

ШІ-АГЕНТИ В КІБЕРПРОСТОРІ: ЗАСТОСУВАННЯ В АТАКАХ ТА СИСТЕМАХ ЗАХИСТУ ІНФОРМАЦІЇ

Мотайленко О.О.
omotaylenko@gmail.com
Черкаський державний фаховий бізнес-коледж
Захарова М.В.
м. Черкаси, Україна

У сучасному цифровому суспільстві кіберпростір став одним із ключових середовищ функціонування інформаційних систем, цифрових сервісів та комунікацій. Швидкий розвиток інформаційних технологій, хмарних обчислень, Інтернету речей (IoT) та глобальних мереж сприяє значному зростанню обсягів даних і кількості підключених пристроїв, а разом із цим збільшується і кількість кіберзагроз, що створює необхідність у використанні нових технологій для забезпечення ефективного захисту інформації. Одним із перспективних напрямів розвитку кібербезпеки є використання систем штучного інтелекту та автономних ШІ-агентів, здатних аналізувати інформацію, приймати рішення та виконувати складні завдання без постійного контролю людини [1–3].

Метою даної роботи є дослідження ролі та можливостей ШІ-агентів у сучасному кіберпросторі, аналіз особливостей їх застосування як у процесі здійснення кібератак, так і в системах захисту інформації. У роботі передбачено розглянути сутність ШІ-агентів, основні напрями їх використання в кібербезпеці, переваги та ризики впровадження, а також перспективи розвитку цих технологій у сфері кіберзахисту з урахуванням технічних, етичних і правових аспектів.

ШІ-агенти - це програмні системи, які можуть автономно взаємодіяти з цифровим середовищем, збирати дані, аналізувати їх та виконувати певні дії відповідно до заданих алгоритмів або моделей машинного навчання. Завдяки використанню методів штучного інтелекту, таких як машинне навчання, глибоке навчання та аналіз великих даних, ці системи здатні обробляти величезні масиви інформації та знаходити закономірності, які складно виявити традиційними методами. Як показано на рис. 1, у сфері кібербезпеки ШІ-агенти можуть

виконувати функції моніторингу мережевого трафіку, аналізу поведінки користувачів, виявлення аномалій та автоматичного реагування на інциденти безпеки [2–4].

З іншого боку, ШІ-агенти активно використовуються для підвищення рівня кіберзахисту. Сучасні системи інформаційної безпеки дедалі частіше застосовують алгоритми машинного навчання для аналізу мережевого трафіку та виявлення підозрілої активності. Наприклад, системи виявлення вторгнень (IDS) та системи запобігання вторгненням (IPS) використовують моделі штучного інтелекту для аналізу поведінкових моделей користувачів і пристроїв у мережі. Якщо система виявляє відхилення від нормальної поведінки, вона може автоматично сигналізувати про потенційну атаку або навіть заблокувати підозрілу активність [1–4].



Рисунок 1 – Схема застосування ШІ-агентів у кіберпросторі

Особливо важливим напрямом є створення автономних систем кіберзахисту, здатних самостійно реагувати на загрози. Такі системи можуть автоматично блокувати шкідливі IP-адреси, ізолювати заражені пристрої у мережі або аналізувати нові типи шкідливого програмного забезпечення. Деякі

сучасні платформи кібербезпеки використовують концепцію self-healing systems, тобто систем, здатних самостійно відновлювати нормальну роботу після атаки та адаптувати механізми захисту до нових загроз. Це дозволяє значно скоротити час реагування на інциденти та зменшити навантаження на фахівців з кібербезпеки [2; 6; 7].

