



УДК 004.89:004.056.55:004.738.1
DOI: <https://doi.org/10.17721/ISTS.2025.9.42-46>



Сергій ТОЛЮПА, д-р техн. наук, проф.
ORCID ID: 0000-0002-1919-9174
e-mail: serhii.toliupa@knu.ua

Київський національний університет імені Тараса Шевченка, Київ, Україна

Яніна ШЕСТАК, канд. техн. наук
ORCID ID: 0000-0002-1703-0316
e-mail: yanina.shestak@knu.ua

Київський національний університет імені Тараса Шевченка, Київ, Україна

Сергій ДАКОВ, канд. техн. наук
ORCID ID: 0000-0001-9413-3709
e-mail: serhii.dakov@knu.ua

Київський національний університет імені Тараса Шевченка, Київ, Україна

ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ ЗАБЕЗПЕЧЕННЯ БЕЗПЕКИ ЦЕНТРІВ ОБРОБЛЕННЯ ДАНИХ

Вступ. В сучасному світі кіберзагрози для центрів оброблення даних (ЦОД) стали значною проблемою через їхню зростаючу складність та адаптивність. Штучний інтелект (ШІ) здатний значно покращити процеси моніторингу та захисту, забезпечуючи виявлення та реагування на загрози в режимі реального часу. Метою дослідження було оцінити ефективність методів ШІ для підвищення рівня безпеки ЦОД і продемонструвати практичне застосування цих методів.

Методи. Для досягнення цілей дослідження використано два підходи: аналіз поведінкових аномалій і моделювання на основі глибоких нейронних мереж. Дані для навчання та тестування моделей включали інформацію про кіберінциденти за три останні роки (2021–2023), що охоплювали різні типи атак, такі як DDoS, фішингові атаки й атаки нульового дня. Обладнання включало сервери з процесорами Intel Xeon, графічний процесор NVIDIA A100 та програмне забезпечення на базі Python із бібліотеками TensorFlow та Scikit-learn.

Результати. Використання методу аналізу поведінкових аномалій показало точність 89 % у виявленні підозрілих активностей, а глибокі нейронні мережі продемонстрували точність до 92 % у прогнозуванні нових загроз. Середній час реагування на потенційні атаки скоротився з 25 до 8 секунд, що забезпечило своєчасне блокування підозрілих дій. Практичне застосування результатів дослідження включає інтеграцію моделей у системи моніторингу, що дозволяє автоматично виявляти та нейтралізувати загрози, зменшуючи залежність від людського фактора та знижуючи ймовірність помилкових спрацювань.

Висновки. Дослідження підтвердило ефективність ШІ як інструменту для забезпечення високого рівня кібербезпеки ЦОД. ШІ забезпечує швидке та точне виявлення загроз, що дозволяє запобігати їхній реалізації та мінімізувати шкоду. Проте для повного використання потенціалу ШІ необхідно враховувати потребу в якісних даних для навчання, підтримці обчислювальних ресурсів і забезпеченні прозорості алгоритмів. Подальші дослідження мають бути спрямовані на вдосконалення моделей для підвищення їхньої стійкості до маніпуляцій та адаптивності до нових типів загроз.

Ключові слова: штучний інтелект, кібербезпека, центри оброблення даних, аналіз аномалій, нейронні мережі, виявлення загроз.

Вступ

У сучасному цифровому світі центри оброблення даних (ЦОД) є основою багатьох інформаційних систем, забезпечуючи зберігання та оброблення критично важливих даних. Однак разом зі збільшенням їхньої ролі у глобальних процесах значно зросла й кількість кібератак, спрямованих на компрометацію безпеки ЦОД. Згідно з останніми дослідженнями, традиційні методи захисту вже не можуть повною мірою забезпечити необхідний рівень безпеки, оскільки атаки стають дедалі складнішими та більш адаптивними.

Одним із найперспективніших рішень для забезпечення безпеки ЦОД є впровадження штучного інтелекту (ШІ), який дозволяє автоматизувати процеси моніторингу й аналізу даних, виявляючи аномалії та можливі загрози в режимі реального часу. Застосування алгоритмів машинного навчання та глибокого аналізу забезпечує здатність систем до адаптації, що дозволяє виявляти навіть невідомі загрози та запобігати їхній реалізації.

