



DOI 10.28925/2663-4023.2025.29.821

УДК 004.8:004.056:004.414

Скоринович Богдан Володимирович

аспірант кафедри «Захист інформації»

Національний університет «Львівська політехніка», Львів, Україна

ORCID ID: 0009-0007-8669-9172

bohdan.v.skorynovych@lpnu.ua**Кулик Юрій Андрійович**

аспірант кафедри «Захист інформації»

Національний університет «Львівська політехніка», Львів, Україна

ORCID ID: 0009-0003-9249-2697

yurii.a.kulyk@lpnu.ua**Лакх Юрій Володимирович**

к.ф.-м.н., доцент кафедри «Захист інформації»

Національний університет «Львівська політехніка», Львів, Україна

ORCID ID: 0000-0003-4153-8125

yurii.v.lakh@lpnu.ua

ОЦІНКА ПОТЕНЦІАЛУ ЗАСТОСУВАННЯ МОДЕЛЕЙ ШТУЧНОГО ІНТЕЛЕКТУ ТА МАШИННОГО НАВЧАННЯ ДЛЯ ЗАБЕЗПЕЧЕННЯ БЕЗПЕКИ ХМАРНИХ СЕРЕДОВИЩ І АВТОМАТИЗОВАНИХ СИСТЕМ УПРАВЛІННЯ КОНТЕЙНЕРИЗОВАНИМИ ЗАСТОСУНКАМИ

Анотація. Хмарні обчислення та контейнеризовані середовища стали основою цифрової інфраструктури, забезпечуючи масштабованість і гнучкість для бізнесу. Водночас динамічність цих систем створює нові загрози кібербезпеці, зокрема аномалії в трафіку, атаки типу DDoS, прихований криптомайнінг (cryptojacking) та компрометацію облікових записів. Традиційні засоби захисту, засновані на сигнатурах або ручному налаштуванні, не забезпечують належної ефективності в таких умовах. Метою дослідження є вивчення потенціалу штучного інтелекту (ШІ) та машинного навчання (МН) у вирішенні ключових проблем безпеки хмарних платформ і контейнерних кластерів. Зокрема, розглянуто ефективність AI/ML-моделей у завданнях виявлення аномалій, класифікації атак, виявлення срутпоjacking та DDoS-активності, реалізації desertion-підходів і зниження хибнопозитивних спрацювань. Методологія базується на огляді актуальних наукових публікацій 2023–2025 років. Проведено порівняльний аналіз експериментальних рішень: гібридних моделей (XGBoost, CNN, LSTM) у системах IDS; використання eBPF для аналізу системних викликів у контейнерах; ML-класифікаторів для пріоритизації вразливостей у DevSecOps; платформ для активного захисту з приманками та адаптивним моніторингом (MAPE-K). Результати дослідження демонструють, що сучасні AI-рішення досягають точності виявлення понад 99%, зменшують кількість хибнопозитивних сигналів до $\approx 2\%$ і забезпечують роботу в реальному часі без критичних втрат продуктивності. Особливо ефективними виявилися системи, що поєднують кілька моделей та використовують дані низького рівня (системні виклики, мережеві патерни) для прогнозування загроз. Таким чином, впровадження ШІ у кіберзахист хмарної інфраструктури має вирішальне значення для забезпечення її безперервної та безпечної роботи. У статті також окреслено ключові виклики для подальших досліджень: необхідність пояснюваності моделей (XAI), обмеженість якісних даних для тренування та ризику adversarial-атак. Отримані висновки є корисними для дослідників і практиків у галузі кібербезпеки, які прагнуть впровадити інтелектуальні механізми захисту в динамічних середовищах.

Ключові слова: штучний інтелект; машинне навчання; безпека хмарних середовищ; виявлення аномалій; Kubernetes; DevSecOps.



ВСТУП

Хмарні технології стали невід'ємною складовою сучасних IT-інфраструктур, дозволяючи організаціям динамічно масштабувати ресурси та знижувати витрати. Провідні хмарні провайдери — Amazon Web Services (AWS), Microsoft Azure, Google Cloud тощо — надають широкий спектр сервісів для бізнесу різних галузей [1]. За оцінками Gartner, до 2025 року 85% підприємств реалізуватимуть стратегію cloud-first, тобто більшість їхніх IT-систем розміщуватиметься в хмарі [2]. Разом із перевагами, така міграція породжує критичні проблеми безпеки: динамічність і масштабованість хмари ускладнюють гарантування цілісності даних і безперервності сервісів [3]. Одним із найактуальніших завдань є виявлення аномалій та вторгнень у реальному часі, адже традиційні методи (статистичні пороги, фіксовані правила тощо) не справляються з розмаїттям і швидкістю змін у хмарних середовищах [3].

Особливо вразливими є *контейнеризовані системи* управління на зразок Kubernetes. Kubernetes став основою сучасних хмарних середовищ, забезпечуючи автоматизоване розгортання та масштабування контейнерів, однак його динамічна природа наражає інфраструктуру на складні багатовекторні атаки — підвищення привілеїв, приховане сканування (reconnaissance), розподіленої атаки відмови в обслуговуванні (DDoS) тощо [4]. Контейнерні кластери є *ефемерними*: сервіси швидко змінюються, з'являються й зникають, що утруднює застосування класичних підходів моніторингу [5]. Наприклад, розрізнення справжньої DDoS-атаки від легітимного сплеску навантаження на мікросервіси потребує аналізу тонких поведінкових ознак у телеметрії контейнерів [5].

Крім DDoS, у хмарно-контейнерних середовищах поширилися інші загрози. Однією з найсерйозніших є прихований криптомайнінг (cryptojacking) — приховане використання ресурсів хмарних машин або контейнерів для майнінгу криптовалют без згоди власника [6]. За даними дослідження Google, близько 86% атак у контейнерних середовищах припадає саме на криптомайнінг [6]. Такі атаки непомітно вичерпують обчислювальні ресурси, погіршуючи продуктивність легітимних сервісів та підвищуючи витрати для власників інфраструктури [6]. Інші поширені загрози включають компрометацію облікових записів хмари, несанкціонований доступ до конфіденційних даних та впровадження шкідливого коду в контейнерні образи.

