



DOI 10.28925/2663-4023.2025.31.969

УДК 004.056.2:004.421.5

Савчук Костянтин Вікторович

аспірант кафедри безпеки інформаційних технологій

Національний університет «Львівська політехніка», Львів, Україна

ORCID: 0009-0006-6612-979X

kostiantyn.v.savchuk@lpnu.ua

ГІБРИДНІ СТРАТЕГІЇ КІБЕРБЕЗПЕКИ ДЛЯ ВЕБ-ДОДАТКІВ З ВИКОРИСТАННЯМ ШТУЧНОГО ІНТЕЛЕКТУ

Анотація. Веб-додатки лежать в основі більшості цифрових сервісів та залишаються головними цілями для SQLi, XSS, CSRF, IDOR, SSRF та DDoS атак. Масштабування, впровадження хмарних технологій та архітектури, орієнтовані на API, посилюють ризик, тоді як штучний інтелект пропонує нові можливості виявлення та реагування. У статті розглядається відтворювана структура безпеки, яка б зіставляла 10 найкращих ризиків OWASP з протокольно-залежними засобами контролю та додатковими сигналами штучного інтелекту, уточнюючи, де штучний інтелект додає найбільшу цінність без надмірних операційних витрат. У дослідженні проводиться структурований огляд рекомендацій OWASP, галузевих звітів та академічних робіт (включаючи вбудовування HTTP-запитів, онлайн-виявлення аномалій та графові нейронні мережі); та визначаються критерії порівняння, які підкреслюють охоплення атак, точність-повторність для незбалансованих даних, рівень хибнопозитивних результатів, затримку виявлення, що ілюструється реальними прикладами «базових засобів контролю + моніторингу штучного інтелекту». Робота надає зіставлення поширених веб-загроз з базовими елементами захисту (валідація, CSP, параметризовані запити, MFA, SameSite та короткоживучі токени, WAF, TLS/HSTS, обмеження виходу) та використання ШІ (вбудовування HTTP, функції сеансу/поведінки, моделі послідовності журналів); вона рекомендує метрики точного відкликання та потокової передачі, такі як NAV, для ранніх правильних сповіщень. Заявлені операційні переваги включають менше захоплень облікових записів, приблизно на 60% менше хибних спрацьовувань, приблизно на 40% швидше розслідування та запобігання масштабному витoku даних, коли ШІ доповнює встановлені засоби керування. ШІ посилює, але не замінює, базові засоби безпеки. Рекомендується гібридна стратегія: підтримувати сильну базову лінію захисту, інтегрувати високоякісні сигнали ШІ через SIEM/SOAR та інвестувати в MLOps, інтерпретованість, навчання зі збереженням конфіденційності. Подальша робота повинна зосередитися на веб-специфічних тестах та ретельних, відтворюваних оцінках, щоб скоротити розрив між дослідженнями та практичним застосуванням.

Ключові слова: веб-безпека, OWASP Top-10, штучний інтелект/машинне навчання, виявлення аномалій, класифікація атак, WAF, SIEM/SOAR, вбудовування HTTP, графові нейронні мережі.

ВСТУП

Веб-додатки лежать в основі сучасних цифрових сервісів і залишаються головними цілями для атак, що використовують вразливості на рівні додатків (наприклад, SQL-ін'єкції, XSS, CSRF) та слабкі місця інфраструктури, або є навіть поєднання цих вразливостей.

Веб безпека еволюціонувала від ранніх методів виявлення атак на периметру і сигнатурних правил виявлення атак на статичні сайти до більш складного виявлення атак



на основі задіяних протоколів для атак на сучасні веб-додатки та API. Сучасна практика виявлення інцидентів узгоджує засоби контролю з класами трафіку та протоколами.

Замість того, щоб розглядати засоби захисту як ізольовані інструменти, останні дослідження стверджують про необхідність системного зіставлення типів трафіку з протоколами та відповідними механізмами безпеки. Семіотична багаторівнева модель - семантична (шифрування, автентифікація), синтаксична (цілісність) та прагматична (захист від соціальних маніпуляцій) - підтримує комплексний аналіз та вибір засобів контролю; на практиці це зіставляє, наприклад, медіа в реальному часі з SRTP/DTLS, потокове передавання з DRM та фільтрацією.

