлю спільної відповідальності та принципами Zero Trust.
- відібрати й обґрунтувати ключові контролі та сформувану матрицю «загроза–контроль–вплив–складність».
- скласти стислу дорожню карту впровадження з мінімальним набором метрик ефективності.

Методи та моделі

Робота має прикладний характер і поєднує:

1. Скопінг-огляд регуляторних вимог і стандартів управління безпекою;
2. Моделювання загроз для корпоративних хмар;
3. Мапінг «загроза → контроль → відповідність» на базі хмарних контрольних рамок;
4. Експертне оцінювання очікуваного впливу та складності впровадження з ранжуванням пріоритетів.

Регуляторну базу становлять стандарти — ДСТУ ISO/IEC 27001:2023 [1], ISO/IEC 27017 [2], ISO/IEC 27018 [3], директива NIS2 [4] та регламент GDPR [5], спеціальна публікація NIST 800-207 [6], Закон України №2163-VIII [7], постанова Кабінету Міністрів України №518 [8], а також CSA Cloud Controls Matrix (матриця хмарних контролів, CCM) v4 як хмаро-специфічна система контролів [9].

Модель загроз

Розглядаються шість класів загроз:

- компрометація ідентичностей/токенів;
- помилки конфігурацій;
- бічний рух у мережі;
- витік/несанкціонований доступ до даних;
- ланцюг постачання ПЗ/IaC;
- периметр/DDoS.

Вибір класів обґрунтовується вимогами NIS2/GDPR [4, 5] щодо управління ризиками та заходів безпеки, а також хмарними доменами ISO/IEC [2] і контрольними доменами CCM [9].

Побудова матриці «загроза–контроль–відповідність»

1. Для кожної загрози визначається набір контрольних заходів із CCM (v4) [9] та відповідних пунктів ISO/IEC [2, 3] (для даних/хмарних ролей).
2. До кожного контролю додається посилання на відповідність: ДСТУ ISO/IEC 27001:2023 [1] (клаузи/Annex A), NIS2 [4] (вимоги до ризик-менеджменту, моніторингу, інцидент-репортування), GDPR [5] (принципи безпеки обробки).
3. Для українського контексту маркуються контролі, що підтримують виконання постанови КМУ №518 [8] (організаційно-технічні вимоги).

Результат — конденсована матриця, що уніфікує технічні й процесні контролі з нормативною прив'язкою.

Оцінювання та пріоритезація

Для кожної пари «загроза–контроль» проводиться якісно-кількісна оцінка за шкалою 1–3:

- Вплив на ризик (Impact) — очікуване зниження імовірності/наслідків;
 - Операційна складність (Complexity) — інтеграція, зміни процесів, експлуатаційні накладки;
 - Покриття (Coverage) — ширина застосовності в мультихмарі/сервісних моделях.
- Ранжування пріоритетів:

$$\text{Priority} = \frac{\text{Impact} \times \text{Coverage}}{\text{Complexity}}, \quad (1)$$

З огляду на Zero Trust [6], найвищий пріоритет отримують контролі ідентичностей/доступу, автоматизований контроль конфігурацій і базове шифрування з керованими ключами.

Метрики ефективності та валідація

Метрики: MTDD/MTTR (виявлення/усунення), частка критичних порушень політик (CSPM), відсоток надмірних прав (CIEM), час виправлення хибних конфігурацій, відсоток покриття активів логуюванням.

Валідація: порівняння базових значень до/після впровадження першої хвилі контролів; відповідність вимогам NIS2/GDPR [4, 5] щодо заходів безпеки та управління інцидентами; підтвердження узгодженості з системою управління інформаційною безпекою (ISMS) [1]. Для технічних перевірок орієнтиром слугують CSA CCM (домени журналювання та моніторингу (LOG), керування ідентичностями та доступом (IAM), безпека застосунків та інтерфейсів (AIS)) [9] і профілі CIS Benchmarks (перелік технічних контролів).

Обмеження та відтворюваність

Оцінки «вплив/складність» мають експертний характер і залежать від зрілості DevSecOps-процесів та моделі сервісів. Щоб підвищити відтворюваність, надається: чітке визначення загроз; перелік контрольних категорій і джерел відповідності; прозора формула пріоритезації; вказівки щодо метрик і періодичності вимірювань.

Це узгоджується з практикою інформаційної безпеки за 27001 та хмарними настановами 27017/27018 і CCM v4 [1–3, 9].

Результати

Внаслідок мапінгу «загроза → контроль → відповідність» і якісно-кількісного оцінювання (шкала 1–3 для впливу, складності, покриття. Пріоритет ((Вплив×Покриття)/Складність)) отримано компактний перелік першочергових контролів для корпоративних хмар.

Пріоритизація контролів

В таблиці 1 узагальнено найважливіші пари «загроза–контроль» і надано орієнтовний пріоритет. Вищий пріоритет мають контролі ідентичностей/доступу, автоматизований контроль конфігурацій і базовий захист даних.

Пріоритет орієнтовно: ((Вплив×Покриття)/Складність) у якісній шкалі: H (high) — високий, M (medium) — середній, L (low) — низький.

Ключові висновки:

- швидший та найбільший зиск дають IAM/MFA/CIEM і CSPM/IaC-гейти (високий вплив за прийнятною складністю).

- шифрування + KMS/HSM і керування секретами критично знижують ризики витоку даних та підтримують вимоги GDPR/NIS2 [4, 5].
- мікросегментація/кінцеві точки (private endpoints) істотно зменшують бічний рух, але потребують зрілої мережевої моделі.
- WAF/DDoS швидко підвищують стійкість сервісів, особливо для публічних інтерфейсів.

Таблиця 1. Пріоритизація контролів для корпоративних хмар

Загроза	Ключові контролі	Вплив	Складність	Покриття	Пріоритет
Компрометація ідентичностей/токенів	Найменші привілеї, MFA (FIDO2/OTP), PAM, CIEM, розділення ролей	Н	М	Н	Н
Помилки конфігурацій	CSPM/CNAPP, політики IaC, гейти в CI/CD, централізоване журналювання	Н	М	Н	Н
Витік/несанкціонований доступ до даних	Шифрування в русі/на зберіганні, KMS/HSM, DLP, класифікація, менеджер секретів	Н	М	Н	Н
Бічний рух у мережі (VPC/VNet)	Сегментація/мікросегментація, private endpoints, контроль egress	М	Н	М	М
Ланцюг постачання ПЗ/IaC	Підпис артефактів, SBOM, SLSA, SCA/скан залежностей, розмежування середовищ	М	М	М	М
Периметр DDoS	WAF, керований DDoS-захист, CDN, ZTNA/SASE, обмеження частоти запитів	М	Л	М	М

Мінімальна дорожня карта впровадження (≈ 90 днів)

0–30 днів (швидкі перемоги): увімкнути багатофакторну автентифікацію; провести інвентаризацію ідентичностей і прав; запустити базовий контроль стану конфігурацій хмари; централізувати журналювання; зашифрувати критичні сховища керованими ключами; впровадити керування секретами.

