

Список використаних джерел:

1. NVIDIA. NVIDIA ACE for Games: Bringing Digital Humans to Life. 2024. URL: <https://developer.nvidia.com/ace-for-games> (дата звернення: 19.03.2026).
2. Unity. Generative AI in Game Development: How It's Changing the Industry . – 2024. – URL: <https://blog.unity.com/engine-platform/generative-ai-in-game-development> (дата звернення: 19.03.2026).
3. Inworld AI. The Future of NPCs: How Inworld AI is Changing Gaming. 2025. URL: <https://inworld.ai/blog/the-future-of-npcs> (дата звернення: 19.03.2026).
4. Game Developer. AI-Powered Storytelling: The Next Frontier in Game Narrative. 2024. URL: <https://www.gamedeveloper.com/design/ai-powered-storytelling> (дата звернення: 19.03.2026).

УДК 004.8:004.49

ВПЛИВ ШІ НА ЕВОЛЮЦІЮ ШКІДЛИВОГО ПРОГРАМНОГО ЗАБЕСПЕЧЕННЯ ТА АНТИВІРУСІВ

*Федоров Є.О.
egorfedorov0306@gmail.com
Черкаський державний фаховий бізнес-коледж
Немченко В.Ю.
м. Черкаси, Україна*

Стрімке впровадження технологій штучного інтелекту та машинного навчання кардинально змінило глобальний ландшафт кібербезпеки. Сучасні інформаційні системи стикаються з викликами, які вимагають принципово нових підходів до захисту даних. Якщо раніше протистояння між зловмисниками та антивірусними компаніями базувалося на статичному аналізі файлів та сигнатурах, то сьогодні ця боротьба перетворилася на високотехнологічну алгоритмічну війну. Штучний інтелект став одночасно і найпотужнішим щитом, і найнебезпечнішою зброєю.

Використання машинного навчання відкрило нові вектори атак, головним з яких є створення інтелектуального поліморфного коду. Сучасне шкідливе ПЗ

здатне динамічно переписувати власну структуру безпосередньо під час виконання.

Зберігаючи свій первинний деструктивний функціонал, такий вірус після кожного нового зараження генерує унікальну цифрову сигнатуру. Це робить кожен його екземпляр абсолютно унікальним, зводячи нанівець ефективність класичних антивірусних сканерів.

Крім того, кіберзлочинці активно застосовують ШІ для автоматизації пошуку вразливостей нульового дня. Спеціально навчені нейромережі здатні за лічені години проаналізувати мільйони рядків коду операційних систем, знаходячи приховані архітектурні хиби, на пошук яких у людей пішли б місяці.

Ще одним критичним досягненням "темного" ШІ стало "розумне ухилення" (AI-driven Evasion). Інтелектуальні віруси навчилися самостійно аналізувати середовище свого виконання. Якщо такий алгоритм розпізнає ознаки віртуальної пісочниці (Sandbox) або фіксує активний моніторинг, він миттєво призупиняє шкідливу активність. Вірус починає ідеально імітувати поведінку звичайного системного процесу, доки не переконається у власній безпеці.

Розглядаючи еволюцію кіберзагроз, неможливо оминати ще один інноваційний та вкрай небезпечний вектор – цілеспрямовані атаки на самі моделі машинного навчання, які лежать в основі сучасних антивірусів. Цей науковий напрямок отримав назву змагального машинного навчання (Adversarial Machine Learning). Зловмисники усвідомили: замість того, щоб намагатися непомітно обійти "розумний" захист, набагато ефективніше змусити помилятися на фундаментальному математичному рівні.

Однією з тактик є "отруєння даних". Її суть полягає в тому, що хакери тривалий час і дуже обережно "згодовують" системі захисту EDR (Endpoint Detection and Response) спеціально модифіковані, зовні безпечні зразки активності. Головна мета цього процесу – змусити нейромережу змінити своє уявлення про "нормальну поведінку" системи. Якщо атакувальнику вдасться поступово змістити базову лінію норми у математичній моделі, штучний

інтелект почне ігнорувати реальні шкідливі процеси, помилково класифікуючи їх як легітимні дії користувача.

Іншим передовим методом є використання змагального шуму (Adversarial Perturbations). Зловмисники вносять у код шкідливої програми мінімальні, функціонально незначущі зміни – своєрідний невидимий "цифровий шум", який абсолютно не впливає на роботу вірусу. Проте для математичної моделі антивірусу цей шум кардинально спотворює сприйняття об'єкта. Через такі маніпуляції нейромережа стикається з оптичною ілюзією на рівні коду і може впевнено класифікувати небезпечний руткіт як безпечний системний драйвер або текстовий документ. Це створює парадоксальну ситуацію в сучасній кібербезпеці: чим складніший і "чутливіший" штучний інтелект захисної системи, тим вразливішим він може виявитися до таких специфічних математичних маніпуляцій.