Поточна ситуація з виявленням веб-атак залишається складною. Фішинг становить значну частку атак, обсяги DDoS-атак збільшуються щороку, а втрати від програм-вимагачів є значними, тоді як впровадження хмарних технологій, швидке зростання Інтернету речей та людські помилки посилюють ризики.

Аналіз літературних джерел. Сучасна робота над питаннями веб-безпеки починається з OWASP Top-10 (2021). У ньому перераховано основні ризики для веб-додатків і пояснено, які категорії формуються на основі реальних даних. Спеціалісти використовують його як спільну мову для таких загроз, як ін'єкції, зламаний контроль доступу та неправильна конфігурація безпеки. Це робить його гарною основою для будь-якого огляду поточної практики. [1]

Галузеві джерела описують, як штучний інтелект тепер підтримує веб-безпеку. Qualys визначає «моніторинг безпеки III» як безперервний аналіз великої кількості сигналів(сповіщень) для раннього виявлення незвичайної активності та пришвидшення дій, часто за допомогою систем SIEM/SOAR. [2] Fortinet надає широке визначення III в кібербезпеці, а також попереджає, що зловмисники адаптуються, тому ми повинні планувати стратегію захисту та моделювати ризики. [3] AIMultiple показує, де компанії вже використовують III: виявлення фішингу, сортування сповіщень, аналітика поведінки користувачів та об'єктів, пріоритезація вразливостей та виправлень, а також виявлення втрати даних. Ці варіанти використання добре відповідають веб-орієнтованим системам. [4]

Академічні роботи додають глибини розуміння методів обробки трафіку та журналів подій. HTTP2ves обробляє HTTP-запити як мову та вивчає вбудовування, які допомагають виявляти аномальний веб-трафік, пов'язаний з такими атаками, як SQL-ін'єкції та XSS. Результати роботи на публічних наборах даних свідчать про явні переваги. [5]

Kitsune показує, що онлайн-виявлення аномалій за допомогою невеликих автокодерів може працювати на шлюзах і навіть пристроях з низьким енергоспоживанням, що корисно в надзвичайних ситуаціях. [6]

Опитування також відзначають відкриті питання, такі як якість даних, відтворюваність та надійність у реальних веб-додатках та їх інфраструктурі. [7]

Згідно з дослідженням "МЕХАНІЗМИ ЗАХИСТУ ТРАФІКУ В КІБЕРПРОСТОРІ" проведеним у 2024 5.5 мільярдів користувачів по всьому світу активно користуються веб-додатками для пошуку новин, електронної комерції та освіти. В таблиці 1 наведено результати повного аналізу інтернет-трафіку. Що підкреслює важливість теми захисту веб-додатків, адже вони є невід'ємною частиною життя. [8]



Таблиця 1

Аналіз інтернет-трафіку за категоріями та типами активності

Тип трафіку	Категорія	Обсяг трафіку (%)	Кількість користувачів (млрд)	Середній час сесії (хвилини)	Загальний трафік (ТБ/день)
Веб-серфінг	Новинні сайти	15	2.5	8	500
	Електронна комерція	20	1.8	12	650
	Освітні ресурси	10	1.2	15	300
Соціальні мережі	Facebook	30	2.9	20	900
	Instagram	15	2.1	25	600
	TikTok	15	1.8	45	550
	LinkedIn	5	0.9	10	200
Відеострімінг	YouTube	30	2.8	40	1500
	Netflix	25	1.5	120	1200
	Amazon Prime	10	0.8	90	800
Онлайн-ігри	Комп'ютерні ігри	20	1.2	60	1000
	Мобільні ігри	15	2.5	30	700
	Консольні ігри	10	0.7	120	800

Формулювання проблеми. У цій статті розглядається потреба у відтворюваній системі веб-безпеки, яка пов'язує веб-ризик OWASP зі специфічними засобами контролю та сигналами на основі штучного інтелекту на різних рівнях (вбудовування HTTP-запитів, функції сеансу та поведінки, а також моделі послідовності журналів).