31–60 днів: запровадити політики інфраструктури як коду та контрольні шлюзи у конвеєрах розгортання; налаштувати перші кореляції та оповіщення в системі журналювання; організувати регулярну ротацію ключів і сертифікатів; перевести керовані сервіси на приватні кінцеві точки.

61–90 днів: виконати мікросегментацію критичних зон; увімкнути контроль вихідного трафіку; розгорнути захист вебзастосунків і базовий захист від відмови в обслуговуванні; додати відомість складових ПЗ та підпис артефактів у конвеєр; провести перегляд доступів і скоротити надмірні привілеї.

Очікувані ефекти та метрики

Безпосередні: зменшення критичних CSPM, вилучення «твердих» секретів із коду/CI, скорочення надмірних прав.

Операційні: покращення MTTD/MTTR, підвищення покриття активів журналюванням, регулярні ротації ключів/таємниць.

Узгодженість із вимогами

IAM/MFA/SIEM, журнали/SIEM підтримують Закон України про кібербезпеку [7], КМУ №518 [8], NIS2 (ризик-менеджмент, інциденти) [4] та ISO/IEC 27001 (Annex A: контроль доступу, журналювання, реагування) [1].

CSPM/IaC, шифрування/KMS/HSM узгоджуються з ISO/IEC 27017/27018 [2, 3], GDPR Art. 32 [5], доменами CSA CCM [9].

Мікросегментація, WAF/DDoS, egress-контроль відповідають вимогам мережевої безпеки (Annex A) [1] та підсилюють захист критичної інфраструктури відповідно до КМУ №518 [8] і NIS2 [4].

Обмеження застосовності

Оцінки «вплив/складність» мають експертний характер і залежать від зрілості DevSecOps-процесів, моделі сервісів (IaaS/PaaS/SaaS) і мультихмарних інтеграцій. Для відтворюваності визначено прозорі критерії пріоритизації та набір метрик (зокрема MTBD/MTTR, порушення політик, надмірні привілеї).

Обговорення

Отримані результати — пріоритизований набір контролів, мінімальна дорожня карта та узгодження з українськими вимогами і базовими міжнародними стандартами — показують, що найбільший «швидкий» вплив дають заходи навколо ідентичностей, конфігурацій та базового захисту даних.

Самокритика отриманих результатів

Запропонована пріоритизація спирається на спрощену якісну шкалу впливу, складності та покриття, тому неминуче узагальнює багатофакторні реалії впровадження (вартість, технічний борг, кадрову готовність, залежності між командами й сервісами, тощо). Оцінки мають експертний характер і не підкріплені широкою статистикою інцидентів по галузях. Мультихмарні відмінності й різниця між IaaS/PaaS/SaaS відображені лише індикатором «покриття», тож нюанси конкретних архітектур можуть бути згладжені. Так само дорожня карта є евристичною: вона не враховує цикли закупівель, формальні зміни в процесах та можливі затримки інтеграцій. Відповідність вимогам розглядається як орієнтир, а не як формальне підтвердження аудиту.

Обмеження дослідження

Робота свідомо вендор-нейтральна і не деталізує особливості конкретних провайдерів або окремих SaaS-продуктів, тож практичні кроки можуть потребувати адаптації під обрані платформи. Не проводилася вартісна оцінка (TCO/ROI), що обмежує застосування результатів у бюджетному плануванні. Галузеві відмінності (фінанси, енергетика, охорона здоров'я) та специфічні регуляторні вимоги опрацьовано лише на загальному рівні. Технічні параметри на кшталт пропускну здатності засобів захисту, впливу мікросегментації на продуктивність або сценаріїв транскордонної обробки даних залишилися поза експериментальною перевіркою. Значна частина успіху залежить від зрілості процесів DevSecOps/ITSM і культури взаємодії команд — за їх відсутності автоматизація може давати «шум» і затримки у реагуванні.

Рекомендації до використання

Результати варто застосовувати як стартову «першу хвилю» в межах формальної системи управління безпекою організації: визначити власний контекст і ризики, приземлити пріоритети на реальні архітектури та дорожні карти змін. На початку доцільно зосередитись на «швидких перемогах» — фішинг-стійкій MFA, інвентаризації ідентичностей і прав, базовому CSPM, централізованих журналах і шифруванні з керованими ключами — це швидко знижує ризики та створює основу для наступних кроків. Далі варто планово поглиблювати мережевий захист (мікросегментація, приватні кінцеві точки, egress-контроль), підсилювати периметр (WAF/DDoS) і ланцюг постачання

(SBOM, підпис артефактів), паралельно інституціоналізуючи метрики (MTTD/MTTR, частка порушення політик, надмірні права, покриття журналюванням) з регулярним переглядом пріоритетів, наприклад, щоквартально.

Висновки

У роботі запропоновано стисле, практично орієнтоване впровадження безпеки корпоративних хмар: з моделі «загроза → контроль → відповідність» виведено пріоритизацію на основі критеріїв «impact», «complexity» та «coverage» (пріоритет пропорційний (Вплив×Покриття)/Складність). Найвищий пріоритет отримали контролю ідентичностей і доступу (найменші привілеї, MFA, CIEM), автоматизованого контролю конфігурацій (CSPM/CNAPP та контрольні шлюзи (gates) в IaC або CI/CD), а також базового захисту даних (шифрування з KMS/HSM, керування секретами) і централізованого журналювання з виявленням (SIEM). Сформовано мінімальну 90-денну дорожню карту: «швидкі перемоги» першого місяця, закріплення змін у другому та поглиблення контролів у третьому, а також окреслено другу хвилю (мікросегментація, приватні кінцеві точки, egress-контроль, WAF/DDoS, SBOM/підпис артефактів).

Надалі планується кількісна валідація запропонованої пріоритизації на вибірках організацій різних галузей із урахуванням вартості (TCO/ROI).

Подяки та гранти

Результати наукових досліджень були виконані в рамках науково-дослідної роботи: «Криптографічні моделі та методи підвищення безпеки корпоративних хмарних обчислень» (КСУБГ, м. Київ).