Актуальність цієї теми підтверджується численними дослідженнями та практичними реалізаціями, які демонструють ефективність використання ШІ для захисту інфраструктури ЦОД. Наприклад, у статтях Saxena et al. (2022) і Pant, Anand, & Onthoni (2023) підкреслено важливість застосування багаторівневих підходів до кіберзахисту, що включають інтеграцію ШІ

для виявлення та реагування на загрози. Інші дослідники, такі як Sharma, H., Sharma, G., & Kumar, N. (2024), розглядають методи передачі даних у наступному поколінні мереж із використанням ШІ, що може бути адаптовано для ЦОД для підвищення їхнього захисту.

Мета цієї статті – розглянути сучасні підходи до використання ШІ для забезпечення безпеки ЦОД, аналізуючи існуючі методи виявлення загроз і реагування на них, а також визначити ключові переваги та виклики, пов'язані з впровадженням ШІ.

Огляд літератури. Впровадження штучного інтелекту для забезпечення безпеки центрів оброблення даних є темою численних досліджень, де фахівці аналізують його потенціал та обмеження. Наприклад, робота Saxena et al. (2022) наголошує на важливості використання мультиоб'єктивних підходів для розподілу віртуальних машин у ЦОД, що сприяє зменшенню навантаження на інфраструктуру та підвищенню її стійкості до атак. Автори аргументують, що інтеграція ШІ дозволяє оптимізувати розміщення ресурсів, мінімізуючи ризики вразливостей. Водночас дослідження потребують подальшого аналізу ефективності цих методів у реальних умовах.

Інше значне дослідження Pant, Anand, & Onthoni (2023) розглядає захищені інформаційні системи та підкреслює важливість розроблення адаптивних

© Толіупа Сергій, Шестак Яніна, Даків Сергій, 2025



стратегій для захисту даних. Їхня робота підтверджує, що впровадження ШІ дозволяє не лише виявляти атаки, але й проактивно запобігати їм. Проте вони зазначають, що ефективність таких систем сильно залежить від якості навчальних даних, що може бути потенційною проблемою у випадках використання нерепрезентативних наборів даних.

Sharma, H., Sharma, G., & Kumar, N. (2024) у своїй статті підкреслюють важливість використання ШІ для забезпечення безпечної передачі даних у сучасних гетерогенних мережах (HetNets). Вони стверджують, що інтелектуальні методи значно покращують захист завдяки аналізу поведінкових аномалій. Висновки авторів актуальні для ЦОД, оскільки їхні технології можуть бути адаптовані для аналізу мережної активності. Це дослідження акцентує увагу на важливості інтеграції ШІ з існуючими системами моніторингу для досягнення максимального рівня безпеки.

Цікаво, що Ferencz, & Büki (2022) зосереджуються на аспектах повторного використання захищених даних і відповідних методах безпеки у європейських центрах оброблення даних. Хоча їхній акцент зосереджено більше на правових та організаційних аспектах, вони також визнають потенціал ШІ для підтримки контролю доступу та захисту інформації. Однак автори підкреслюють, що будь-яке впровадження технологій потребує додаткового регулювання для уникнення конфліктів у сфері захисту даних.

Taylor (2021) описує концепцію ультразахищених хмарних сховищ і вказує на використання ШІ для створення "бункерних" моделей безпеки. Хоча їхнє дослідження переважно зосереджене на фізичних заходах безпеки, включення ШІ дозволяє досягти нових рівнів захисту за рахунок виявлення потенційних атак на рівні даних. Це важливо для ЦОД, які шукають шляхи посилення своїх систем безпеки.

Відповідно до Balakrishnan, & Surendran (2020), важливим елементом є стратегія доступу до інформації у віртуальних ЦОД. Це дослідження демонструє, як ШІ може забезпечити безпечний доступ і управління даними в таких середовищах, що є необхідним для запобігання несанкціонованому доступу та порушенню цілісності даних. Автори вказують на можливість використання методів машинного навчання для підвищення ефективності таких систем.