З огляду на ці виклики, зростає інтерес до використання штучного інтелекту для зміцнення безпеки у хмарі. У науковій літературі спостерігається вибух публікацій, присвячених застосуванню методів машинного навчання та глибокого навчання у кібербезпеці хмари. Згідно з недавнім оглядом, станом на кінець 2023 р. було виявлено понад 4050 наукових робіт у цій галузі, що підкреслює її стрімкий розвиток [7]. Основні напрямки включають автоматизоване виявлення аномалій, інтелектуалізацію реагування на інциденти та впровадження новітніх технологій (хмарний SecOps, блокчейн, штучний інтелект речей тощо) для покращення безпеки [7]. Натомість залишаються невирішеними питання масштабованості, приватності даних та пояснюваності моделей ШІ, які потребують уваги дослідників [7].

Таким чином, актуальною науковою задачею є оцінка потенціалу моделей ШІ/МН у вирішенні означених проблем, а саме — наскільки ефективно інтелектуальні алгоритми можуть виявляти та попереджувати загрози в хмарних і контейнеризованих середовищах, і як їх інтегрувати в існуючі процеси забезпечення безпеки.



Постановка проблеми. Забезпечення безпеки у хмарних та контейнерних середовищах стикається з низкою специфічних проблем:

- *Складність виявлення аномалій.* Хмарні системи генерують величезні обсяги різнотипних даних (журнали аудиту, логи мережевих операцій, телеметрія контейнерів). Виявлення у них ознак атак чи відхилень від норми у режимі реального часу є нетривіальним завданням. Статичні правила обмеження запитів та класичні системи виявлення вторгнень (IDS), розроблені для локальних мереж, демонструють незадовільну ефективність у динамічних середовищах (хмарах) — вони або не виявляють нові атаки, або генерують надто багато хибних тривог [3]. Необхідні методи, здатні автоматично навчатися нормальній поведінці сервісів і помічати аномалії без ручного налаштування обмежень.
- *Класифікація різних типів атак.* Сучасні зловмисники використовують широкий спектр технік — від мережевих сканувань і експлойтів до атак на рівні додатків та облікових даних. Система захисту повинна не лише виявити факт аномалії, але й класифікувати тип загрози (наприклад, DDoS, SQL-ін'єкція, крадіжка облікових даних, шкідливий контейнерний образ тощо) для належного реагування. Класичні підходи на основі сигнатур (антивіруси, правила IDS) не розпізнають нові або динамічні атаки. Потрібні інтелектуальні моделі, здатні узагальнювати ознаки атак і виконувати багатокласову класифікацію в режимі, наближеному до реального часу [4].
- *Протидія DDoS-атакам і криптомайнінгу.* Атаки типу відмови в обслуговуванні залишаються однією з найбільших загроз для хмарних сервісів. У контейнерних кластерах DDoS-атаки можуть маскуватися під очікуване різке зростання навантаження (наприклад, автоматичне масштабування подів у Kubernetes). Аналогічно, прихований криптомайнінг важко відрізнити від ресурсомістких легітимних процесів. Проблема полягає в своєчасному виявленні аномального використання ресурсів: перевищення нормальних профілів ЦП/Пам'яті, підозрілих послідовностей системних викликів тощо — без внесення значного перевантаження на моніторингові системи [5]. Важливо мінімізувати вплив засобів моніторингу на продуктивність хмари, зокрема за рахунок технологій на кшталт eBPF (вбудованих в ядро фільтрів пакетів) для відслідковування низькорівневих подій з мінімальним навантаженням [8].
- *Техніки обману і активний захист.* Досвід показує, що досвідчені зловмисники можуть обходити пасивні засоби захисту. В таких умовах актуальним є застосування технологій-приманок: створення фальшивих цілей (honeypots, honeytokens, пасток) для відволікання і виявлення зловмисників. У контейнерних середовищах це, наприклад, розгортання контейнерів-приманок або служб (decoy pods), які імітують вразливі компоненти системи [4]. Проблема полягає у динамічному керуванні такими приманками та їх інтеграції з системами моніторингу: як адаптивно розгортати/згортати обманні середовища у відповідь на ворожу активність і як аналізувати взаємодію зловмисника з ними, щоб отримати розвіддані.
- *Хибнопозитивні спрацювання.* Ефективність системи моніторингу визначається не лише здатністю виявляти атаки, а й кількістю помилкових тривог. В умовах хмари велика частка хибнопозитивних спрацювань призводить до перевантаження аналітиків безпеки та ігнорування



попереджень. Натомість хибнонегативні спрацювання (пропущені атаки) можуть спричинити прямі збитки. Традиційні IDS, налаштовані занадто чутливо, часто страждають від лавини хибних сповіщень. Отже, стоїть завдання знизити рівень хибнопозитивних спрацювань без втрати чутливості до реальних атак. Існує потреба у таких моделях виявлення, які самонавчаються на основі підтверджених інцидентів та корегують свої рішення, мінімізуючи помилки класифікації [3].

- *Інтеграція в конвеєри DevSecOps.* Сьогодні багато організацій переходять до практики DevSecOps, коли безпека впроваджується на всіх етапах циклу розробки та експлуатації програм. Однак додавання численних перевірок безпеки може сповільнити випуск оновлень. Проблема полягає в тому, як інтегрувати розумні механізми безпеки (сканування вразливостей, аналіз конфігурацій, тестування на проникнення) у CI/CD-конвеєри без значних затримок релізів. Необхідні автоматизовані інструменти на базі ШІ, що здатні працювати у фоновому режимі і надавати розробникам своєчасні попередження про ризики [9]. Така інтеграція вимагає крос-взаємодію моделей з DevOps-інструментами та високу швидкість обробки даних пов'язаних з безпекою. Важливою є і довіра до рекомендацій ШІ: команди повинні розуміти результати, щоб оперативно вносити зміни (тут знову постає питання пояснюваності рішень ШІ).

Таким чином, проблема полягає в комплексному забезпеченні проактивного, адаптивного та точного захисту хмарних і контейнерних платформ від широкого спектра загроз, мінімізуючи при цьому помилкові спрацювання і вбудовуючи засоби захисту в швидкісний ритм DevSecOps.

Аналіз останніх досліджень і публікацій. Для розв'язання окреслених проблем дослідники активно застосовують методи штучного інтелекту. Проведемо огляд ключових публікацій за напрямками, що відповідають вищезгаданій підзадачі.