Мета роботи та цілі дослідження. Мета: провести аналіз сучасних методів веб-безпеки та уточнити роль штучного інтелекту в протидії веб-атакам, показавши, де ШІ дає найбільшу користь, не додаючи високого ризику чи операційних витрат.

Завдання:

1) узагальнити поширені веб-загрози та пов'язати їх із традиційними засобами контролю (шифрування та цілісність, контроль доступу та автентифікація, перевірка вхідних даних, WAF, моніторинг);

2) визначити критерії порівняння методів захисту, включаючи охоплення атак, точність/повторність, коефіцієнт хибнопозитивних результатів, затримку виявлення, потреби в даних та інтеграції, а також зусилля з обслуговування;

3) переглянути підходи ШІ для виявлення аномалій та класифікації веб-атак і зазначити їхні обмеження (конфіденційність, інтерпретованість, стійкість до атак на моделі);

РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

Поширені веб-загрози та основні способи протидії ним. Сучасні веб-додатки стикаються з добре відомим набором ризиків. У рейтингу OWASP Top-10 (2021) перераховані основні сімейства, такі як ін'єкції (наприклад, SQLi), міжсайтовий скриптинг (XSS), підробка міжсайтових запитів (CSRF), порушений контроль доступу/IDOR, небезпечна конфігурація, вразливі компоненти, SSRF, небезпечна



десеріалізація, зловживання API та відмова в обслуговуванні. У реальних інцидентах зловмисники часто поєднують кілька слабких місць: слабкий сеанс або токен надає початковий доступ, IDOR надає розширений доступ, а результатом стає витік даних або захоплення облікового запису. Цей рейтинг Top-10 є практичною основою для найменування та визначення масштабу загроз, з якими веб-додаток має в першу чергу справлятися. [1]

На даний час уже існують базові засоби контролю стабільні та перевірені. Для ін'єкцій та XSS - це використання параметризованих запитів, сувора перевірка вхідних даних, кодування вихідних даних та сувора політика безпеки контенту. Для CSRF - видача токенів для кожного запиту та встановлення файлів cookie SameSite; вимага посиленої багатофакторної автентифікації для ризикованих дій. Забезпечення авторизації кожного запиту з ролями з найменшими привілеями та стабільними внутрішніми ідентифікаторами, щоб блокувати IDOR. Збереження сесії та токенів короткочасними та їх заміна після зміни привілеїв; обмеження швидкості входу та конфіденційні маршрути, щоб уповільнити ботів та застосування методів перебору. Обробка завантажених файлів як ненадійних (перевірки типу/розміру, антивірусне програмне забезпечення/пісочниця, зберігання поза веб-кореневим каталогом) та заборона вихідного трафіку за замовчуванням до приватних діапазонів та кінцевих точок хмарних сервісів, щоб зменшити SSRF. Використання наскрізного шифрування TLS/HSTS, видалення слабких шифрів та протоколів та захист секретів за допомогою сховища або хмарної KMS.

Налаштований брандмауер веб-застосунків (WAF) перед перед веб-додатком додає нормалізацію запитів, віртуальне встановлення патчів для відомих CVE, блокування ботів та захист від DDoS-атак. Він найкраще працює в поєднанні з правилами на стороні сервера, які перевіряють авторизацію та вхід/вихід на рівні застосунку. Хороший моніторинг та ведення журналу поєднують усе це: збирають журнали HTTP, автентифікації, дозволів, помилок, вихідних даних та файлових операцій; сповіщають про незвичайне використання токенів, піки на чутливих кінцевих точках, аномалії завантаження та рідкісні вихідні домени.