У відповідь на ці безпрецедентні виклики індустрія безпеки здійснила перехід від реактивного до проактивного захисту. Сучасні системи класу EDR більше не покладаються виключно на бази відомих загроз. Ядром новітніх антивірусів стали ансамблі моделей машинного навчання, які здійснюють глибокий поведінковий аналіз у режимі реального часу. Такі системи безперервно збирають телеметрію, відстежуючи звернення програм до пам'яті та мережі.

На основі зібраних даних штучний інтелект формує унікальний профіль нормальної поведінки для кожної машини. Будь-яке відхилення від цієї норми (наприклад, спроба масового шифрування файлів) миттєво класифікується як аномалія. Це дозволяє системі захисту приймати автономні рішення. Вона може автоматично заблокувати підозрілий процес та ізолювати пристрій від загальної мережі ще до того, як комп'ютеру буде завдано фактичної шкоди.

Окремою і надзвичайно важливою тенденцією є взаємне використання генеративно-змагальних мереж (GAN) обома сторонами конфлікту. Зловмисники застосовують GAN для генерації мільйонів варіацій шкідливого коду, тренуючи їх обходити комерційні антивіруси. Своєю чергою, розробники

систем захисту використовують ці ж технології у своїх лабораторіях для створення "штучних вірусів". На цих згенерованих зразках вони превентивно тренують свої моделі виявлення, навчаючи їх розпізнавати загрози майбутнього.

Таким чином штучний інтелект назавжди змінив фундаментальні правила інформаційної безпеки. Відмова від класичних парадигм на користь машинного навчання стала єдиним шляхом виживання в сучасному цифровому світі.

Ефективність сучасного захисту сьогодні вимірюється не розміром вірусних баз, а обчислювальною потужністю його математичних моделей. У найближчому майбутньому кіберпростір остаточно перетвориться на арену автономних алгоритмічних битв, де участь людини буде зведена до контролю високорівневих політик безпеки.

Але незважаючи на революційний вплив штучного інтелекту, його впровадження не є абсолютною панацеєю від усіх сучасних кіберзагроз.

Список використаних джерел:

1. Звіт про глобальні кіберзагрози (Global Threat Report 2024). CrowdStrike. URL: <https://www.crowdstrike.com/global-threat-report/> (Дата звернення: 07.04.2026)
2. Звіт про цифровий захист (Microsoft Digital Defense Report 2024). Microsoft. URL: <https://www.microsoft.com/en-us/security/security-insider/threat-landscape/microsoft-digital-defense-report-2024> (Дата звернення: 07.04.2026)
3. AI в кібербезпеці: як штучний інтелект захищає і атакує цифровий світ. *Maxnet Blog*. URL: <https://maxnet.ua/blog/ai-v-kiberbezpeci-yak-shtuchnij-intelekt-zahishaye-i-atakuye-cifrovij-svit/> (Дата звернення: 07.04.2026)
4. ENISA. Artificial Intelligence Cybersecurity Challenges. European Union Agency for Cybersecurity, 2020. URL: <https://www.enisa.europa.eu/publications/artificial-intelligence-cybersecurity-challenges> (Дата звернення: 07.04.2026)
5. Gallagher S., Posey B. AI and the Future of Cybersecurity: Defensive and Offensive Applications. O'Reilly Media, 2022. 210 p.

6. IBM Security. *Cost of a Data Breach Report 2023* URL:
<https://www.ibm.com/security/data-breach> (Дата звернення: 07.04.2026)

УДК 004.738.5

ПРОГНОЗУВАННЯ ПСИХОЛОГІЧНОГО СТАНУ КОРИСТУВАЧА НА ОСНОВІ ЦИФРОВИХ СЛІДІВ

Рибалко А. Д.
anna10rybalko@gmail.com

Черкаський державний фаховий бізнес-коледж
Люта М. В.
м. Черкаси, Україна

У цифрову епоху кожна дія людини в інформаційному просторі залишає «цифрові сліди» – сукупність даних про онлайн-активність, які відкривають нові можливості для аналізу ментального здоров'я. Прогнозування психологічного стану на основі цих даних є актуальним інструментом для превентивної психології та персоналізованої підтримки користувача [1].

Цифрові сліди включають широкий спектр інформації: тексти публікацій та коментарів у соціальних мережах, історію пошукових запитів, вподобання контенту, часові інтервали активності та навіть метадані фотографій [2]. Аналіз семантики повідомлень дозволяє виявляти ранні ознаки депресивних станів, тривожності або професійного вигорання.

Технологічною основою такого прогнозування є моделі обробки природної мови (NLP) та нейронні мережі [3]. Алгоритми здатні обробляти неструктуровані дані, розпізнавати складні патерни поведінки та зіставляти їх із психологічними маркерами.

Це дозволяє створювати динамічний профіль стану користувача, який оновлюється в режимі реального часу, на відміну від традиційних опитувальників [4].

Наприклад, зміна кола інтересів, часте відвідування ресурсів із песимістичним контентом або зміна нічного режиму активності можуть свідчити про порушення психологічної рівноваги [5]. Водночас аналіз цифрових слідів