Проведений огляд літератури показує, що використання ШІ в контексті забезпечення безпеки ЦОД має багатообіцяючий потенціал, але також стикається з певними викликами. Це стосується не лише технічних аспектів, таких як оброблення великих обсягів даних і захист алгоритмів від зовнішніх втручань, але й правових і етичних питань. Враховуючи ці аспекти, варто зазначити, що комплексна стратегія із застосуванням ШІ може бути ефективним рішенням для захисту ЦОД, але її впровадження вимагає ретельного планування і узгодження з існуючими стандартами безпеки.

Методи

Для дослідження ефективності застосування штучного інтелекту із забезпечення безпеки центрів оброблення даних використано два основні методи:

Метод аналізу поведінкових аномалій. Цей метод передбачає застосування алгоритмів машинного навчання для моніторингу мережного трафіка та виявлення відхилень від нормальної поведінки. Алгоритми аналізували дані в реальному часі, ідентифікуючи потенційні загрози порівнянням із відомими патернами.

Для цього використовували такі інструменти, як кластеризація та методи аналізу часових рядів.

Моделювання на основі нейронних мереж.

Застосовували глибинні нейронні мережі для навчання системи на великих масивах даних з історією кіберінцидентів. Цю модель використовували для прогнозування загроз і визначення їх імовірності. Результати моделювання порівнювали з реальними атаками для оцінювання точності й ефективності.

Дані для навчання та тестування методів були отримані з наборів даних, що містять інформацію про мережні аномалії та кіберінциденти.

Результати

Для дослідження використано великий масив даних, який містив інформацію про мережні аномалії та кіберінциденти за останні три роки (2021–2023). Дані включали статистику про DDoS-атаки, фішингові атаки й атаки нульового дня, отримані від мережних провайдерів і відкритих репозиторіїв кібербезпеки.

Для оброблення й аналізу даних використовували таке.

Обладнання з високою обчислювальною потужністю:

- Сервер: HP ProLiant DL380 Gen10, оснащений двома процесорами Intel Xeon Gold 6258R, 256 ГБ оперативної пам'яті.
- Сховище: масиви SSD на 8 ТБ для швидкого оброблення великих обсягів даних.
- Графічний процесор: NVIDIA A100 для роботи з глибинними нейронними мережами.

Програмне забезпечення:

- Мови програмування: Python 3.9 із бібліотеками для аналізу даних (Pandas, NumPy) і машинного навчання (TensorFlow, Scikit-learn).
- Інструменти візуалізації: Matplotlib і Seaborn для побудови графіків та аналізу результатів.
- Система управління базами даних: PostgreSQL для зберігання й оброблення даних.
- Операційна система: Ubuntu Server 20.04 LTS.

Під час дослідження здійснено навчання моделей на тестових наборах даних, що включали реальні випадки атак і симуляції потенційних загроз.

Методи аналізу поведінкових аномалій дозволив виявляти підозрілі зміни у мережному трафіку в режимі реального часу. Алгоритми кластеризації натреновано на наборах даних, що містять інформацію про нормальну та підозрілу активність. Використання цього методу дозволило знизити час реагування на загрозу до 10 секунд, що значно скоротило можливість шкоди для ЦОД (табл. 1).

Глибинні нейронні мережі використовували для прогнозування загроз та аналізу мережного трафіка. Модель була навчена на історичних даних про різні типи атак, що дозволило передбачати загрози з високою точністю. Алгоритм демонстрував здатність розпізнавати нові патерни атак, які раніше не були представлені у навчальних наборах, з точністю 92 %. Це стало можливим завдяки глибинному навчанню та використанню алгоритмів розпізнавання аномалій (табл. 2).

Розроблено модель для інтеграції в існуючі системи моніторингу. Під час тестування моделі на реальних мережних даних вдалося ідентифікувати кілька спроб проникнення й атак, які не були виявлені традиційними системами безпеки. Модель працювала в реальному часі та дозволяла автоматично блокувати підозрілі дії, знижуючи навантаження на фахівців із кібербезпеки (табл. 3).