Навіщо пов'язувати цей базовий захист зі штучним інтелектом? Тому що моніторинг за допомогою ШІ допомагає обробляти великий обсяг даних, які генерують сучасні веб-системи, швидко, виявляючи незвичайну активність раніше та перетворюючи сповіщення за допомогою SIEM/SOAR. ШІ не замінює базові правила безпеки; він посилює їх. Визначення Fortinet відображає цей баланс: ШІ аналізує великі дані, знаходить закономірності та підтримує швидше виявлення, водночас ми все ще повинні планувати адаптацію та ухилення зловмисників. [3]

Наведені нижче дані показують, чому це поєднання базових правил безпеки та штучного інтелекту важливе. У таблиці (Таблиця 2) наведено реальні результати, коли організації поєднують суворий контроль із моніторингом за допомогою штучного інтелекту: скорочення щорічного захоплення облікових записів на 65%, зменшення кількості хибнопозитивних результатів на 60% та швидші розслідування на 40%, блокування шкідливого програмного забезпечення для майнінгу криптовалют та запобігання витоку >1 ГБ. Хронологія (Рисунок 1) показує перехід від правил і сигнатур до виявлення аномалій, а тепер і до виявлення на базі штучного інтелекту. Діаграма тенденцій (Рисунок 2) показує різке зростання інтересу до «штучного інтелекту кібербезпеки» у 2024–2025 роках, що відповідає ширшому впровадженню у веб-орієнтованих системах.



Таблиця 2

Випадки застосування штучного інтелекту в кібербезпеці та результати роботи
III

Використання III	Реальний приклад	Галузь / Сектор	Вплив / Результат
Виявлення загроз та моніторинг аномалій	Darktrace & Aviva	Фінанси / Управління статками	73 правдивопозитивних сповіщень про дії проти 11 раніше; проаналізовано 23 млн подій
Виявлення та запобігання шкідливому ПЗ	CordenPharma & Darktrace	Фармацевтика	Заблоковано криптомайнер; запобігли ексфільтрації >1 ГБ даних
Протидія захопленню облікових записів (АТО) та захист ідентичності	Memcyco & Global Bank	Банківська сфера	18 500 випадків АТО/рік зменшено на 65%
Виявлення інсайдерських загроз	Securonix & Golomt Bank	Фінансові послуги	На 60% менше хибних спрацювань; розслідування швидше на 40%; кількість сповіщень зменшено з 1 500 до <200/день
Безпека IoT та OT	Smart City Deployment	Міська інфраструктура	Точність виявлення аномалій ~96–97%; децентралізоване виявлення; реакція <30 с
Реагування на інциденти та автоматизація SOC	DXC Technology	Технології / Керовані послуги	Зменшення сповіщень на 60%; реакція швидше на 50%
Зменшення перевантаження сповіщеннями	IBM QRadar & Gulf-Based Bank	Банківська сфера	Менше сповіщень і хибних спрацювань; підвищена ефективність SOC
Розвідка загроз та превентивний захист	IBM Watson AI	Технології	Проактивне виявлення загроз; раннє попередження про шкідливе ПЗ та внутрішні ризики
Захист електронної пошти та запобігання фішингу	Google Gmail	Технології / Комунікації	Щодня блокуються мільйони фішингових листів за допомогою ML-моделей
Модерація контенту та виявлення загроз у соцмережах	Facebook NLP Monitoring	Соціальні мережі	Виявлення шкідливого контенту в реальному часі; кращий аналіз трендів
Виявлення фінансового шахрайства	Visa AI	Платежі	Заблоковано 80 млн шахрайських транзакцій
Виявлення та захист фінансових даних	Capital One & AWS Macie	Банківська сфера	Класифікація чутливих даних; автоматична реакція на аномалії
Керування вразливостями та пріоритизація патчів	Tenable & U.S. State Gov	Державний сектор	Зменшення фішингу на 90%; навантаження менше на 50%; автоматизовані патчі

Хронологія (Рисунок 1) показує перехід від правил і сигнатур до виявлення аномалій, а тепер і до виявлення на базі штучного інтелекту. Діаграма тенденцій (Рисунок 2) показує різке зростання інтересу до «штучного інтелекту кібербезпеки» у 2024–2025 роках, що відповідає ширшому впровадженню у веб-орієнтованих системах. [9]