Виявлення аномалій у хмарних середовищах. У низці досліджень підтверджено, що алгоритми машинного навчання перевершують традиційні підходи у виявленні нестандартної поведінки в хмарі [3]. Зокрема, Nwachukwu та ін. (2024) провели ґрунтовний огляд AI-орієнтованих методів виявлення аномалій у хмарних обчисленнях [3]. Вони зазначають, що контрольоване навчання (supervised learning) дозволяє будувати класифікатори для відомих типів аномалій за наявності розмічених даних, тоді як неконтрольовані методи (unsupervised learning) і кластери можуть виявляти раніше невідомі аномалії на основі статистичних відхилень [3]. Набирають популярності й гібридні моделі, що поєднують обидва підходи. За рахунок глибоких нейронних мереж такі системи здатні розпізнавати комплексні багатовимірні шаблони, недоступні для людини, і виявляти приховані атаки, які не мають явних сигнатур [3]. Alzoubi та ін. (2024) відзначають серед актуальних трендів застосування підходів на основі аномалій до захисту хмари, автоматизацію аналізу безпеки та впровадження ШІ для реагування на інциденти [7]. Таким чином, у літературі закріпилась думка, що AI/ML є необхідною складовою для моніторингу безпеки у великих хмарних інфраструктурах, дозволяючи адаптивно «вчитися» поведінці системи та знаходити відхилення в режимі реального часу. Водночас вказується на труднощі: доступність повноцінних даних для тренування (журнали з реальними атаками), масштабованість алгоритмів на рівні хмарних дата-центрів та потреба у поясненні отриманих результатів для довіри з боку аналітиків [3].



Дослідники пропонують вирішувати ці проблеми через федеративне навчання (для обміну знаннями без обміну даними між організаціями), оптимізацію моделей для роботи в розподілених середовищах та використання методів пояснюваного ШІ (Explainable AI, XAI) для інтерпретації виявлених аномалій [3].

Інтелектуальні системи виявлення та класифікації атак. Значну кількість праць присвячено розвитку систем виявлення вторгнень (IDS) з використанням МН/глибокого навчання. Sajid та ін. (2024) запропонували гібридний підхід до IDS, що поєднує алгоритми машинного і глибокого навчання [10]. В їхньому рішенні екстремальний градієнтний бустинг (XGBoost) та згорткові нейронні мережі (CNN) застосовуються для вилучення суттєвих ознак мережевого трафіку, після чого рекурентні нейронні мережі (LSTM) виконують класифікацію атак [10]. Модель протестовано на кількох загальноприйнятих наборах даних (CIC IDS 2017, UNSW-NB15 тощо) як у режимі бінарної класифікації «атака/норма», так і багатокласової класифікації типів атак [10]. Результати демонструють високу точність виявлення відомих атак (понад 99% для деяких датасетів) та низький рівень хибного спрацювання — т.зв. *False Acceptance Rate* незначний [10]. Це свідчить, що поєднання традиційних алгоритмів (дерева рішень, бустинг) із глибокими нейронними мережами дозволяє ефективно обробляти великі масиви мережевих даних і виявляти складні вторгнення. Подібні дослідження підтверджують: багаторівневі моделі (ensemble, stacking) покращують надійність IDS, зменшуючи ймовірність пропуску нових атак та підвищуючи стійкість до шуму в даних [10]. Однак важливо враховувати, що навчання таких моделей потребує значних обчислювальних ресурсів, а також ретельної обробки даних (нормалізації, балансування вибірки тощо) для запобігання зміщенню.

Виявлення DDoS-атак. З розвитком мікросервісних та хмарних архітектур дедалі більшого значення набуває ефективне виявлення розподілених атак типу «відмова в обслуговуванні» (DDoS). Дослідження Pasupathi та ін. (2025) присвячене створенню проактивного механізму виявлення DDoS-атак із застосуванням методів маркування пакетів, поглибленого аналізу мережевого трафіку та алгоритмів машинного навчання (ML). Автори запропонували інтегрований підхід, що базується на комбінації ідентифікації і позначення підозрілих мережевих пакетів (packet marking), аналізу поведінкових патернів трафіку, та ML-класифікації для оперативного виявлення атак. Для збору даних автори використали ретельно побудований набір трафіку, що містив як нормальні, так і згенеровані атакуювальні сценарії з різними типами DDoS-активності. На основі отриманого набору даних було протестовано кілька алгоритмів машинного навчання, серед яких найбільш ефективними виявились ансамблеві алгоритми та методи на основі дерев рішень і нейронних мереж. Автори наголошують на тому, що запропоновані моделі показали надзвичайно високу точність, яка сягала 98%, що дозволило впевнено розрізняти звичайні сценарії інтенсивного трафіку та шкідливу DDoS-активність. Особливістю даного дослідження стало ефективне зниження хибнопозитивних спрацювань за рахунок поєднання маркування пакетів і аналізу поведінки трафіку, що суттєво підвищило надійність виявлення атак у реальному часі. Автори зазначають, що така інтеграція не лише дозволяє швидко ідентифікувати атаки на ранніх стадіях, але й мінімізує кількість помилкових тривог, що є ключовим для підтримки високого рівня доступності сервісів. Результати дослідження свідчать про високий потенціал впровадження запропонованих ML-рішень на мережевому рівні контейнеризованих та хмарних середовищ, де вони можуть працювати як автономні



агенти, що безперервно аналізують мережевий трафік і оперативно повідомляють про виявлені загрози. Це дозволяє реалізувати самоадаптивний механізм захисту, здатний автоматично реагувати на зміни в поведінці мережі, ізолюючи шкідливі потоки або динамічно змінюючи конфігурацію мережі для зменшення наслідків атаки [5].

Виявлення шкідливого криптомайнінгу в контейнерах. Ще одним спеціалізованим напрямом є застосування ІІІ для боротьби з криптомайнінговими атаками у хмарі. Kim та ін. (2025) запропонували оригінальну систему для виявлення криптоджекінгу у контейнерних середовищах на основі трасування системних викликів за допомогою eBPF та подальшого аналізу методами МН [6]. Використання eBPF дозволило знизити накладні витрати: інструмент Tetragon інтегровано в ядро Linux для моніторингу системних викликів контейнерів без значного уповільнення їх роботи [6]. Зібрані послідовності системних викликів і метрики мережевих потоків конвертуються у вектори ознак (із застосуванням технік на основі n-грам, частотних характеристик тощо) та подаються на вхід класифікаторам машинного навчання [6]. Моделі тренувалися розрізняти нормальну активність контейнера та характерні шаблони, притаманні майнерським атакам (наприклад, серії системних викликів, пов'язані з інтенсивними обчисленнями, або мережеві з'єднання до пулів майнінгу). Запропонований підхід продемонстрував надзвичайно високу ефективність — точність класифікації досягла 99.75%, при цьому накладні витрати на моніторинг залишалися помірними [6]. Автори відзначають, що значна частина зловмисної активності у контейнерах наразі припадає на майнінг криптовалют, тому впровадження подібних засобів є нагальною потребою для захисту хмарних платформ. Тим більше, звіт Sysdig 2022 показав, що образи контейнерів із вбудованим майнером є найпоширенішим типом шкідливих образів у публічних репозиторіях — їх знаходять вдвічі частіше, ніж наступний за поширеністю тип загроз [6]. Отже, застосування ІІІ для моніторингу поведінки контейнерів (через системні виклики, ресурси, мережу) дозволяє своєчасно виявляти та знешкоджувати криптомайнери до того, як вони завдадуть значної шкоди.

