

адаптивної системи безпеки організації. Більше можливостей демонструє впровадження інноваційних рішень, якими є використання штучного інтелекту у напрямках: виявлення аномалій, щоб відстежувати поведінку мережі та автоматично реагувати на загрози; прогнозування кібератак і використання сучасних технологій протидії різним кібератакам. Активно використовуються сучасні інноваційні технології: SOAR для автоматизації реагування, XDR для розширеного виявлення загроз, Digital Twins для оцінки інфраструктури, Cyber Range для використання як кіберполігонів у дослідженнях.

1. Звіт ДДКЗ Держспецзв'язку про роботу Системи виявлення вразливостей і реагування на кіберінциденти та кібератаки (СВВ) за 2025 рік. URL:
2. <https://cip.gov.ua/services/cm/api/attachment/download?id=73033>
3. Ільєнко А.В., Телющенко В.А., Дубчак О.В. Сучасні кіберзагрози критичної інфраструктури України та світу № 3 (27), 2025. DOI 10.28925/2663-4023.2025.27.719
4. Постанова КМУ від 13.11.2025 р. № 1470. Загальні вимоги до кіберзахисту об'єктів критичної інфраструктури (ОКІ) в новій редакції.

Трасування безпекових вимог у системах предиктивної аналітики

УДК 004.89:339.138 (043.2)

Дмитро Яценко¹, Володимир Садовенко²

*Державний університет інформаційно-комунікаційних технологій,
¹d.yatsenko@stud.duikt.edu.ua, ²v.sadovenko@duikt.edu.ua*

Системи предиктивної аналітики (СПА) у цифровому маркетингу обробляють масивні поведінкові, демографічні та транзакційні дані користувачів, що формують специфічну поверхню атаки. Поряд із класичними загрозами інформаційній безпеці такі системи зазнають впливу атак, специфічних для машинного навчання (МН): data poisoning, evasion, model extraction, membership inference, model inversion [1, 2, 3]. Чинні стандарти управління ризиками штучного інтелекту, зокрема ISO/IEC 23894:2023 [4] та NIST AI 100-2 E2025 [1], формують принципи високого рівня, проте не пропонують архітектурного інструментарію проєктування. Безпекові механізми впроваджуються на пізніх етапах життєвого циклу системи, що знижує стійкість і ускладнює верифікацію її властивостей.

Мета роботи — підвищення стійкості СПА до атак на МН шляхом розроблення методу трасування безпекових вимог до архітектурних компонентів, механізмів контролю та метрик верифікації, що враховує особливості маркетингових даних — відкритість каналів збору поведінкових сигналів та схильність навчальних вибірок до забруднення через клік-фрод.

Наукова новизна. Уперше для класу СПА у галузі цифрового маркетингу запропоновано метод трасування безпекових вимог за схемою «вимога — вектор загрози — архітектурний компонент — механізм контролю — метрика верифікації». На відміну від універсальних стандартів управління ризиками AI

[1, 4], метод враховує специфіку маркетингових даних: їхню поведінкову природу, чутливість до приватності та доступність каналів збору для зловмисних впливів. На відміну від каталогів загроз [2, 3], трасування інтегрується безпосередньо в етап архітектурного проектування.

Формальне подання методу. Введемо скінченні множини: R — безпекових вимог до СПА; C — архітектурних компонентів СПА; T — векторів загроз за NIST AI 100-2 [1], MITRE ATLAS [3], OWASP ML Top 10 [2]; M — механізмів контролю; V — метрик верифікації. Метод трасування визначається як п'ятимісне відношення:

$$\Phi \subseteq R \times C \times T \times M \times V, \quad (1)$$

де кортеж $(r, c, t, m, v) \in \Phi$ задає трасований ланцюжок «вимога r локалізована на компоненті c , протистоїть загрозі t через механізм m із вимірюванням ефективності метрикою v . Властивість повноти трасування формулюється як:

$$\forall r \in R \exists (c, t, m, v): (r, c, t, m, v) \in \Phi, \quad (2)$$

Тобто, кожна вимога має хоча б один ланцюжок до метрики верифікації, що слугує критерієм верифікованості архітектури СПА. Кількісним показником якості трасування виступає коефіцієнт покриття загроз $Cov(T) = \frac{|\pi_T(\Phi)|}{|T|}$, де π_T є проєкцією відношення на множину загроз. Введене формальне подання дозволяє верифікувати архітектуру СПА через перевірку умови (2) та обчислення коефіцієнта покриття на множині референсних загроз.