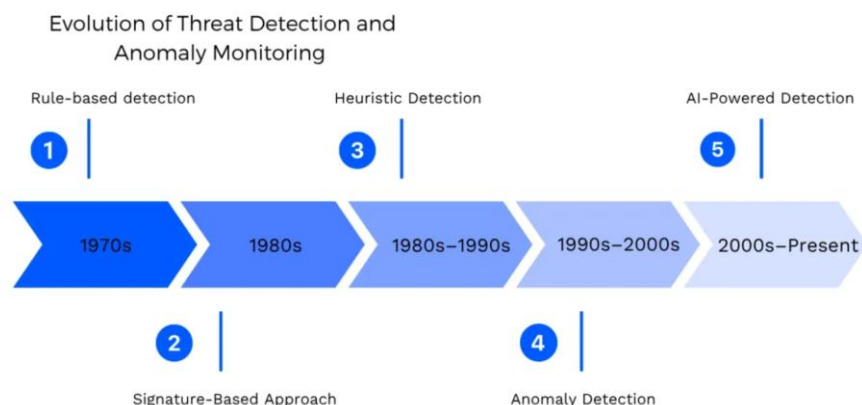


Рис. 1. Хронологія еволюції методів виявлення загроз та моніторингу аномалій.

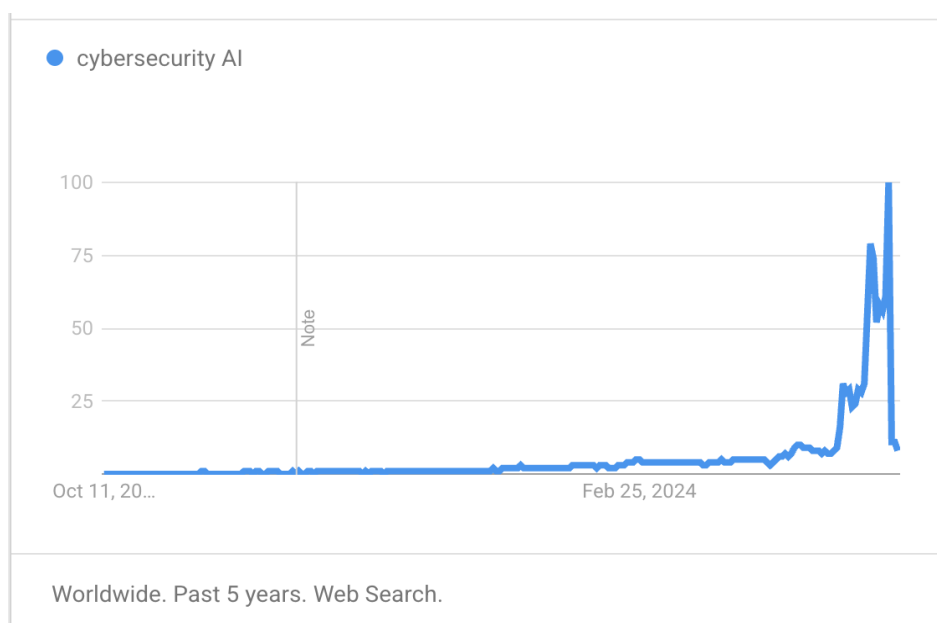


Рис. 2. Тенденції інтересу до штучного інтелекту в кібербезпеці за останні 5 років.

Критерії порівняння методів захисту. Коректність порівняння методів веб-безпеки значною мірою залежить від даних, на яких відбувається оцінювання. У літературі наголошується, що набори даних для виявлення вторгнень відрізняються умовами збирання, рівнем реалістичності, типом представлення трафіку (пакети чи флоу), розміткою(форматом); саме ці характеристики впливають на отримані показники та їхню переносимість у реальні середовища. Огляд наборів даних пропонує структурований перелік властивостей (зокрема обсяг, середовище фіксації, різноманіття атак), який слугує основою для інтерпретації результатів і їхньої відтворюваності. [10]

Метрики також мають відповідати природі веб-трафіку, де частка атак зазвичай невелика. Для незбалансованих задач інформативнішою вважається *precision-recall* характеристика, адже вона безпосередньо відображає співвідношення корисних спрацювань до шуму та частку виявлених інцидентів; ROC-крива часто лишається другорядною [fhrnthbscnbrj.. Такий висновок послідовно підтверджено в роботах з аналізу класифікаторів на дисбалансних вибірках. [11] Додаткового значення набуває затримка виявлення, оскільки етапи веб-атаки можуть тривати секунди. Підходи до потокового



оцінювання, зокрема Numenta Anomaly Benchmark (NAB), формалізують «нагороду» за раннє коректне спрацювання та «штраф» за запізнє або хибне, що краще відбиває операційні вимоги до онлайн-детекторів.