Техніки обману та проактивна оборона. Оскільки кібератаки стають дедалі витонченішими, науковці звертаються до концепції «active defense» — активного протистояння зловмисникам. Одним із ключових елементів такої стратегії є «cyber deception», тобто дезінформація і відволікання зловмисника. У контексті хмар і Kubernetes це може бути реалізовано як згадано вище — шляхом розгортання спеціальних приманок (фіктивних секретів, фальшивих контейнерів зі штучними вразливостями тощо). Дослідження Ali та ін. (2025) поєднало методи «adaptive deception» з алгоритмами МН у єдиній рамковій системі захисту Kubernetes кластерів [4]. Автори розробили платформу, яка інтегрує компонент KubeDeceive для гнучкого керування мережею honeypot-контейнерів та компонент KServe для розгортання моделей МН, які класифікують тип загрози за спостережуваною активністю [4]. Вся система працює під управлінням циклу MAPE-K (Monitor-Analyze-Plan-Execute with knowledge), що забезпечує безперервну адаптацію стратегії обману на основі даних моніторингу. Тестування цієї концепції показало багатообіцяючі результати: по-перше, мультимодельний класифікатор на базі нейромереж досяг 91% точності у розпізнаванні різних видів атак (включаючи спроби привілейованого доступу, сканування, DDoS) [4]. По-друге, успішність обману — частка випадків, коли атакувальник «клонує» на приманку і взаємодіє із фальшивим ресурсом — склала ~93% [4]. Це свідчить про ефективність налаштування приманок: більшість зловмисників були спрямовані на хибні цілі, що дозволило вивчити їхні дії та убезпечити реальні ресурси. Схема продемонструвала ефективну взаємодію з атакувальником — система не лише пасивно



фіксує факт атаки, а й активно реагує, розгортаючи додаткові пастки, збираючи дані про методи противника і таким чином виграючи час для захисту справжніх сервісів [4]. Дослідження підтверджує: об'єднання ШІ (для аналізу та прийняття рішень) і технологій обману (для протидії зловмисникам) відкриває новий рівень кібероборони, де система може проактивно нейтралізувати загрози ще до того, як вони завдадуть шкоди.

Інтеграція ШІ в DevSecOps та управління вразливостями. Останніми роками з'явилися роботи, що спрямовані на застосування машинного навчання для автоматизації процесів кібербезпеки в рамках DevSecOps. Прикладом є дослідження Raju і Nadella (2024), яке стосується управління вразливостями у хмарному середовищі за допомогою МН [2]. Вони зосередилися на задачі аналізу великих масивів журналів безпеки та звітів про вразливості з метою пріоритезації знайдених проблем і прогнозування появи нових. Запропонована модель навчалась на датасеті з 500 000 записів (логи, знайдені вразливості, дані інцидентів) і досягла ~92% точності у передбаченні наявності вразливості в тому чи іншому компоненті системи [2]. Важливо, що алгоритм не лише виявляє проблеми, але й оцінює їх критичність, завдяки чому можна автоматично визначати пріоритети для виправлення. Запровадження такої інтелектуальної системи дало суттєвий вигрaш у ефективності: частка хибних спрацювань при сповіщенні про вразливості знизилась до 2% [2], а середній час реакції на інцидент скоротився приблизно на 30% [2]. Крім того, автоматизація процесу аналізу дозволила оптимізувати використання ресурсів безпеки (~20% економії трудовитрат) і знизити операційні витрати на ~15% [2]. Ці результати демонструють, що ML-алгоритми можуть успішно інтегруватися у SDLC (Software Development Life Cycle), виконуючи роль «розумного» аналітика, який постійно сканує код, конфігурації та середовище на наявність вразливостей. Більш загально, Fu та ін. (2024) у своєму огляді відзначають 12 ключових завдань безпеки в процесах DevSecOps, для яких вже запропоновано AI-рішення (автоматичний аналіз коду на вразливості, динамічне тестування на проникнення, перевірка конфігурацій хмари, моніторинг контейнерів тощо) [9]. Вони підкреслюють, що штучний інтелект дозволяє поєднати швидкість DevOps із належним рівнем захисту, автоматизувавши багато рутинних перевірок [9]. У результаті компанії можуть частіше розгортати оновлення, не жертвуючи безпекою, оскільки ШІ-агенти невинно слідкують за потенційними ризиками на кожному кроці конвеєра розробки.

В цілому, аналіз сучасної літератури підтверджує високий потенціал моделей ШІ та МН у вирішенні різноманітних завдань забезпечення безпеки хмари і контейнерів. Вже створені прототипи, що демонструють близьку до 100% точність у виявленні певних типів атак (DDoS, cryptojacking) [5], гібридні IDS-системи, здатні класифікувати багато класів загроз з високою достовірністю [10], а також засоби для інтелектуального аналізу вразливостей, інтегровані у робочі процеси DevSecOps [2].

Мета статті. Метою статті є оцінка потенціалу застосування моделей штучного інтелекту і машинного навчання для підвищення рівня безпеки в хмарних середовищах та автоматизованих системах управління контейнеризованими застосунками. Зокрема, дослідження спрямоване на виявлення, наскільки ефективними є сучасні AI/ML-рішення у вирішенні ключових завдань кібербезпеки (виявлення аномалій, класифікація атак, виявлення DDoS і криптомайнінгу, обманні технології, зниження хибнопозитивних спрацювань) та як їх можна інтегрувати в практики DevSecOps для проактивного захисту динамічних інфраструктур. Для досягнення цієї мети проаналізовано наявні наукові підходи і результати, наведені в останніх публікаціях, та проведено узагальнення перспектив і обмежень застосування ШІ/МН у даній сфері.



РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

На основі проведеного аналізу літератури можна зробити декілька узагальнень щодо ефективності та перспектив застосування ШІ/МН для безпеки хмар і контейнерів. Нижче наведено ключові результати та переваги, підтверджені сучасними дослідженнями.

Покращення виявлення аномалій і невідомих атак. AI-моделі суттєво підвищують можливість виявляти відхилення від норми у складних, динамічних середовищах, що не під силу традиційним методам. Застосування алгоритмів неконтрольованого навчання дозволяє фіксувати нові, раніше не бачені типи атак. Так, в роботі [3] підкреслюється здатність МН-моделей виявляти «складні, раніше невідомі» аномалії у хмарі, що забезпечує проактивний захист від zero-day експлойтів. В експериментальних сценаріях моделі типу автоенкодерів успішно виявляли аномальні шаблони доступу до хмарних API, які не могли бути визначені жодною сигнатурою. Це підтверджує, що потенціал ШІ у забезпеченні хмарної безпеки полягає в його адаптивності — здатності навчатися на поведінці системи та помічати будь-які статистично значущі зміни. У реальних умовах це означає своєчасне виявлення навіть тих атак, які ще не відомі спільноті (наприклад, нова техніка обходу автентифікації), що дає адміністраторам системи критично важливий вииграш у часі для реакції на інцидент.

Висока точність і повнота виявлення загроз. Сучасні дослідження демонструють, що правильно налаштовані ML-моделі здатні досягати майже еталонних показників якості виявлення. Для прикладу, гібридна IDS-модель Sajid та інші показала близько 99% точності на кількох датасетах мережових атак [10]. завданні проактивного розпізнавання DDoS-атак дослідження Pasupathi та ін. (2025) продемонструвало 98% точності завдяки інтеграції маркування пакетів, аналізу поведінки мережового трафіку та алгоритмів машинного навчання [5], що дозволило оперативно ідентифікувати шкідливу активність практично без помилок. Для задачі криптомайнінгу — показник 99.75% точності класифікації майнерів проти нормальних контейнерів [6]. Такі результати підтверджують: моделі ШІ здатні виявляти більшість атак з мінімальною кількістю помилок, значно перевершуючи за ефективністю як традиційні системи, так і людський моніторинг (оскільки людина фізично не може обробити такі обсяги даних з подібною швидкістю). Особливо важливо, що висока Recall (повнота) досягається без критичного збільшення False Positive — тобто системи виявляють практично всі атаки і при цьому майже не турбують аналітиків помилковими сповіщеннями. Це знімає давнє протиріччя «чутливість vs точність» у системах безпеки.

Суттєве зниження хибних спрацювань. Інтеграція ML-моделей у процес оповіщення про інциденти дозволяє радикально зменшити кількість хибнопозитивних тривог. Наприклад, впровадження інтелектуального аналізатора журналів безпеки в хмарній інфраструктурі (Raju & Nadella, 2024) дало змогу скоротити частку таких сповіщень до $\approx 2\%$ [2], тоді як до цього більшість спрацювань не підтверджувалась реальними інцидентами. Висока точність моделей, що класифікують події, означає менше непотрібних алертів — отже, інженери безпеки можуть зосередитись на дійсно критичних питаннях. Зниження «шуму» особливо актуальне для великих хмарних платформ, де генеруються тисячі подій за хвилину: AI-фільтрація відсіює понад 90–98% зайвих сигналів. Таким чином, ШІ сприяє підвищенню ефективності SecOps-команд, позбавляючи їх «синдрому перевтоми від оповіщень» та дозволяючи швидше реагувати на реальні атаки. Зафіксовано і зменшення середнього часу реакції на інциденти — в згаданому дослідженні він скоротився на $\sim 30\%$ після автоматизації аналізу логів [2]. Отже, ML не тільки точніше сигналізує про загрози, а й прискорює весь цикл «виявлення-реагування».



Виявлення спеціалізованих атак (DDoS, cryptojacking) у режимі реального часу. Области, що раніше були надто складними для автоматизації (як-то відокремлення DDoS від легітимного трафіку, або пошук майнерів у середовищі з багатьма клієнтами), тепер успішно охоплюються AI-модулями. Наприклад, дослідження Pasupathi та ін. (2025) продемонструвало досягнення $\approx 98\%$ точності при проактивному виявленні DDoS-атак завдяки комбінації маркування пакетів, аналізу поведінкових патернів трафіку та машинного навчання [5]. Це дозволяє системі на льоту ідентифікувати шкідливу активність, ініціювати масштабування захисних ресурсів чи перенаправлення трафіку в хмарній інфраструктурі. Аналогічно, виявлення майнерів [6] дає змогу постійно моніторити тисячі контейнерів і миттєво ізолювати той, в якому стартував несанкціонований майнінг. Такі результати були б неможливі без інтелектуального аналізу: класичні скрипти моніторингу або фіксовані обмеження запитів не впоралися б з розмаїттям нормальних режимів роботи та атакуючих шаблонів. Отже, ІІ розширює можливості захисту на ті види загроз, де раніше автоматичне виявлення було ненадійним або вимагало значних ресурсів. Наприклад, методики з eBPF і ML, запропоновані для cryptojacking [6], показали, що навіть на рівні ядра можна здійснювати інтелектуальний аналіз без суттєвого впливу на продуктивність — це відкриває шлях до створення «розумного» захисного шару безпосередньо в хмарній платформі (вбудованого в гіпервізор або ОС хоста). Практична цінність цього — провайдери та адміністратори отримують інструменти для автоматичної нейтралізації вузькоспеціалізованих атак (атак на доступність, крадіжку ресурсів тощо), які раніше могли залишатися непоміченими протягом тривалого часу.

Адаптивний захист та deception-стратегії. Комбінація ІІ з техніками обману (honeypots, decoys) приводить до появи самоадаптивних систем кібербезпеки. У експериментальному фреймворку Aly та інші (2025) інтегровано цикл автоматичного навчання та планування дій (MAPE-K) [4]: модель МН аналізує дії зловмисника, класифікує тип атаки і приймає рішення, якому саме обманному сценарію слідувати (наприклад, розгорнути додаткові приманки або активувати пастку у вигляді уразливого контейнера). Випробування показали, що така система ефективно заманює зловмисників у штучне середовище (93% успіху) і утримує контроль над ситуацією, не дозволяючи атаці паралізувати реальні сервіси [4]. Це якісно новий рівень захисту: раніше засоби кібербезпеки переважно діяли реактивно (виявили — сповістили — заблокували), тоді як тут реалізовано проактивний підхід (виявили — заманили нападника — зібрали розвіддані — нейтралізували). Роль ІІ критична, адже без нього неможливо оперативно аналізувати поведінку нападника і приймати рішення у реальному часі. У перспективі такі системи адаптивного обману можуть стати невід'ємною частиною хмарних платформ, автоматично розбудовуючи «ландшафт обману» в міру розвитку атаки. Потенціал ІІ в цьому напрямі — створення гнучкої, самонавчальної оборони, що підлаштовується під стратегію супротивника, мінімізуючи його шанси на успіх.