Запропонований підхід. Сформовано матрицю трасування (табл. 1), що пов'язує безпекові вимоги до СПА з цільовими архітектурними компонентами, релевантними векторами загроз за NIST AI 100-2 E2025 [1], MITRE ATLAS [3] та OWASP ML Top 10 [2], типовими механізмами контролю та метриками верифікації. Табл. 1 подає скінченну реалізацію відношення (1) для базового набору вимог потужністю $|R| = 6$.

Таблиця 1

Матриця трасування безпекових вимог СПА

Вимога	Компонент	Вектор загрози	Механізм контролю	Метрика
Цілісність навчальних даних	Шар прийому даних	Data Poisoning (NIST; ATLAS AML.T0020; OWASP ML02)	Детекція аномалій, санітизація даних	Частка виявлених забруднених записів
Конфіденційність персональних даних	Сховище даних	Membership Inference (NIST; OWASP ML04)	Диференційна приватність, керування доступом	ε-бюджет приватності
Цілісність моделі	Репозиторій моделей	Model Poisoning (OWASP ML10)	Криптографічне підписування артефактів	Повнота перевірки підпису
Контрольованість прогнозів	Сервіс прогнозування	Evasion / Output Integrity (NIST; OWASP ML01, ML09)	Моніторинг дрейфу, валідація входу	Відхилення розподілу прогнозів

Вимога	Компонент	Вектор загрози	Механізм контролю	Метрика
Доступність сервісу прогнозування	Сервіс прогнозування	Availability / Energy-latency Attacks (NIST)	Обмеження частоти, таймаути обчислень	SLA доступності, p95 latency
Авторизованість запитів	Шар інтеграції	Model Extraction (NIST; OWASP ML05)	Автентифікація API, rate limiting	Частка авторизованих запитів

Висновки. Запропонований метод дозволяє розглядати безпеку СПА як архітектурно інтегровану нефункціональну характеристику, що підвищує верифіковність архітектурних рішень. Подальші дослідження передбачають експериментальну валідацію методу на прототипі СПА для задач прогнозування відтоку клієнтів та динамічного ціноутворення.

1. Vassilev A., Oprea A., Fordyce A., Anderson H., Davies X., Hamin M. Adversarial Machine Learning: A Taxonomy and Terminology of Attacks and Mitigations. NIST AI 100-2 E2025. Gaithersburg: NIST, 2025. DOI: 10.6028/NIST.AI.100-2e2025.
2. OWASP Machine Learning Security Top 10. 2023 edition. URL: <https://mltop10.info/> (дата звернення: 06.05.2026).
3. MITRE ATLAS: Adversarial Threat Landscape for Artificial-Intelligence Systems. URL: <https://atlas.mitre.org/> (дата звернення: 06.05.2026).
4. ISO/IEC 23894:2023. Information technology — Artificial intelligence — Guidance on risk management. Geneva: ISO, 2023. 28 p.

Оцінювання рівня інформаційної безпеки державних інформаційних ресурсів

УДК 004.056

Валентина Ящук¹, Діана Рівняк²

Львівський державний університет безпеки життєдіяльності,

¹valentina.lender@gmail.com, ²dianarivnak@gmail.com

У сучасних умовах цифровізації державного управління та зростання інтенсивності кіберзагроз проблема забезпечення належного рівня інформаційної безпеки державних інформаційних ресурсів набуває стратегічного значення. Державні інформаційні ресурси є основою функціонування електронного урядування, а їх компрометація може призвести до значних соціально-економічних і політичних наслідків. Це обумовлює необхідність розроблення ефективних методів оцінювання рівня їх захищеності. Метою роботи є розроблення математичної моделі оцінювання рівня інформаційної безпеки державних інформаційних ресурсів на основі інтегрального підходу з використанням системи індикаторів та вагових коефіцієнтів.

Аналіз існуючих підходів до оцінювання інформаційної безпеки показує, що більшість із них базується на якісних або експертних оцінках, що знижує об'єктивність результатів. Перспективним є застосування кількісних моделей,