Таблиця 1

Виявлення аномалій на основі кластеризації

Параметр	Значення для нормальної активності	Значення для підозрілої активності
Обсяг трафіка, МБ/с	100–200	500–800
Кількість пакетів	500–1000	>1500
Час пікової активності	14:00–16:00	03:00–05:00

Таблиця 2

Порівняння ефективності нейронних мереж для різних типів загроз

Тип загрози	Виявлення, %	Прогнозування, %	Час оброблення, секунди
DDoS-атака	95	90	7
Фішингова атака	88	85	12
Атака нульового дня	92	89	10

Таблиця 3

Результати тестування моделі на реальних даних

Показник	До інтеграції ШІ	Після інтеграції ШІ
Виявлені загрози, шт.	150	210
Час реагування, секунди	25	8
Відсоток помилкових спрацювань	15	7

Модель, яка використовувала метод аналізу поведінкових аномалій, дозволила ідентифікувати незвичні патерни поведінки, зокрема атаки на рівні мережі та втручання через уразливості програмного

забезпечення. В результаті, автоматизовані дії з нейтралізації, такі як блокування IP-адрес і сегментування мережі, привели до підвищення стійкості системи (табл. 4).

Таблиця 4

Порівняння кількості успішно нейтралізованих загроз

Тип загрози	Кількість атак	Нейтралізація традиційними методами, %	Нейтралізація методами ШІ, %
Втручання на рівні мережі	50	80	95
Вразливості ПЗ	30	70	88

Загалом, результати показали, що застосування ШІ для захисту ЦОД значно підвищує рівень безпеки, забезпечуючи своєчасне виявлення загроз та мінімізацію їх впливу.

Практичне застосування отриманих результатів цього дослідження може бути реалізовано у кількох ключових напрямках:

1. Автоматизовані системи моніторингу та захисту ЦОД. Результати дослідження підтверджують ефективність використання ШІ для виявлення та нейтралізації кіберзагроз у режимі реального часу. Впровадження алгоритмів, розроблених у дослідженні, дозволить операторам ЦОД підвищити рівень безпеки, зменшивши ризики порушення цілісності даних і роботи системи.

2. Інтеграція в існуючі платформи безпеки. Моделі на основі глибоких нейронних мереж можна інтегрувати в існуючі системи кіберзахисту для підвищення ефективності їхньої роботи. Це дозволить знизити кількість помилкових спрацювань і підвищити точність і швидкість виявлення загроз.

3. Адаптація в корпоративних мережах. Виявлені методи можуть бути застосовані для захисту не тільки

великих ЦОД, але й корпоративних мереж, що зберігають конфіденційні дані. Використання ШІ для виявлення поведінкових аномалій і швидкого реагування допоможе зменшити вплив атак на малий і середній бізнес.

4. Розроблення навчальних програм. Одержані результати можуть бути використані для навчання фахівців у галузі кібербезпеки, демонструючи практичні аспекти застосування ШІ для захисту інформаційної інфраструктури.

Дискусія і висновки

Дослідження використання штучного інтелекту для забезпечення безпеки центрів оброблення даних показало значний потенціал для підвищення ефективності моніторингу, виявлення та реагування на кіберзагрози. Практична перевірка застосування методів аналізу поведінкових аномалій і моделювання на основі глибоких нейронних мереж підтвердила високу точність і швидкість реагування на загрози, що особливо важливо для захисту критичних інформаційних систем.



Основними перевагами застосування ШІ є його здатність до безперервного навчання й адаптації, що дозволяє системам не лише виявляти відомі атаки, але й прогнозувати нові. Під час експериментів встановлено, що глибинні нейронні мережі можуть виявляти атаки типу нульового дня, використовуючи аналіз поведінкових патернів, що недоступно для традиційних систем захисту. Крім того, застосування ШІ значно скоротило час реакції на загрози, що дозволило автоматизувати процеси нейтралізації атак і зменшити навантаження на фахівців із кібербезпеки.

Проте впровадження ШІ у системи безпеки ЦОД стикається з низкою викликів. До них включено високу залежність від якості навчальних даних, які повинні бути репрезентативними й оновлюватися, щоб забезпечити належну роботу алгоритмів. Крім того, процес інтеграції ШІ у вже існуючі системи безпеки потребує значних ресурсів і належного планування. Іншим важливим аспектом є забезпечення прозорості й інтерпретованості алгоритмів ШІ, щоб уникнути помилкових спрацювань і підвищити довіру до їхньої роботи.

Практичне значення дослідження полягає у можливості інтеграції розроблених методів у сучасні системи кібербезпеки. Це дозволить підвищити рівень захищеності ЦОД завдяки автоматизації процесів виявлення та реагування на загрози. Отримані результати можуть бути використані для адаптації алгоритмів у корпоративних мережах і розроблення нових навчальних програм для фахівців із кібербезпеки.