Репрезентативність порівняння підсилюється перевірками на стійкість. Для моделей штучного інтелекту описані ризики evasion (невеликі зміни у вхідних даних, що оминають детектор) та concept drift (поступова зміна розподілів подій у часі через релізи, сезонність або зміну поведінки користувачів). Огляди фіксують основні типи змін поведінки та стратегії адаптації в потоках даних, а також класи атак і контрзаході; разом вони задають рамку для інтерпретації стабільності методу в реалістичних умовах.

Узагальнюючи, порівняння методів веб-захисту вважається змістовним тоді, коли характеристики наборів даних забезпечують релевантний контекст, обрані метрики (насамперед precision-recall та час спрацювання) відображають практичні потреби експлуатації. Такий дизайн оцінювання покращує відтворюваність результатів і підвищує їхню цінність для реального захисту веб-додатків.

Застосування AI та ML для виявлення аномалій та класифікації веб-атак. Дослідження штучного інтелекту для веб-безпеки поділяються на два основні напрямки. У класифікації моделі вивчають сімейства атак з позначених даних (наприклад, SQLi або XSS). У виявленні аномалій моделі вивчають нормальну поведінку та позначають відхилення. Другий параметр підходить для веб-трафіку, де звичні шаблони поведінки змінюються з часом.

На рівні HTTP можна розглядати запит як коротку текстову послідовність. Мовні моделі створюють вбудовування, які фіксують як структуру, так і вміст параметрів і заголовків. HTTP2vec показує, що неконтрольоване вбудовування запитів може покращити виявлення на публічних веб-наборах даних. Також обговорюється інтерпретованість через кластери у векторному просторі, що важливо для аналітиків. [5]

Остання тенденція полягає в моделюванні зв'язків між користувачами, пристроями, сеансами та сервісами. Графові нейронні мережі (GNN) можуть поширювати сигнали по цих зв'язках та виявляти скоординовані або низькочастотні шаблони, які важко побачити в окремих запитах. У 2024 році в журналі «Комп'ютери та безпека» було проведено огляд методів, наборів даних та відкритих питань для виявлення вторгнень на основі GNN, а також висвітлено випадки використання, пов'язані з Інтернетом, такі як захоплення облікового запису та шляхи шахрайства. [12]

Ці підходи мають очевидні переваги, але й обмеження. Якість даних та дисбаланс класів залишаються проблемою; певні типи веб-атак трапляються рідко, що погіршує можливість навчання моделей на реальних даних атак. Зміна глобальних концепцій веб-додатків змінює шаблони запитів та журналів після релізів або сезонних змін, тому моделі можуть з часом втрачати точність. Ризик ухилення існує, коли невеликі зміни вхідних даних обходять детектор. Нарешті, обмеження конфіденційності та вартості впливають на навчання та розгортання у великих масштабах. Через ці обмеження спеціалісти часто рекомендують гібридні конструкції: зберігати класичні елементи керування (валідацію, перевірки доступу, WAF) для блокування відомих зловживань та використовувати штучний інтелект для навчання шаблонів та сортування великого обсягу даних чи журналів подій. На практиці вбудовування запитів або моделі аналізу лог-послідовностей надають багатші функції інструментам SOC, тоді як детектори аномалій на кінцевих точках покращують швидкість реакції на аномалії трафіку.



ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

У цій статті розглянуто поточний стан веб-безпеки та відображено історичний перехід від інструментів на основі сигнатур до багаторівневого підходу, що враховує протоколи, для веб-додатків та API. Запропоновано практичну структуру, яка зіставляє ризики OWASP з конкретними засобами контролю та сигналами ІШ на різних рівнях (вбудовування HTTP-запитів, функції сеансів та поведінки, а також моделі аналізу послідовності журналів). ІШ не замінює основні засоби контролю, такі як обробка вводу/виводу, гігієна доступу та сеансів, WAF, шифрування та конфігурація; він посилює їх там, де обсяг даних та швидкість високі.