Список джерел

1. Стандарт ДСТУ ISO/IEC 27001:2023. — Режим доступу: https://online.budstandart.com/ua/catalog/doc-page?id_doc=103359
2. Стандарт ISO/IEC 27017:2015. — Режим доступу: <https://amnafzar.net/files/1/ISO%2027000/ISO%20IEC%2027017-2015.pdf>
3. Стандарт ISO/IEC 27018:2019. — Режим доступу: <https://amnafzar.net/files/1/ISO%2027000/ISO%20IEC%2027018-2019.pdf>
4. Директива ЄС «Про заходи щодо забезпечення високого загального рівня кібербезпеки на всій території Союзу». (2022). — Режим доступу: <https://eur-lex.europa.eu/eli/dir/2022/2555/oj>
5. Регламент ЄС «Про захист фізичних осіб у зв'язку з обробкою персональних даних та про вільний рух таких даних». (2016). — Режим доступу: <https://eur-lex.europa.eu/eli/reg/2016/679/oj>
6. Rose, S., Borchert, O., Mitchell, S., Connelly, S. (2020). Zero Trust Architecture. NIST Special Publication 800-207. Gaithersburg, MD: National Institute of Standards and Technology. DOI: 10.6028/NIST.SP.800-207.
7. Закон України «Про основні засади забезпечення кібербезпеки України» № 2163-VIII (2017). — Режим доступу: <https://zakon.rada.gov.ua/laws/show/2163-19>
8. Постанова Кабінету Міністрів України № 518 «Про затвердження Загальних вимог до кіберзахисту об'єктів критичної інфраструктури» (2019). — Режим доступу: <https://zakon.rada.gov.ua/laws/show/518-2019-%D0%BF#Text>
9. Cloud Security Alliance «Cloud Controls Matrix (CCM) v4.» (2021). — Режим доступу: <https://cloudsecurityalliance.org/artifacts/cloud-controls-matrix-v4/>

Розробка тестового середовища для побудови системи захисту даних на рівні додатків в інформаційно-комунікаційній системі корпоративної мережі

Євгеній Скуратовський

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. Загрози на рівні додатків становлять один із векторів атак у корпоративних інформаційно-комунікаційних системах. Інструменти автоматизованого тестування не завжди здатні виявляти всі можливі вразливості, що зумовлює необхідність створення контрольованого середовища для оцінки ефективності засобів захисту. Розроблено віртуалізовану тестову інфраструктуру з чітким зонуванням на базі VMware Workstation Pro. Середовище включає три логічні зони: DMZ, внутрішню мережу та інструментальну зону. Проведено порівняльний аналіз чотирьох SIEM-систем за п'ятьма критеріями. Splunk продемонстрував показники інтеграції 0,9, реакції 1,0 та масштабованості 1,0.

Ключові слова: безпека додатків, тестове середовище, SIEM, DevSecOps, Zero Trust.

Вступ

Загрози безпеки додатків виникають через недостатню увагу до безпеки під час розробки програмного забезпечення, використання застарілих технологій та вразливих бібліотек, відсутність механізмів автентифікації та контролю доступу. До основних ризиків належать SQL-ін'єкції, XSS, CSRF, використання вразливих компонентів та неправильна реалізація контролю доступу.

Традиційні методи забезпечення безпеки, зокрема мережеві брандмауери та антивірусне програмне забезпечення, не забезпечують достатнього захисту від атак, спрямованих на вразливості у веб-додатках, API-інтерфейсах та мікросерверній архітектурі. Інструменти автоматизованого тестування не завжди здатні виявляти складні логічні помилки або специфічні вразливості, адже орієнтовані на пошук поширених загроз.

Метою дослідження є розробка ізольованої тестової інфраструктури для об'єктивної оцінки ефективності засобів захисту на рівні додатків у корпоративних інформаційно-комунікаційних системах без впливу на продуктивні сервіси та дані.

Методи та моделі

Тестове середовище побудовано на основі віртуалізації з використанням VMware Workstation Pro. Середовище включає сім віртуальних машин, розподілених між трьома логічними зонами згідно з принципами Zero Trust Architecture (див. рис. 1).

DMZ містить веб-сервер на базі Ubuntu Server 22.04 з Apache для імітації зовнішньодоступних додатків та SIEM-сервер для централізованого збору подій безпеки. Внутрішня мережа включає доменний контролер Windows Server 2019 з Active Directory, файловий сервер з SMB-службою та клієнтську машину Windows 10 для моделювання типової корпоративної інфраструктури. Інструментальна зона базується на Kali Linux та містить набір засобів для моделювання атак: Burp Suite, OWASP ZAP, sqlmap, Hydra, Metasploit Framework.

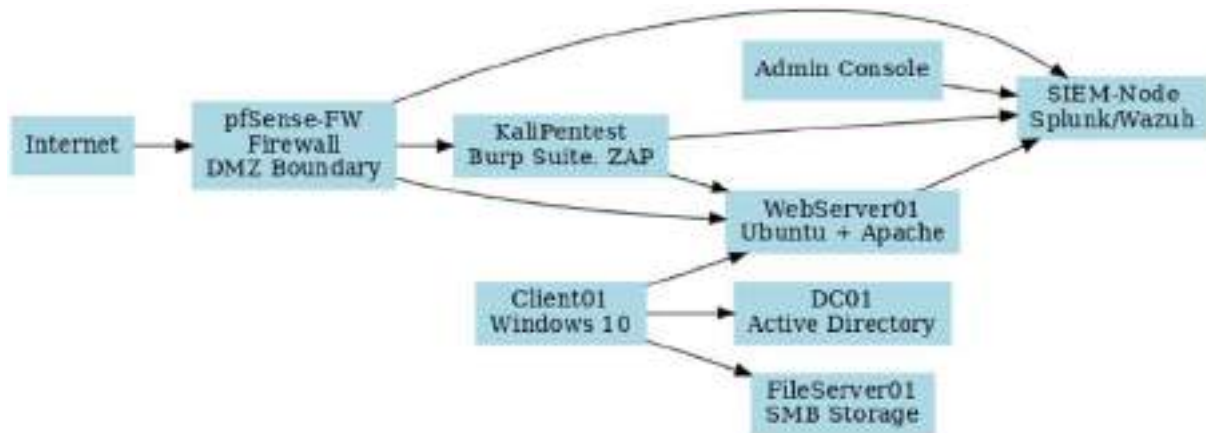


Рис. 1. Схема взаємодії компонентів тестового середовища

Розмежування між зонами реалізовано через віртуальний фаєрвол pfSense, що забезпечує контрольований доступ та дозволяє імітувати політики безпеки корпоративних мереж. Така архітектура відповідає принципам сегментації мережі згідно з ISO/IEC 27001 (розділ A.13) та NIST SP 800-53 (сімейство SC). Трафік з Internet надходить через pfSense-фаєрвол, що виконує функцію DMZ boundary. Інструментальна зона KaliPentest моделює атаки на веб-сервер та інші компоненти.

Дослідження враховує підходи до забезпечення безпеки на всіх етапах життєвого циклу розробки програмного забезпечення (див. рис. 2).

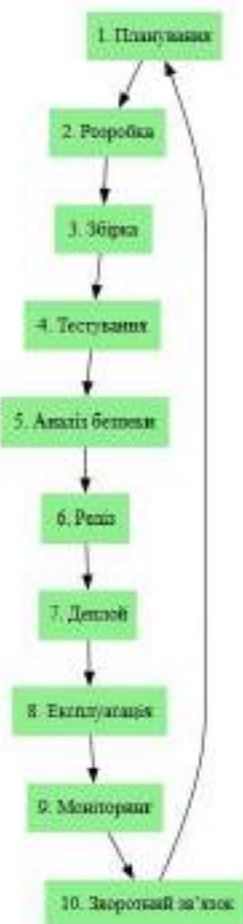


Рис. 2. Етапи процесу DevSecOps у життєвому циклі розробки

Процес DevSecOps включає десять послідовних етапів від планування до зворотного зв'язку, що дозволяє виявляти та виправляти вразливості на ранніх стадіях розробки. Тестове середовище інтегрує інструменти статичного аналізу коду, динамічного тестування безпеки та перевірки залежностей. Автоматизація процесів безпеки реалізована через механізми безперервного моніторингу та збору логів за допомогою агентів Wazuh та централізованої консолі Splunk з кореляційними правилами.