Інтеграція в DevSecOps і підвищення швидкості безпечної розробки. Нарешті, результати досліджень свідчать, що впровадження ML-алгоритмів у процеси DevOps реально допомагає поєднати високу швидкість розробки з вимогами безпеки. Алгоритми на кшталт запропонованого Raju і Nadella класифікатора вразливостей можуть працювати у фоні CI/CD-процесів, аналізуючи кожен новий реліз або конфігурацію на предмет ризиків безпеки. Практичний ефект — скорочення часу на перевірки та реагування. У традиційному процесі релізу патчу безпеки команди могли витратити дні на аналіз журналів та оцінку пріоритетів; з AI-асистентом це займає години або й автоматизується повністю. За рахунок зниження кількості інцидентів, що просочуються в робоче середовище, і швидшого їх усунення, компанії підтримують репутацію та



уникають фінансових втрат від витоків. Більш того, огляд Fu та інші (2024) показав, що компанії-лідери вже застосовують 65 різних метрик і бенчмарків для оцінки AI-рішень у DevSecOps, і за багатьма з них спостерігається позитивна динаміка [9]. Це свідчить про зрілість напрямку: AI у DevSecOps переходить від експериментів до стандартизованих практик. Наприклад, у хмарній платформі AWS нові можливості GuardDuty з 2023 р. включають ML-моделі для аналізу логів Kubernetes — вони автоматично будують поведінковий профіль кластера і сповіщають про аномальні дії (несанкціонований доступ до секретів, запуск контейнерів з нетиповими образами тощо) [1]. Усе це інтегровано у фреймворк AWS Security Hub і може брати участь у автоматичних пайплайнах реагування (AWS Lambda автоматизацій). Таким чином, ШІ стає невід’ємним компонентом інструментарію DevSecOps, забезпечуючи більш розумну і швидку безпеку «за замовчуванням».

Зведені результати не залишають сумнівів: впровадження штучного інтелекту в системи захисту хмарних і контейнеризованих платформ значно підсилює їхню безпеку, роблячи захист більш адаптивним, масштабованим та проактивним. Для наочності наведемо узагальнену таблицю деяких розглянутих рішень та досягнутих ефектів:

Таблиця 1

**Узагальнення прикладів застосування AI/ML
для різних аспектів хмарної і контейнерної безпеки**

Завдання безпеки	AI/ML-рішення (джерело)	Досягнуті результати
Виявлення аномалій у хмарі	Огляд AI-методів для аномалій (Nwachukwu та інші, 2024) [3]	Виявлення складних невідомих аномалій; адаптація до динаміки хмари; підвищена проактивність моніторингу.
Класифікація мережових атак (IDS)	Гібридна модель XGBoost+CNN+LSTM (Sajid та інші, 2024) [10]	>99% точність і високий F1 на декількох наборах даних; низький рівень False Alarm; покращене виявлення багатьох типів атак.
Виявлення DDoS у контейнерних та мережових середовищах	Інтеграція маркування пакетів, аналізу мережевого трафіку та алгоритмів МН (Pasupathi та інші, 2025) [5]	~98% точності; оперативне виявлення DDoS-атак із мінімальною затримкою; суттєве зниження кількості хибнопозитивних спрацювань.
Виявлення прихованого криптомайнінгу	eBPF + ML аналіз системних викликів (Kim та інші, 2025) [6]	~99.75% точність класифікації cryptojacking; мінімальне навантаження на систему; оперативне виявлення прихованого майнінгу в контейнерах.
Адаптивний обман (deception)	ML-класифікатор + KubeDeceive honeypots (Aly та інші, 2025) [4]	~91% точність розпізнавання типу атаки; ~93% успішність залучення атакувальників у пастки; активне стримування загроз.
Управління вразливостями (DevSecOps)	ML-модель пріоритезації вразливостей (Raju & Nadella, 2024) [2]	~92% точність прогнозу вразливостей; зниження False Positives до ~2%; прискорення реакції на інциденти (~30%).
Інтеграція в конвеєри DevOps	Огляд 99 AI-рішень для DevSecOps (Fu та інші, 2024) [9]	Визначено 12 задач DevSecOps, де AI успішно застосовується; автоматизація безпеки без втрати швидкості розробки.

Варто підкреслити, що наведені результати отримані в контрольованих експериментальних умовах або пілотних впровадженнях. Реальний ефект залежатиме від правильності впровадження, але тенденція очевидна: ШІ та МН докорінно покращують показники кібербезпеки у порівнянні з традиційними підходами.



ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Отже, проведений аналіз підтвердив високий потенціал застосування моделей штучного інтелекту та машинного навчання для забезпечення безпеки хмарних середовищ і автоматизованих систем управління контейнеризованими застосунками. Сучасні AI-рішення демонструють здатність суттєво підсилити всі ключові компоненти кіберзахисту: від моніторингу та виявлення загроз до реагування і проактивної оборони. Зокрема, інтелектуальні алгоритми дозволяють:

- Оперативно виявляти аномалії та нові види атак навіть у масштабних і динамічних інфраструктурах, забезпечуючи більш раннє попередження про інциденти і знижуючи ризик прихованих порушень безпеки;
- Підвищувати точність класифікації атак (аж до ~99% у контрольованих експериментах) та охоплювати широкий спектр загроз, що зменшує ймовірність як пропущених атак, так і помилкових тривог;
- Автоматизувати виявлення специфічних атак (DDoS, криптомайнінг тощо) у реальному часі без суттєвого впливу на продуктивність системи, тим самим швидко нейтралізуючи загрози, які раніше могли залишатися непоміченими;
- Реалізовувати адаптивні стратегії захисту (техніки-приманки), де AI керує приманками і активно протидіє зловмисникам, вводячи їх в оману та вигравуючи час для захисту критичних ресурсів;
- Значно скорочувати кількість хибнопозитивних спрацювань, фільтруючи “шум” у даних безпеки і зменшуючи навантаження на аналітиків, що підвищує ефективність роботи команд SecOps;
- Інтегруватися в практики DevSecOps, автоматично забезпечуючи безпеку на всіх етапах життєвого циклу розробки і експлуатації, без уповільнення випуску оновлень, — таким чином, безпека стає органічною частиною процесу розробки, посиленою швидкістю та масштабом ШІ.