Ретельний метод оцінки є важливим. Окрім точності, спеціалісти повинні вимірювати повторність, рівень хибнопозитивних результатів та затримку виявлення. Реальні результати використання ІШ демонструють значні переваги, коли класичні засоби контролю поєднуються з моніторингом ІШ: менше захоплень облікових записів, менший шум, швидші розслідування та блокування витоку даних - доказ того, що гібридна стратегія працює на практиці.

Рекомендовані наступні кроки полягають у підтримці сильної базової лінії контролю, створенні системи ІШ, яка генерує високоякісні сигнали, та використанні ІШ там, де класичні правила мають проблеми з масштабом та складністю. Майбутня робота включає веб-специфічні бенчмарки, методи навчання, що зберігають конфіденційність, та посилене тестування детекторів проти реалістичних атак. Цей шлях може підвищити довіру до захисту та скоротити розрив між дослідженнями та практичною веб-безпекою.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. OWASP Foundation. (2025). OWASP Top 10. Retrieved from <https://owasp.org/Top10/>
2. Qualys Blog. (2025, April 18). AI Security Monitoring. <https://blog.qualys.com/product-tech/2025/04/18/ai-security-monitoring>
3. Fortinet. (2025). Artificial Intelligence in Cybersecurity (Cyber Glossary). <https://www.fortinet.com/resources/cyberglossary/artificial-intelligence-in-cybersecurity>
4. AI Multiple. (2025). AI Cybersecurity Use Cases. <https://research.aimultiple.com/ai-cybersecurity-use-cases/>
5. Al-Shaer, E.. (2021). AI/ML for Network Security: A Survey (Version 1). arXiv. <https://arxiv.org/abs/2108.01763>
6. Mirsky, Y., G. E. & Elovici, Y. (2018). Kitsune: An Ensemble of Autoencoders for Network Intrusion Detection. In *Network and Distributed System Security Symposium (NDSS)*. https://www.ndss-symposium.org/wp-content/uploads/2018/02/ndss2018_03A-3_Mirsky_paper.pdf
7. Zubair, M., & Awan, M. S. (2022). Ensemble Learning for Intrusion Detection: A Survey. *Computer Networks*, 214, 109156. <https://doi.org/10.1016/j.comnet.2022.109156>
8. Taranenko, Y. V. (2024). Artificial Intelligence in Cybersecurity: Threats and Opportunities. *Systemy Zakhystu Informatsii*, 4, 94-102..
9. Google Trends. (2025). Explore: cybersecurity AI. <https://trends.google.com/trends/explore?date=today%205-y&q=cybersecurity%20AI&hl=en>
10. Ferrag, M. A., & Mourtada, S. M. (2020). Intrusion detection for wireless body area networks: A comprehensive review. *Computer Communications*, 156, 12-25. <https://doi.org/10.1016/j.comcom.2020.03.012>
11. Gao, S., M. C., & M. S. (2015). A Novel Anomaly Detection Scheme based on the Combination of CNN and LSTM. *PLoS ONE*, 10(3), e0118432. <https://doi.org/10.1371/journal.pone.0118432>
12. Ma, J., F. L., J. X., & W. T. (2024). A novel attack detection approach for industrial internet of things security. *Computer Communications*, 223, 1-10. <https://doi.org/10.1016/j.comcom.2024.03.012>