Реалізація контролю доступу базується на моделі Zero Trust (див. рис. 3), що передбачає багаторівневу перевірку на кожному етапі взаємодії користувача з ресурсами.



Рис. 3. Модель контролю доступу на основі Zero Trust Architecture

Після ідентифікації та автентифікації користувача через сервер політик доступу здійснюється детальний контроль доступу до ресурсів та даних. Система забезпечує логуювання всіх спроб доступу та моніторинг активності користувачів у режимі реального часу. Такий підхід дозволяє не лише запобігати несанкціонованому доступу, а й швидко виявляти аномальну поведінку, що може свідчити про компрометацію облікових записів.

Реалізовано сім сценаріїв тестування, що охоплюють основні типи атак на рівні додатків: SQL-ін'єкції з використанням sqlmap та Burp Suite, XSS з використанням OWASP ZAP та Burp Suite, CSRF-атаки, Brute Force за допомогою Hydra, SMB enumeration, мережеве сканування з використанням Nmap та аналіз логів через Splunk і Wazuh.

Результати

Проведено порівняльний аналіз ефективності чотирьох SIEM-систем за п'ятьма критеріями: інтеграція, аналіз, реакція, масштабування та зручність використання (див. рис. 4).

Оцінка ефективності виконувалася за нормованою шкалою від 0 до 1. Splunk продемонстрував показники інтеграції 0,9, реакції на загрози 1,0 та масштабованості 1,0. Wazuh показав результати у масштабуванні 1,0 та інтеграції 0,8, що підтверджує доцільність його використання як альтернативи з відкритим вихідним кодом. LogRhythm виявив збалансовані характеристики з акцентом на реакції 0,99, тоді як SolarWinds SEM продемонстрував нижчу ефективність у більшості категорій, зокрема в аналізі 0,6 та інтеграції 0,7.

Адміністративна консоль та SIEM-вузол забезпечують централізований моніторинг усіх компонентів системи. Внутрішня мережа, що включає клієнтську машину Windows 10, доменний контролер Active Directory та файловий сервер SMB Storage, захищена додатковим рівнем ізоляції та контролю доступу.

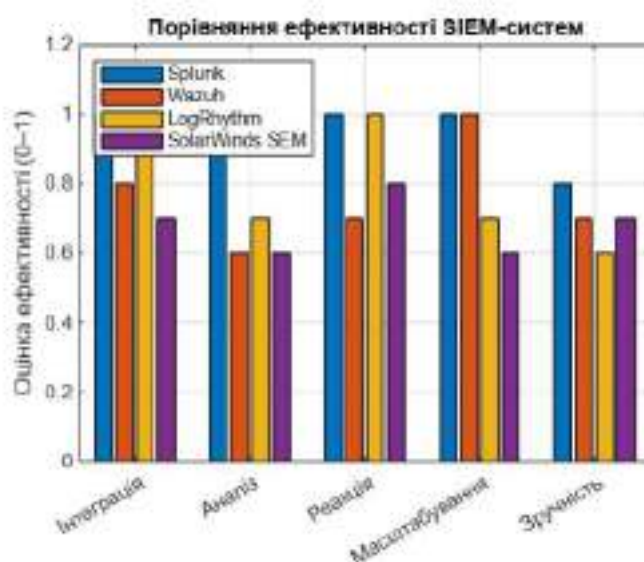


Рис. 4. Порівняльний аналіз ефективності SIEM-систем

Система логування реалізована за допомогою агентів Wazuh та централізованої консолі Splunk, що дозволяє формувати статистичні висновки щодо ефективності захисту. Налаштовані механізми охоплюють глибоке виявлення вразливостей, контроль цілісності системних компонентів, моніторинг у режимі реального часу та автоматизоване реагування на інциденти.

Обговорення

Розроблене середовище забезпечує повну мережеву ізоляцію, відтворюваність експериментів за рахунок функціоналу snapshots, різноманітність операційних систем та ролей. Модульна архітектура, повна ізоляція та підтримка обох типів тестування створюють валідну основу для експериментальної перевірки ефективності методів захисту даних.

Порівняння результатів з попередніми дослідженнями показує, що інтеграція DevSecOps-принципів у тестове середовище дозволяє виявляти вразливості на ранніх стадіях розробки. Обмеженням дослідження є використання лише чотирьох SIEM-систем для порівняльного аналізу. Рекомендується застосовувати розроблене середовище для навчальних цілей та проведення регулярних аудитів безпеки корпоративних додатків.

Висновки

Розроблено віртуалізовану тестову інфраструктуру з чітким зонуванням для оцінки ефективності засобів захисту на рівні додатків. Середовище включає три логічні зони та сім віртуальних машин, що дозволяє моделювати реалістичні сценарії атак. Проведено порівняльний аналіз чотирьох SIEM-систем, що показав перевагу Splunk за показниками інтеграції 0,9, реакції 1,0 та масштабованості 1,0. Інтеграція DevSecOps-принципів та механізмів Zero Trust забезпечує виявлення вразливостей на ранніх стадіях розробки.

Подальші дослідження передбачають розширення кількості сценаріїв тестування та впровадження механізмів машинного навчання для автоматизованого виявлення аномалій.

Джерела

1. OWASP Foundation (2024). OWASP Top Ten Security Risks. <https://owasp.org/www-project-top-ten/>
2. National Institute of Standards and Technology (2024). NIST SP 800-53. <https://csrc.nist.gov/publications/detail/sp/800-53/>
3. ISO/IEC 27001 (2022). Information security management systems.
4. Splunk Inc. (2024). Enterprise Security. https://www.splunk.com/en_us/products/enterprise-security.html
5. Wazuh Inc. (2024). Open Source Security Platform. <https://wazuh.com/>
6. VMware Inc. (2024). Workstation Pro Documentation.
7. Kali Linux (2024). Penetration Testing Tools. <https://www.kali.org/tools/>
8. SANS Institute (2024). DevSecOps Integration. <https://www.sans.org/cyber-security-courses/devsecops/>

Дослідження ефективності системи акустичного зашумлення від витоку акустичної інформації

Валентин Романюк^[0009-0004-8965-3098]

Артем Платоненко^[0000-0002-2962-5667]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. У роботі проведено експериментальне дослідження ефективності системи активного акустичного зашумлення (СААЗ) для захисту мовної інформації від витоку віброакустичним каналом. На модельному стенді, що імітує типові умови, було продемонстровано, що без активного захисту словесна розбірливість мови за межами приміщення становить 97%. Застосування СААЗ дозволило знизити цей показник до 9%, що підтверджує високу ефективність системи та її відповідність нормативним вимогам технічного захисту інформації.