Дослідження підтвердило, що застосування ШІ є перспективним напрямом для покращення захисту інформаційних систем від сучасних кіберзагроз. Його інтеграція дозволяє не лише оперативніше реагувати на загрози, але й передбачати їх, що забезпечує високий рівень безпеки та надійності роботи ЦОД. Проте для ефективного використання необхідно враховувати технічні вимоги та постійно вдосконалювати алгоритми відповідно до змін у загрозах та інфраструктурі.

Застосування ШІ у кібербезпеці відкриває нові горизонти для досліджень, зокрема й у напрямі інтеграції з іншими системами безпеки та розроблення алгоритмів, стійких до зовнішніх маніпуляцій. Подальші дослідження можуть бути спрямовані на створення більш адаптивних і захищених від маніпуляцій моделей, що дозволить забезпечити ще вищий рівень захисту для ЦОД та інших критичних інформаційних систем.

Внесок авторів: Сергій Толюпа – концептуалізація; методологія; Яніна Шестак – аналіз джерел, підготування огляду літератури або теоретичних засад дослідження; Сергій Даков – збір емпіричних даних та їхня валідація; емпіричне дослідження.

Джерела фінансування. Це дослідження не отримало жодного гранту від фінансової установи в державному, комерційному або некомерційному секторах.

Список використаних джерел

- Rathod, T., Jadav, N. K., Tanwar, S., Polkowski, Z., Yamsani, N., Sharma, R., Fayed A., & Gafar, A. (2023). AI and Blockchain-Based Secure Data Dissemination Architecture for IoT-Enabled Critical Infrastructure. *Sensors*, 23(21), 8928. <https://doi.org/10.3390/s23218928>
- Sharma, H., Sharma, G., & Kumar, N. (2024). AI-assisted secure data transmission techniques for next-generation HetNets: A review. *Computer Communications*. Elsevier B.V., 215, 74–90. <https://doi.org/10.1016/j.comcom.2023.12.015>
- Ferencz, B., & Büki, B. (2022). Three Ways of Secure Data Reusability in Europe: German Research Data Centres, Finnish Findata and the French Secure Access Data Centre. *ELTE Law Journal*, 1, 81–108. <https://doi.org/10.54148/ELTELJ.2022.1.81>
- Taylor, A. R. E. (2021). Future-proof: bunkered data centres and the selling of ultra-secure cloud storage. *Journal of the Royal Anthropological Institute*, 27, 76–94. <https://doi.org/10.1111/1467-9655.13481>
- Saxena, D., Gupta, I., Kumar, J., Singh, A. K., & Wen, X. (2022). A Secure and Multiobjective Virtual Machine Placement Framework for Cloud Data Center. *IEEE Systems Journal*, 16(2), 3163–3174. <https://doi.org/10.1109/JSYST.2021.3092521>
- Balakrishnan, S., & Surendran, D. (2020). Secure information access strategy for a virtual data centre. *Computer Systems Science and Engineering*, 35(5), 357–366. <https://doi.org/10.32604/CSSE.2020.35.357>
- Pant, P., Anand, K., & Onthoni, D. D. (2023). Secure Information and Data Centres: An Exploratory Study. In *Studies in Computational Intelligence* (Vol. 1065, pp. 61–88). Springer Science and Business Media Deutschland GmbH. https://doi.org/10.1007/978-981-19-6290-5_4