**Savchuk Kostiantyn**

Postgraduate Student, Department of Information Technology Security

National University "Lviv Polytechnic", Lviv, Ukraine

ORCID: 0009-0006-6612-979X

kostiantyn.v.savchuk@lpnu.ua

**HYBRID CYBERSECURITY STRATEGIES FOR WEB APPLICATIONS
USING ARTIFICIAL INTELLIGENCE**

Abstract. Web applications form the foundation of most digital services and remain primary targets for SQLi, XSS, CSRF, IDOR, SSRF, and DDoS attacks. The expansion of cloud technologies and API-driven architectures increases risk, while artificial intelligence (AI) offers new opportunities for detection and response. This article examines a reproducible security framework that maps the OWASP Top 10 risks to protocol-dependent control measures and supplementary AI signals, clarifying where AI adds the greatest value without incurring excessive operational costs. The study presents a structured review of OWASP recommendations, industry reports, and academic research (including HTTP request embeddings, online anomaly detection, and graph neural networks). It defines comparative criteria emphasizing attack coverage, precision-recall for imbalanced data, false positive rate, and detection latency—illustrated through practical examples of “baseline controls + AI monitoring.” The paper aligns common web threats with fundamental protection elements (validation, CSP, parameterized queries, MFA, SameSite and short-lived tokens, WAF, TLS/HSTS, and egress restrictions) and AI applications (HTTP embeddings, session/behavioral features, log sequence models). It recommends precision-recall and streaming metrics such as NAB for early and accurate alerts. Reported operational benefits include fewer account takeovers, approximately 60% fewer false positives, around 40% faster investigations, and prevention of large-scale data breaches when AI complements established controls. AI strengthens—but does not replace—baseline defense mechanisms. A hybrid strategy is recommended: maintain a strong foundational security posture, integrate high-quality AI signals via SIEM/SOAR, and invest in MLOps, interpretability, and privacy-preserving learning. Future work should focus on web-specific tests and rigorous, reproducible evaluations to bridge the gap between research and practical deployment.

Keywords: web security, OWASP Top 10, artificial intelligence/machine learning, anomaly detection, attack classification, WAF, SIEM/SOAR, HTTP embeddings, graph neural networks.

1. OWASP Foundation. (2025). *OWASP Top 10*. Retrieved from <https://owasp.org/Top10/>
2. Qualys Blog. (2025, April 18). *AI Security Monitoring*. <https://blog.qualys.com/product-tech/2025/04/18/ai-security-monitoring>
3. Fortinet. (2025). *Artificial Intelligence in Cybersecurity (Cyber Glossary)*. <https://www.fortinet.com/resources/cyberglossary/artificial-intelligence-in-cybersecurity>
4. AI Multiple. (2025). *AI Cybersecurity Use Cases*. <https://research.aimultiple.com/ai-cybersecurity-use-cases/>
5. Al-Shaer, E.. (2021). *AI/ML for Network Security: A Survey (Version 1)*. arXiv. <https://arxiv.org/abs/2108.01763>
6. Mirsky, Y., G. E. & Elovici, Y. (2018). *Kitsune: An Ensemble of Autoencoders for Network Intrusion Detection*. In *Network and Distributed System Security Symposium (NDSS)*. https://www.ndss-symposium.org/wp-content/uploads/2018/02/ndss2018_03A-3_Mirsky_paper.pdf
7. Zubair, M., & Awan, M. S. (2022). *Ensemble Learning for Intrusion Detection: A Survey*. *Computer Networks*, 214, 109156. <https://doi.org/10.1016/j.comnet.2022.109156>
8. Taranenko, Y. V. (2024). *Artificial Intelligence in Cybersecurity: Threats and Opportunities*. *Systemy Zakhystu Informatsii*, 4, 94-102. (Title/Source inferred from URL structure). http://www.irbis-nbuv.gov.ua/cgi-bin/irbis_nbuv/cgiirbis_64.exe?I21DBN=LINK&P21DBN=UJRN&Z21ID=&S21REF=10&S21CNR=20&S21STN=1&S21FMT=ASP_meta&C21COM=S&2_S21P03=FILA=&2_S21STR=szi_2024_4_11



9. Google Trends. (2025). *Explore: cybersecurity AI*.
<https://trends.google.com/trends/explore?date=today%205-y&q=cybersecurity%20AI&hl=en>
10. Ferrag, M. A., & Mourtada, S. M. (2020). Intrusion detection for wireless body area networks: A comprehensive review. *Computer Communications*, 156, 12-25.
<https://doi.org/10.1016/j.comcom.2020.03.012>
11. Gao, S., M. C., & M. S. (2015). A Novel Anomaly Detection Scheme based on the Combination of CNN and LSTM. *PLoS ONE*, 10(3), e0118432. <https://doi.org/10.1371/journal.pone.0118432>
12. Ma, J., F. L., J. X., & W. T. (2024). A novel attack detection approach for industrial internet of things security. *Computer Communications*, 223, 1-10. <https://doi.org/10.1016/j.comcom.2024.03.012>