Ключові слова: технічний захист інформації, віброакустичний канал, активне зашумлення, розбірливість мови.

Research into the effectiveness of an acoustic noise suppression system against acoustic information leakage

Valentyn Romaniuk^[0009-0004-8965-3098]

Artem Platonenko^[0000-0002-2962-5667]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. The paper presents an experimental study of the effectiveness of the active acoustic noise suppression system (AAS) for protecting speech information from leakage through a vibroacoustic channel. On a model stand simulating typical conditions, it was demonstrated that without active protection, speech intelligibility outside the room is 97%. The use of AAS allowed to reduce this figure to 9%, which confirms the high effectiveness of the system and its compliance with regulatory requirements for technical information protection.

Keywords: technical information protection, vibroacoustic channel, active noise reduction, speech intelligibility.

Вступ

Актуальність теми зумовлена критичною важливістю захисту мовної інформації в сучасному світі та поширенням високочутливих засобів технічної розвідки, здатних перехоплювати акустичні коливання через огорожувальні конструкції будівель. Віброакустичний канал витоку є одним з найбільш небезпечних, оскільки не вимагає фізичного доступу до приміщення, що захищається, і може бути реалізований через стіни, вікна, інженерні комунікації [1, 2].

Наукова новизна роботи полягає у проведенні комплексного моделювання експерименту та кількісної оцінки ефективності СААЗ для типової монолітної залізобетонної конструкції. На основі експериментальних даних уточнено залежність

нормативних показників захищеності від параметрів системи зашумлення, що підвищує точність прогнозування ефективності захисту на етапі проектування.

Метою дослідження є кількісна оцінка ефективності сучасної системи активного акустичного зашумлення для блокування віброакустичного каналу витоку.

Об'єктом дослідження є процес поширення мовної інформації у віброакустичному каналі та процес її маскування завадою.

Предметом дослідження є ефективність СААЗ, що оцінюється за критеріями зниження розбірливості мови.

Завданням для даного дослідження є: аналіз теоретичних засад, розробку методики експерименту, проведення вимірювань рівнів сигналу та завади, розрахунок показників захищеності та формулювання висновків і рекомендацій.

Методи та моделі

Дослідження ґрунтується на аналізі фізичних механізмів формування віброакустичного каналу витоку інформації, який виникає внаслідок перетворення акустичної енергії мовного сигналу в механічні коливання огорожувальних конструкцій [3–6]. Попередні дослідження підтверджують, що ефективність цього каналу залежить від маси, жорсткості та герметичності конструкцій, причому пасивних методів захисту (звукоізоляції) часто буває недостатньо [7, 8].

Активні методи, зокрема застосування СААЗ, передбачають генерацію маскуючої завади, яка робить неможливим виділення корисного сигналу. Ефективність таких систем оцінюється за нормативними критеріями, встановленими в Україні (НД ТЗІ), що базуються на акустичному показнику – розбірливості мови.

В основі методики лежить інструментально-розрахунковий формантний метод оцінки, регламентований НД ТЗІ 2.3-017-08 «Методика контролю захищеності мовної інформації від витоку акустичним та віброакустичним каналами» [9]. Ключовим фізичним параметром є відношення сигнал/шум (SNR), що розраховується в п'яти октавних смугах частот (250, 500, 1000, 2000, 4000 Гц). На його основі обчислюється інтегральний індекс артикуляції (R) за формулою:

$$R = \sum_{i=1}^5 W_i L_i \quad (1)$$

де W_i – вагові коефіцієнти значущості кожної октавної смуги для розбірливості мови, а L_i – коефіцієнт сприйняття, що є функцією від SNR в цій смузі:

$$L_i = \begin{cases} 0, & SNR_i \leq -15\text{дБ} \\ \frac{SNR_i + 15}{30}, & -15\text{дБ} < SNR_i < +15\text{дБ} \\ 1, & SNR_i \geq +15\text{дБ} \end{cases} \quad (2)$$

Далі, за допомогою стандартних емпіричних кривих, значення індексу артикуляції R перераховується в кінцевий критерій – словесну розбірливість (W , %).

Для проведення дослідження було змодельовано експериментальний стенд: приміщення для переговорів, огорожене монолітною залізобетонною стіною товщиною 200 мм. Всередині приміщення розміщувалось джерело тестового сигналу ("рожевий шум" з рівнем звукового тиску 70 дБ), а на внутрішню поверхню стіни монтувались вібровипромінювачі СААЗ (PIAC–2ВП, PIAC–2ГС), що генерували мовленнєво подібний шум. Вимірювання рівнів віброприскорення сигналу (L_s) та завади (L_n) проводились на зовнішній поверхні стіни за допомогою прецизійного акселерометра та аналізатора спектра.

Результати

Перед проведенням дослідження було виміряно індекс артикуляції (R) та словесну розбірливість (W) у приміщенні при вимкненій СААЗ. За результатами розрахунку отримано, що $R = 0.85$, що відповідає словесній розбірливості $W \approx 97\%$.

Це свідчить про те, що без застосування СААЗ розбірливість мовлення через огорожувальну конструкцію є дуже високою, тобто можливий витік мовної інформації за межі приміщення. Тепер проведемо обрахунки при увімкненій СААЗ та отримаємо дані про рівні віброприскорення в контрольній точці на зовнішній поверхні стіни. Узагальнені результати вимірювань та розрахунків наведені в таблиці 1.

Таблиця 1. Вимірювання значень дослідів

Частота, Гц	Рівень сигналу L_s , дБ	Рівень завади, L_n , дБ	Відношення сигнал/шум, SNR, дБ	Вагові коефіцієнти, W_i	Коефіцієнт сприйняття, L_i
250	38	52	-14	0,113	0.0333
500	35	48	-13	0,205	0.0667
1000	29	41	-12	0,207	0.1
2000	21	35	-14	0,275	0.0333
4000	14	29	-15	0,200	0

Після вмикання СААЗ, за тих самих умов експерименту, індекс артикуляції знизився до $R = 0.0473 \approx 0.05$, що відповідає словесній розбірливості $W \approx 9\%$.

Обговорення

Отримані результати наочно демонструють критичну різницю в стані захищеності об'єкта до та після застосування активних засобів захисту, які наведені в таблиці 2.

Таблиця 2. Оцінка станів захищеності об'єкта

Стан системи захисту	Індекс артикуляції R	Словесна розбірливість W , %	Висновок про захищеність
СААЗ вимкнено	0,85	97	Захист відсутній
СААЗ увімкнено	0,05	9	Захист забезпечено

Результат у 97% розбірливості за відсутності СААЗ, незважаючи на масивну бетонну стіну, підтверджує висновки інших дослідників про недостатність виключно пасивних методів захисту. Кардинальне падіння розбірливості до 9% після активації системи пояснюється тим, що СААЗ створює настільки інтенсивну заваду, що відношення сигнал/шум опускається значно нижче, що значно ускладнює розбірливість та сприйняття мови.