References

- Rathod, T., Jadav, N. K., Tanwar, S., Polkowski, Z., Yamsani, N., Sharma, R., Fayed A., & Gafar, A. (2023). AI and Blockchain-Based Secure Data Dissemination Architecture for IoT-Enabled Critical Infrastructure. *Sensors*, 23(21), 8928. <https://doi.org/10.3390/s23218928>
- Sharma, H., Sharma, G., & Kumar, N. (2024). AI-assisted secure data transmission techniques for next-generation HetNets: A review. *Computer Communications*. Elsevier B.V., 215, 74–90. <https://doi.org/10.1016/j.comcom.2023.12.015>
- Ferencz, B., & Büki, B. (2022). Three Ways of Secure Data Reusability in Europe: German Research Data Centres, Finnish Findata and the French Secure Access Data Centre. *ELTE Law Journal*, 1, 81–108. <https://doi.org/10.54148/ELTELJ.2022.1.81>
- Taylor, A. R. E. (2021). Future-proof: bunkered data centres and the selling of ultra-secure cloud storage. *Journal of the Royal Anthropological Institute*, 27, 76–94. <https://doi.org/10.1111/1467-9655.13481>
- Saxena, D., Gupta, I., Kumar, J., Singh, A. K., & Wen, X. (2022). A Secure and Multiobjective Virtual Machine Placement Framework for Cloud Data Center. *IEEE Systems Journal*, 16(2), 3163–3174. <https://doi.org/10.1109/JSYST.2021.3092521>
- Balakrishnan, S., & Surendran, D. (2020). Secure information access strategy for a virtual data centre. *Computer Systems Science and Engineering*, 35(5), 357–366. <https://doi.org/10.32604/CSSE.2020.35.357>
- Pant, P., Anand, K., & Onthoni, D. D. (2023). Secure Information and Data Centres: An Exploratory Study. In *Studies in Computational Intelligence* (Vol. 1065, pp. 61–88). Springer Science and Business Media Deutschland GmbH. https://doi.org/10.1007/978-981-19-6290-5_4

Отримано редакцією журналу / Received: 20.05.25

Прорецензовано / Revised: 27.05.25

Схвалено до друку / Accepted: 30.06.25



Serhii TOLIUPA, DSc (Engin.), Prof.
ORCID ID: 0000-0002-1919-9174
e-mail: serhii.toliupa@knu.ua
Taras Shevchenko National University of Kyiv, Kyiv, Ukraine

Yanina SHESTAK, PhD (Engin.)
ORCID ID: 0000-0002-1703-0316
e-mail: yanina.shestak@knu.ua
Taras Shevchenko National University of Kyiv, Kyiv, Ukraine

Serhii DAKOV, PhD (Engin.)
ORCID ID: 0000-0001-9413-3709
e-mail: serhii.dakov@knu.ua
Taras Shevchenko National University of Kyiv, Kyiv, Ukraine

THE USE OF ARTIFICIAL INTELLIGENCE FOR ENSURING THE SECURITY OF DATA CENTERS

Background. *In today's world, cyber threats to data centers (DCs) have become a significant concern due to their growing complexity and adaptability. Artificial Intelligence (AI) can greatly enhance monitoring and security processes, ensuring real-time threat detection and response. The aim of this study was to evaluate the effectiveness of AI methods for improving DC security and to demonstrate their practical applications.*

Methods. *In today's world, cyber threats to data centers (DCs) have become a significant concern due to their growing complexity and adaptability. Artificial Intelligence (AI) can greatly enhance monitoring and security processes, ensuring real-time threat detection and response. The aim of this study was to evaluate the effectiveness of AI methods for improving DC security and to demonstrate their practical applications.*

Results. *The use of the behavioral anomaly analysis method achieved an accuracy of 89% in detecting suspicious activities, while deep neural networks demonstrated up to 92% accuracy in predicting new threats. The average response time to potential attacks was reduced from 25 to 8 seconds, enabling timely blocking of suspicious actions. Practical applications include integrating these models into monitoring systems, allowing automatic threat detection and mitigation, reducing reliance on human intervention, and minimizing false positives.*

Conclusions. *The study confirmed the effectiveness of AI as a tool for ensuring high levels of DC cybersecurity. AI enables quick and precise threat detection, preventing their realization and minimizing potential damage. However, to fully harness AI's potential, it is essential to consider the need for high-quality training data, computational resources, and algorithm transparency. Future research should focus on refining models to enhance their resistance to manipulation and adaptability to new types of threats.*

Keywords: *artificial intelligence, cybersecurity, data centers, anomaly analysis, neural networks, threat detection.*

Автори заявляють про відсутність конфлікту інтересів. Спонсори не брали участі в розробленні дослідження; у зборі, аналізі чи інтерпретації даних; у написанні рукопису; в рішенні про публікацію результатів.

The authors declare no conflicts of interest. The funders had no role in the design of the study; in the collection, analyses or interpretation of data; in the writing of the manuscript; in the decision to publish the results.