З наукової точки зору, дане дослідження підтвердило, що застосування ШІ/МН у сфері хмарної безпеки вже перейшло від теоретичних пропозицій до практичних реалізацій із вражаючими метриками. Це дає підстави прогнозувати, що в найближчі роки такі системи будуть дедалі ширше впроваджуватися в промислових середовищах. Провайдери хмарних послуг (AWS, Microsoft, Google) активно додають ML-можливості до своїх продуктів безпеки (як-от Amazon GuardDuty, Azure Security Center тощо), а корпоративні DevSecOps-стратегії включають AI-інструменти для сканування коду, управління конфігураціями і реагування на інциденти.

Разом з тим, результати роботи виявили і низку обмежень та викликів, що формують напрямки майбутніх досліджень:

- *Проблема даних*: для навчання високоточних моделей потрібні великі масиви якісних даних (журналів атак, нормальної активності тощо). Збір таких датасетів у хмарних середовищах ускладнений міркуваннями конфіденційності та ізолюваністю клієнтських інфраструктур. Перспективним напрямком є розвиток федеративного навчання та методів генерації синтетичних даних, аби моделі могли навчатися на спільному досвіді багатьох організацій без розголошення чутливої інформації.
- *Пояснюваність і довіра*: складні ML-моделі (наприклад, глибокі нейромережі) часто функціонують як «чорні скриньки». Для впровадження в критично важливі системи необхідно забезпечити пояснюваність AI-рішень — тобто надання аналітикам безпеки інтерпретацій, чому саме спрацював детектор, які



ознаки вказують на атаку. Подальші дослідження повинні зосередитися на ХАІ (eXplainable AI) в домені кібербезпеки, щоб підвищити прозорість прийняття рішень і довіру до рекомендацій ШІ.

- *Стійкість до атак на моделі*: поява адверсаріальних атак — методів введення в оману самих ML-моделей (наприклад, шляхом спеціального формування трафіку, який обманює детектор) — ставить нові виклики. Вже показано, що злоумисники можуть знаходити обхідні шляхи для ML-детекторів, змінюючи поведінку атак так, щоб залишатися під порогом виявлення. Перспективним напрямом є розробка стійких алгоритмів і методів захисту моделей (adversarial training, виявлення адверсаріальних впливів), щоб зберегти ефективність ШІ навіть проти цілеспрямованих обходів.
- *Масштабованість і продуктивність*: хоча сучасні дослідження демонструють високу точність, багато моделей ще потребують оптимізації перед розгортанням у хмарних робочих середовищах. Подальші роботи мають зосередитися на масштабуванні AI-рішень — зокрема, на ефективному розподіленому обчисленні для аналізу петабайтів логів, оптимізації роботи моделей в контейнерних середовищах (наприклад, використання легших архітектур, компресія нейронних мереж) та зниженні латентності прийняття рішень, щоб відповідати вимогам захисту у реальному часі.
- *Інтеграція та стандартизація*: впровадження AI в DevSecOps потребує стандартизованих інтерфейсів і протоколів, аби різні інструменти могли обмінюватися даними про загрози і разом реагувати. Майбутні дослідження можуть зосередитися на створенні єдиних платформ SecOps з підтримкою ШІ, які б легко інтегрувалися з існуючими системами (SIEM, SOAR) і хмарними службами. Також важливим є аспект відповідності нормативним вимогам — розробка AI-систем, що враховують вимоги GDPR, стандарти безпеки (наприклад, ISO/IEC 27001) тощо.

Підсумовуючи, штучний інтелект уже сьогодні здатен кардинально підвищити безпеку хмарних та контейнеризованих середовищ, що підтверджено численними науковими роботами. Його подальше впровадження є неминучим і бажаним кроком у розвитку хмарних технологій. По мірі вирішення згаданих викликів — насамперед, забезпечення якості даних, прозорості моделей та стійкості до протидії — можна очікувати появу цілісних автономних систем кібербезпеки, які поєднують силу штучного інтелекту з гнучкістю хмари. Такі системи стануть фундаментом для безпечного функціонування все більш комплексних і масштабних цифрових інфраструктур майбутнього.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Skorynovych, B.V., Lakh, Y.V. (2025). ANALYSIS OF METHODS FOR MONITORING SECURITY STATUS IN A CLOUD ENVIRONMENT. *Modern Information Security*, 61(1). <https://doi.org/10.31673/2409-7292.2025.012256>
2. Raju, S., & Nadella, D. (2024). Enhancing Cloud Vulnerability Management Using Machine Learning: Advancing Data Privacy and Security in Modern Cloud Environments. *International Journal of Computer Trends and Technology*, 72(9), 137–142. <https://doi.org/10.14445/22312803/ijctt-v72i9p121>
3. Chukwuemeka Nwachukwu, Kehinde Durodola-Tunde & Chukwuebuka Akwiwu-Uzoma. (2024). AI-driven anomaly detection in cloud computing environments. *International Journal of Science and Research Archive*, 13(2), 692–710. <https://doi.org/10.30574/ijrsra.2024.13.2.2184>



4. Aly, A., Hamad, A. M., Al-Qutt, M., & Fayez, M. (2025). Real-time multi-class threat detection and adaptive deception in Kubernetes environments. *Scientific Reports*, 15(1). <https://doi.org/10.1038/s41598-025-91606-8>
5. Pasupathi, S., Kumar, R., & Pavithra, L. K. (2025). Proactive DDoS detection: integrating packet marking, traffic analysis, and machine learning for enhanced network security. *Cluster Computing*, 28(3). <https://doi.org/10.1007/s10586-024-04849-x>
6. Kim, R., Ryu, J., Kim, S., Lee, S., & Kim, S. (2025). Detecting cryptojacking containers using eBPF-based security runtime and machine learning. *Electronics*, 14(6), 1208. <https://doi.org/10.3390/electronics14061208>
7. Alzoubi, Y. I., Mishra, A., & Topcu, A. E. (2024). Research trends in deep learning and machine learning for cloud computing security. *Artificial Intelligence Review*, 57(5). <https://doi.org/10.1007/s10462-024-10776-5>
8. Kulyk Y.A., Lakh, Y.V. (2025). SECURITY ANALYSIS OF KUBERNETES NETWORK PLUGINS. *Modern Information Security*, 61(1). <https://doi.org/10.31673/2409-7292.2025.015886>
9. Fu, M., Pasuksmit, J., & Tantithamthavorn, C. (2025). AI for DevSecOps: A Landscape and Future Opportunities. *ACM Transactions on Software Engineering and Methodology*. <https://doi.org/10.1145/3712190>
10. Sajid, M., Malik, K. R., Almogren, A., Malik, T. S., Khan, A. H., Tanveer, J., & Rehman, A. U. (2024). Enhancing intrusion detection: a hybrid machine and deep learning approach. *Journal of Cloud Computing*, 13(1). <https://doi.org/10.1186/s13677-024-00685-x>