Слід зазначити, що дане дослідження є моделюванням і проводилось для конкретного типу конструкції (монолітний залізобетон 200 мм). Ефективність системи для інших типів конструкцій (наприклад, цегляна кладка, легкі перегородки) може відрізнятися і потребує окремих досліджень.

Висновки

Проведене дослідження дозволило кількісно оцінити ефективність системи активного акустичного зашумлення. Встановлено, що за відсутності активного захисту

віброакустичний канал через монолітну залізобетонну стіну товщиною 200 мм дозволяє перехоплювати мовну інформацію з розбірливістю 97%. Застосування СААЗ забезпечує зниження словесної розбірливості до 9%, що є значно нижчим за нормативне порогове значення (15%) і гарантує надійний захист інформації. Таким чином, доведено, що сучасні системи активного зашумлення є високоефективним та необхідним компонентом комплексної системи технічного захисту інформації на об'єктах інформаційної діяльності.

Майбутні дослідження можуть бути спрямовані на вивчення ефективності адаптивних систем зашумлення, які динамічно змінюють параметри завади залежно від характеристик інформативного сигналу.

Подяки та гранти

Результати наукових досліджень були виконані в рамках науково-дослідної роботи: «Методи та моделі забезпечення кібербезпеки інформаційних систем переробки інформації та функціональної безпеки програмно-технічних комплексів управління критичної інфраструктури» (№ 0122U200483, КСУБГ, м. Київ).

Джерела

1. Гуменюк, І., Костерев, Д., & Шейгас, В. (2024). Методика оптимізації кількості засобів блокування каналів витоку акустичної інформації на об'єкті інформаційної діяльності. *Ukrainian Scientific Journal of Information Security*, 30(2), 289–296. <https://doi.org/10.18372/2225-5036.30.19241>
2. Danilov, V. V., & Kotenko, A. M. (2020). Directions of protection of acoustic information on the object of information activity. *Modern Information Security*, 44(4). <https://doi.org/10.31673/2409-7292.2020.041822>
3. Sarampalis, A., Kalluri, S., Edwards, B., & Hafter, E. (2009). Objective Measures of Listening Effort: Effects of Background Noise and Noise Reduction. *Journal of Speech, Language, and Hearing Research*, 52(5), 1230–1240. [https://doi.org/10.1044/1092-4388\(2009/08-0111\)](https://doi.org/10.1044/1092-4388(2009/08-0111))
4. Романюк, В., & Платоненко, А. (2025). Аналіз якості генераторів рожевого шуму. Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка», 1(29), 789–799. <https://doi.org/10.28925/2663-4023.2025.29.940>
5. Лізунов, С., & Філобок, Є. (2022). Удосконалена система активного придушення звуку. *Ukrainian Information Security Research Journal*, 23(4), 221–225. <https://doi.org/10.18372/2410-7840.23.16768>
6. Романюк, В., & Платоненко, А. (2025). Аналіз якості генераторів рожевого шуму. Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка», 1(29), 789–799. <https://doi.org/10.28925/2663-4023.2025.29.940>
7. Ликов, Ю. В., Морозова, Г. Д., & Кукуш, В. Д. (2017). Influence of features of information leakage channels on intelligibility of eavesdropped voice messages. *Technology Audit and Production Reserves*, 1(2(33)), 4–8. <https://doi.org/10.15587/2312-8372.2017.90571>
8. Prodeus, A. M., Bukhta, K. V., Morozko, P. V., Serhiienko, O. V., Kotvytskyi, I. V., & Dvornyk, O. O. (2018). Automated Subjective Assessment of Speech Intelligibility in Various Listening Modes. *Microsystems, Electronics and Acoustics*, 23(3), 49–57. <https://doi.org/10.20535/2523-4455.2018.23.3.130367>
9. Методика контролю захищеності мовної інформації від витоку акустичним та віброакустичним каналами [Текст]: НД ТЗІ 2.3-017-08. [Чинний від 2008-08-26]. – К., 2008. – 18 с. – (Нормативний документ системи технічного захисту інформації).

Когнітивне моделювання як інструмент інформаційної безпеки

Аріна Гаркушенко^[0009-0003-7568-7900]

Світлана Шевченко^[0000-0002-9736-8623]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. У тезах розглядається когнітивне моделювання як інструмент управління інформаційними ризиками в контексті зростаючої динамічності та складності кіберзагроз. Зокрема розглянуто когнітивні карти як потужний інструмент для аналізу ризиків, виявлення уразливих місць у системі та оцінки ефективності засобів захисту. Розглянуто принцип їх функціонування, побудови, а також можливі переваги і недоліки. Приділено увагу можливості застосування нечітких когнітивних карт для прогнозування загроз та створення сценаріїв захисту. Представлено методику проведення тренінгу з моделювання поведінки атакуючих і внутрішніх користувачів у сценаріях з використанням нечітких когнітивних карт.

Ключові слова: інформаційна безпека, когнітивне моделювання, загрози, нечіткі когнітивні карти.

Cognitive Modeling as an Information Security Tool

Arina Harkushenko^[0009-0003-7568-7900]

Svitlana Shevchenko^[0000-0002-9736-8623]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. The paper considers cognitive modeling as a tool for managing information risks in the context of the growing dynamism and complexity of cyber threats. In particular, cognitive maps are considered as a powerful tool for analyzing risks, identifying vulnerabilities in the system and assessing the effectiveness of protection measures. The principle of their functioning, construction, as well as possible advantages and disadvantages are considered. Attention is paid to the possibility of using fuzzy cognitive maps for threat prediction and creating protection scenarios. A training methodology for modeling the behavior of attackers and internal users in scenarios using fuzzy cognitive maps is presented.

Keywords: information security, cognitive modeling, threats, fuzzy cognitive maps.

Вступ

Під потенційними ризиками ІБ ми розуміємо числову (словесну) функцію, яка описує ймовірність втілення загроз ІБ та величини збитку від їх реалізації внаслідок використання цими загрозами уразливостей активів з метою нанесення шкоди організації [1]. Основним принципом управління ризиками ІБ є безперервна оцінка та адаптація, що має на меті постійне оновлення як оцінок ризиків, так і стратегій їх усунення [2].

Традиційні підходи до оцінювання ризиків ІБ обмежені у можливості врахування нечітких зв'язків, впливу людського чинника та динамічної природи загроз. У цьому

контексті перспективним інструментом стає когнітивне моделювання, яке є корисним інструментом для виявлення уразливих місць в системах безпеки та оцінювання ефективних заходів для їх усунення, а також надає відповідальним особам інструмент для аналізу різних сценаріїв та обґрунтованого прийняття рішень. Візуалізація цього процесу здійснюється за допомогою нечітких когнітивних карт [3-4].

Метою дослідження є підвищення ефективності захисту інформації на основі аналізу ризиків інформаційної безпеки з використанням нечітких когнітивних карт.

Об'єктом дослідження є процес управління ризиками інформаційної безпеки в умовах невизначеності.

Предметом дослідження є методи когнітивного моделювання та їх застосування для аналізу та прогнозування ризиків у кібербезпеці.

Частковими завдання даного дослідження є розкриття сутності когнітивного моделювання з використанням нечітких когнітивних карт та принципів їх побудови.

Методи та моделі

Когнітивна карта складається з вершин (понять) та дуг (впливів), що формують сценарії розвитку ситуацій при зміні окремих факторів. На основі таких моделей можна досліджувати зміну загального рівня ризику в системі з урахуванням взаємозв'язків відповідно до певних загроз, уразливостей чи активів.

Нечіткі когнітивні карти (НKK) – це тип когнітивної карти, який використовує нечітку логіку для представлення та аналізу взаємозв'язків між різними змінними в системі. Основна ідея полягає в моделюванні компонентів системи (змінних або факторів) та того, як вони впливають один на одного, з метою розуміння та прогнозування поведінки системи в цілому [4].

Поняття представлені у вигляді вузлів, а причинно-наслідкові зв'язки представлені у вигляді дуг між вузлами в НKK. Вузли або поняття вказують на інформацію про досліджувану систему, таку як атрибути, характеристики, якості, змінні та стани [5]. Поняття можуть мати прямі або непрямі зв'язки між собою або не мати жодних зв'язків [6]. Зв'язки між вузлами можуть бути позитивними, негативними або нейтральними, виражаючи тип зв'язку між поняттями та ступінь причинно-наслідкового зв'язку. На рис. 1 показано приклад методу НKK з його компонентами.

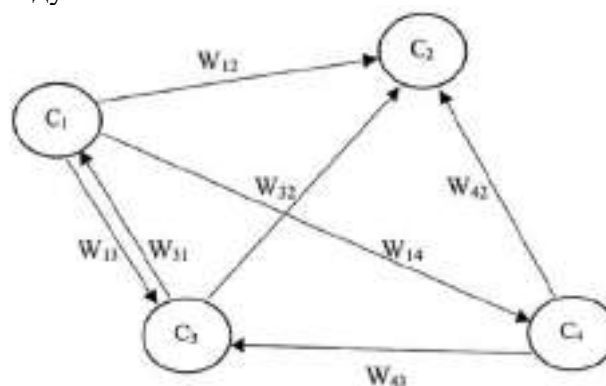


Рис. 1 Зразок НKK

C_i виражає вузли або поняття, пов'язані з зваженими дугами [7]. Кожен зв'язок між поняттями C_i та C_j має вагу W_{ij} (вага між вузлом i та вузлом j), яка може бути позитивною, негативною або нейтральною (вказує на те, що два поняття, що розглядаються, не мають зв'язку).

Після визначення значення вузла необхідно визначити інші вузли, пов'язані з цим вузлом, за допомогою рівняння (1):

$$A_i^{(k+1)} = f \left(A_i^{(k)} + \sum_{\substack{j=1 \\ j \neq i}}^n W_{ij} A_j^{(k)} \right) \quad (1),$$

де $A_i^{(k+1)}$ – значення концепції C_i в ітерації, $A_i^{(k)}$ – значення поняття C_i в ітерації, $f(x)$ – функція перетворення.

Результати

Говорячи про розробку та застосування когнітивних карт у системах управління безпеки, можна виокремити такі практичні напрямки як формування сценаріїв кіберінцидентів на основі когнітивних моделей, оцінка ризиків на базі сценарного аналізу з урахуванням когнітивних факторів, моделювання поведінки атакуючих і внутрішніх користувачів у сценаріях, впровадження когнітивних сценаріїв у процес прийняття рішень (SOC, CERT) або ж автоматизація моніторингу та коригування ризиків за допомогою когнітивних моделей.

У даній роботі ми зупинились на одному з них, а саме створенні тренінгів з моделювання поведінки атакуючих і внутрішніх користувачів у сценаріях з використанням нечітких когнітивних карт. Обґрунтовано це тим, що такі тренінги безпосередньо підвищують готовність команд реагування, переводять теоретичні ризики в відпрацьовані сценарії реагування. Це підходить для більш обмежених даних, що робить реалізацію більш швидкою та реалістичною на початкових етапах, враховує невизначеності та причинно-наслідкові ланцюжки (сценарії внутрішньої взаємодії та поведінка під час атаки часто залежать від ледь помітних людських факторів (мотивація, доступ, стрес) – НКК корисні там, де причинно-наслідкові зв'язки нечіткі та взаємозалежні [8]), враховує інтерпретованість та навчання персоналу, що важливо для вміння практичного застосування і корекції сценаріїв у режимі реального часу в подальшому, а також має відносну простоту впровадження, що дозволяє швидко перейти від дослідження до практичної реалізації.

На початковому етапі важливо сформулювати мету і цілі тренінгу. В нашому випадку це можливість виявляти ознаки інсайдерської або зовнішньої атаки на початкових етапах, відпрацювання стратегії і прийняття рішень та поліпшення координації групи аналітиків у умовах можливого тиску чи невизначеності.

Наступним етапом є визначення концептів НКК. В нашому випадку ми будемо список концептів, розбитих на групи: технічні («вхід з незвичайної IP-адреси», «підозріла активність», «права доступу»), поведінкові/людські («мотивація», «емоційний стрес», «задоволення від роботи») та процесуальні («час реагування», «доступність аудитів», «політика найменших привілеїв»).

Після чого ми переходимо до збирання експертних даних, оцінок та визначення ваг. На цьому кроці проводяться воркшопи з аналітиками та менеджерами безпеки для того, щоб визначати зв'язки між концептами (наприклад, «зменшення задоволеності» призводить до «зростання ймовірності інсайдерської шкоди»), оцінювати силу впливу у нечітких термінах (наприклад, «низький/середній/високий» переходить у конкретні числові ваги або нечіткі множини). Після чого відбувається усереднення та узгодження експертних оцінок, яке можливе, наприклад, порівнянням декількох створених НКК і їх обговоренням на експертній сесії [9].

Наступним кроком ми записуємо, використовуючи вже отримані дані, матрицю W ($n \times n$) ваг зв'язків. Коли ми запускаємо ітерації оновлення стану, ми отримуємо

динамічний сценарій. Для тренінгу ми можемо використовувати кілька конфігурацій початкових станів, які відрізняються в різних сценаріях [10].

Далі ми переходимо до генерації сценаріїв. В нашому випадку ми формуємо сценарій А ((повільне внутрішнє шахрайство), початковими значеннями якого є «задоволеність» = низький, «доступ» = високий, «контроль» = низький. Модель показує поступове збільшення «коефіцієнта витоку» з часом) та сценарій В ((швидка зовнішня компрометація), початковими даними якого є «успішність фішингу» = високий, «спільне використання прав» = середній. Модель показує швидкий стрибок у «доступі» та «експорті даних»). Кожен сценарій відтворюється через ключові ітерації НКК, учасники аналізують проміжні стани та приймають на основі цього рішення.