**Bohdan Skorynovych**

Postgraduate Student of the "Information Protection" Department

Lviv Polytechnic National University, Lviv, Ukraine

ORCID ID: 0009-0007-8669-9172

bohdan.v.skorynovych@lpnu.ua**Yurii Kulyk**

Postgraduate Student of the "Information Protection" Department

Lviv Polytechnic National University, Lviv, Ukraine

ORCID ID: 0009-0003-9249-2697

yurii.a.kulyk@lpnu.ua**Yuriy Lakh**

PhD, Docent of the "Information Protection" Department

Lviv Polytechnic National University, Lviv, Ukraine

ORCID ID: 0000-0003-4153-8125

yurii.v.lakh@lpnu.ua

ASSESSING THE POTENTIAL OF USING ARTIFICIAL INTELLIGENCE AND MACHINE LEARNING MODELS TO ENSURE THE SECURITY OF CLOUD ENVIRONMENTS AND AUTOMATED MANAGEMENT SYSTEMS FOR CONTAINERIZED APPLICATIONS

Abstract. Cloud computing and containerized environments have become foundational components of modern IT infrastructure, offering scalability and agility. However, their dynamic nature introduces significant security challenges, including anomalies in traffic, DDoS attacks, hidden crypto mining (cryptojacking), and credential compromise. Traditional signature-based security mechanisms often fail to address these rapidly evolving threats effectively. The objective of this study is to assess the potential of artificial intelligence (AI) and machine learning (ML) in enhancing cloud and container security. Specifically, it explores the effectiveness of AI/ML models for anomaly detection, threat classification, cryptojacking, and DDoS identification, deception-based defenses, and false positive reduction. The methodology involves a structured literature review of key scientific publications from 2023 to 2025. Comparative analysis is conducted on experimental solutions, including hybrid models (XGBoost, CNN, LSTM) in intrusion detection systems; eBPF-based syscalls tracing for container behavior profiling; ML classifiers for vulnerability prioritization in DevSecOps; and active defense platforms combining honeypots and adaptive monitoring loops (MAPE-K). Findings indicate that AI-powered security systems achieve detection accuracies above 99%, reduce false positive rates to around 2%, and enable real-time responsiveness without degrading system performance. Notably, systems that integrate multiple models and utilize low-level data (e.g., syscalls, network patterns) exhibit superior threat identification and resilience. In conclusion, integrating AI into cloud security architectures is essential for ensuring continuous and proactive defense in dynamic infrastructures. The paper also outlines future research challenges, such as the need for explainable AI (XAI), limited availability of high-quality training datasets, and vulnerability to adversarial inputs. The insights are relevant for cybersecurity researchers and practitioners seeking to deploy intelligent defense mechanisms in cloud-native ecosystems.

Keywords: artificial intelligence; machine learning; cloud security; anomaly detection; Kubernetes; DevSecOps.

REFERENCES (TRANSLATED AND TRANSLITERATED)

1. Skorynovych, B.V., Lakh, Y.V. (2025). ANALYSIS OF METHODS FOR MONITORING SECURITY STATUS IN A CLOUD ENVIRONMENT. *Modern Information Security*, 61(1). <https://doi.org/10.31673/2409-7292.2025.012256>



2. Raju, S., & Nadella, D. (2024). Enhancing Cloud Vulnerability Management Using Machine Learning: Advancing Data Privacy and Security in Modern Cloud Environments. *International Journal of Computer Trends and Technology*, 72(9), 137–142. <https://doi.org/10.14445/22312803/ijctt-v72i9p121>
3. Chukwuemeka Nwachukwu, Kehinde Durodola-Tunde & Chukwuebuka Akwiwu-Uzoma. (2024). AI-driven anomaly detection in cloud computing environments. *International Journal of Science and Research Archive*, 13(2), 692–710. <https://doi.org/10.30574/ijsra.2024.13.2.2184>
4. Aly, A., Hamad, A. M., Al-Qutt, M., & Fayez, M. (2025). Real-time multi-class threat detection and adaptive deception in Kubernetes environments. *Scientific Reports*, 15(1). <https://doi.org/10.1038/s41598-025-91606-8>
5. Pasupathi, S., Kumar, R., & Pavithra, L. K. (2025). Proactive DDoS detection: integrating packet marking, traffic analysis, and machine learning for enhanced network security. *Cluster Computing*, 28(3). <https://doi.org/10.1007/s10586-024-04849-x>
6. Kim, R., Ryu, J., Kim, S., Lee, S., & Kim, S. (2025). Detecting cryptojacking containers using eBPF-based security runtime and machine learning. *Electronics*, 14(6), 1208. <https://doi.org/10.3390/electronics14061208>
7. Alzoubi, Y. I., Mishra, A., & Topcu, A. E. (2024). Research trends in deep learning and machine learning for cloud computing security. *Artificial Intelligence Review*, 57(5). <https://doi.org/10.1007/s10462-024-10776-5>
8. Kulyk Y.A., Lakh, Y.V. (2025). SECURITY ANALYSIS OF KUBERNETES NETWORK PLUGINS. *Modern Information Security*, 61(1). <https://doi.org/10.31673/2409-7292.2025.015886>
9. Fu, M., Pasuksmit, J., & Tantithamthavorn, C. (2025). AI for DevSecOps: A Landscape and Future Opportunities. *ACM Transactions on Software Engineering and Methodology*. <https://doi.org/10.1145/3712190>
10. Sajid, M., Malik, K. R., Almogren, A., Malik, T. S., Khan, A. H., Tanveer, J., & Rehman, A. U. (2024). Enhancing intrusion detection: a hybrid machine and deep learning approach. *Journal of Cloud Computing*, 13(1). <https://doi.org/10.1186/s13677-024-00685-x>