Наступним етапом є інтерактивна сесія тренінгу, тобто тепер учасники отримують початкові спостереження (лог-події) та візуалізації НКК. Протягом ітерацій організатор чи ведучий тренінгу дає додаткові фактори або події, згідно чому учасники приймають рішення (наприклад, блокування, опитування користувачів, ігнорування), після чого обговорюють прогнозовані події і бачать наслідки своїх рішень у оновленій моделі.

Метриками оцінювання успішності та ефективності запропонованого тренінгу є кількість помилкових/позитивних рішень, покращення узгодженості дій у команді (згідно опитування чи оцінку експертів) та час виявлення загроз у різних сценаріях.

Далі ми розглянемо приклад створення простої НКК для тренінгу. У нашій спрощеній демонстраційній моделі ми розглянемо п'ять основних факторів (вершин):

- мотивація атакуючого (A1);
- рівень підготовки атакуючого (A2);
- необережність користувачів (U1);
- політика безпеки та навчання (U2);
- рівень загрози (R).

Після визначення основних компонентів ми переходимо до визначення зв'язків (ребер). У нашому випадку:

A1 → R (чим вища мотивація атакуючого, тим вищий ризик, вага +0,7);

A2 → R (чим більше підготовлений атакуючий, тим вищий ризик, вага +0,9);

U1 → R (необережність користувачів збільшує ризик, вага +0,6);

U2 → U1 (навчання зменшує необережність, вага -0,8);

U2 → R (навчання знижує ризик, вага -0,5).

На основі цих даних створюємо матрицю ваг W для НКК:

	A1	A2	U1	U2	R
A1	0	0	0	0	0,7
A2	0	0	0	0	0,9
U1	0	0	0	0	0,6
U2	0	0	+0,8	0	0,5
R	0	0	0	0	0

На основі цього прописуємо можливі сценарії. Наприклад, у сценарії 1 «недостатнє навчання персоналу» ми маємо на меті показати як дефіцит політики і навчання (U2) підвищує необережність користувачів (U1) і в підсумку збільшує загальний рівень ризику (R). У сценарії 2 «високий рівень підготовки атакуючого» ми маємо на меті відпрацювати реагування та майстерно виконану швидку атаку, коли атакуючий має високі навички і може використати слабкі точки у безпеці системи, навіть якщо персонал має певні знання (U2 середнє). І у сценарії 3 «сильна політика безпеки та навчання» ми

показуємо як висока якість політики безпеки та навчання (U2) знижують необережність користувачів (U1) і у підсумку загальний ризик (R), навіть коли є помірна/висока загроза атакуючого. У цьому сценарії ми вчимо команду фокусуватись не лише на технічних аспектах захисту, а і на профілактичних.

На основі цих даних ми створюємо власне візуальну ілюстрацію карти, граф, де вузли (кола) – це наші фактори (A1, A2, U1, U2, R), стрілки – причинно-наслідкові зв'язки між вагами (позитивні – червоні, негативні – сині). Наприклад, стрілка від U2 до U1 з негативною вагою показує, що підвищення навчання персоналу знижує необережність.

Учасники тренінгу можуть змінювати значення вхідних факторів (A1, A2, U1, U2), запускати модель та спостерігати, як змінюється рівень ризику (R). Це дозволяє зрозуміти вплив поведінки користувачів, побачити різні траєкторії розвитку кіберінциденту та усвідомити важливість навчання персоналу та політик безпеки.

Обговорення

Когнітивне моделювання має низку переваг: дозволяє візуалізувати складні сценарії; забезпечує інтеграцію знань різних експертів. Однак є й обмеження: модель залежить від якості експертних оцінок; складність масштабування для великих систем; необхідна адаптація у разі швидких змін у середовищі загроз.

Порівняно з класичними методами (наприклад, SWOT-аналіз), когнітивні карти є більш динамічними та інформативними.

Висновки

Когнітивне моделювання є ефективним підходом для управління ризиками інформаційної безпеки, особливо в умовах невизначеності, який має свої переваги та недоліки. Моделі, розроблені на основі нечітких когнітивних карт, можуть використовуватись як основи для створення систем підтримки прийняття рішень в управлінні інформаційними ризиками. За допомогою когнітивного підходу можна створити детальну модель системи безпеки, яка показує, як усі її елементи взаємопов'язані, як вони впливають один на одного та які наслідки можуть виникнути в результаті різних подій. Проте треба відзначити, що при побудові необхідно враховувати якість експертних знань та масштабування системи задля уникнення хибних результатів.

Слід також зазначити, що великі моделі потребують значних обчислювальних ресурсів для обробки даних. Тому перспективними напрямками подальших досліджень є розробка адаптивних моделей, здатних до самонавчання, а також інтеграція нечітких когнітивних карт з іншими методами штучного інтелекту.

Подяки та гранти

Результати наукових досліджень були виконані в рамках науково-дослідної роботи: «Методи та моделі забезпечення кібербезпеки інформаційних систем переробки інформації та функціональної безпеки програмно-технічних комплексів управління критичної інфраструктури» (№ 0122U200483, КСУБГ, м. Київ).

Джерела

1. Шевченко, С.М., Жданова, Ю.Д., Спасітелєва, С.О. & Складанний П.М. (2020) «Проведення swot-аналізу оцінювання інформаційних ризиків як засіб формування

практичних навичок студентів спеціальності 125 Кібербезпека». Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка», 2(10) с. 158-168.

2. ДСТУ ISO/IEC 27005:2023 Інформаційні технології. Методи захисту. Управління ризиками інформаційної безпеки

3. Shevchenko S., Zhdanova Y., Kryvytska O., Shevchenko H., Spasiteleva S. (2024) Fuzzy cognitive mapping as a scenario approach for information security risk analysis Cybersecurity Providing in Information and Telecommunication Systems II 2024, 3826. с. 356-362. ISSN 1613-0073

4. Шевченко, С., Жданова, Ю., Складанний, П., & Петренко, Т. (2024). «Нечіткі когнітивні карти як інструмент візуалізації сценаріїв реагування на інциденти в системах безпеки». Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка», 2(26), 417–429.

5. Aguilar J (2003) A dynamic fuzzy-cognitive-map approach based on random neural networks. Int J Computat Cogn 1(4):91–107

6. Kang I, Lee S, Choi J (2004) Using fuzzy cognitive map for the relationship management in airline service. Expert Syst Appl 26(4):545–555

7. Papageorgiou EI, Spyridonos PP, Glotsos DT, Stylios CD, Ravazoula P, Nikiforidis GN, Groumpos PP (2008) Brain tumor characterization using the soft computing technique of fuzzy cognitive maps. Appl Soft Comput 8(1):820–828

8. Sarmiento, I., Cockcroft, A., Dion, A. et al. Fuzzy cognitive mapping in participatory research and decision making: a practice review. Arch Public Health 82, 76 (2024). <https://doi.org/10.1186/s13690-024-01303-7>

9. Yoon, Byung Sung and Jetter, Antoine J., "Comparative Analysis for Fuzzy Cognitive Mapping" (2016). Engineering and Technology Management Faculty Publications and Presentations. 112.

10. Kosko B (1986) Fuzzy cognitive maps. Int J Man Mach Stud 24(1):65–75.