Таблиця 1 – Автоматизація різних видів кібератак

Вид кібератаки	Як відбувається автоматизація	Що це дає зловмисникам	Ризики для систем захисту
Фішинг	ШІ автоматично створює велику кількість правдоподібних повідомлень і підлаштовує їх під різні категорії користувачів	Підвищення масовості атак, економія часу, більша переконливість повідомлень	Користувачам важче відрізнити підроблені листи від справжніх
Соціальна інженерія	Алгоритми аналізують відкриті дані про людину та формують персоналізовані сценарії впливу	Атаки стають більш адресними та психологічно ефективними	Зростає ймовірність успішного обману користувачів
Пошук вразливостей	Автоматизовані системи швидко аналізують програми, мережі й сервіси на наявність слабких місць	Прискорення виявлення потенційних точок проникнення	Скорочується час між появою вразливості та її використанням
Шкідливе програмне забезпечення	ШІ може змінювати поведінку шкідливого коду для обходу стандартних засобів виявлення	Підвищення стійкості шкідливих програм до захисту	Традиційні сигнатурні методи виявлення працюють гірше
DDoS-атаки	Бот-мережі автоматично координують великий потік запитів до сервісів чи сайтів	Масштабність і безперервність перевантаження ресурсів	Сервери та мережеві служби можуть втрачати доступність
Підбір облікових даних	Автоматизовані механізми масово перевіряють викрадені або типові дані входу	Прискорення несанкціонованого доступу до акаунтів	Зростає кількість атак на автентифікацію
Дезінформаційні кампанії	Генеративні моделі швидко створюють тексти, зображення чи повідомлення для масового поширення	Масове охоплення та швидке поширення неправдивої інформації	Ускладнюється перевірка достовірності контенту
Адаптивні атаки	Алгоритми аналізують реакцію системи захисту та змінюють модель поведінки атаки	Підвищення шансів обійти захисні механізми	Системи безпеки повинні постійно адаптуватися до нових сценаріїв

Інтеграція штучного інтелекту з іншими технологіями, такими як блокчейн, хмарні обчислення та Інтернет речей, відкриває нові можливості для створення комплексних систем інформаційної безпеки. Наприклад, поєднання технологій блокчейну та штучного інтелекту може забезпечити більш надійний контроль доступу до даних, а використання ШІ в IoT-мережах дозволяє виявляти аномалії у роботі підключених пристроїв. Однак розширення цифрової інфраструктури також збільшує поверхню потенційних атак, що потребує постійного вдосконалення систем захисту [3; 9].

Разом із перевагами використання ШІ-агентів виникають і нові виклики. Одним із них є проблема прозорості алгоритмів штучного інтелекту, оскільки складні нейронні мережі можуть приймати рішення, які складно пояснити або перевірити. Крім того, існує ризик використання таких технологій у злочинних цілях, що може призвести до появи нових форм кібератак. У зв'язку з цим важливим є розвиток правових механізмів регулювання використання штучного інтелекту у сфері кібербезпеки та створення етичних стандартів його застосування [3; 8; 9].

Таким чином, ШІ-агенти стають важливим елементом сучасного кіберпростору. Вони можуть використовуватися як для здійснення складних кібератак, так і для ефективного захисту інформаційних систем. Подальший розвиток цих технологій потребує комплексного підходу, що поєднує технічні інновації, міжнародні стандарти безпеки та відповідальне використання штучного інтелекту. Лише за умови поєднання наукових досліджень, технологічного розвитку та правового регулювання можливо забезпечити ефективний захист інформації в умовах стрімкого розвитку цифрового середовища.

Список використаних джерел:

1. Microsoft. What is AI for cybersecurity? URL: <https://www.microsoft.com/uk-ua/security/business/security-101/what-is-ai-for-cybersecurity> (дата звернення: 17.03.2026).

2. IBM. Artificial Intelligence in Cybersecurity. URL: <https://www.ibm.com/topics/artificial-intelligence-cybersecurity> (дата звернення: 16.03.2026).
3. ENISA. Artificial Intelligence and Cybersecurity: Opportunities and Challenges. URL: <https://www.enisa.europa.eu/publications/artificial-intelligence-and-cybersecurity> (дата звернення: 18.03.2026).
4. Cisco. AI and Cybersecurity: How AI is Transforming Security. URL: <https://www.cisco.com/c/en/us/products/security/ai-cybersecurity.html> (дата звернення: 15.03.2026).
5. CrowdStrike. Artificial Intelligence in Cybersecurity. URL: <https://www.crowdstrike.com/cybersecurity-101/artificial-intelligence/> (дата звернення: 17.03.2026).
6. Palo Alto Networks. The Role of AI in Cybersecurity. URL: <https://www.paloaltonetworks.com/cyberpedia/what-is-ai-in-cybersecurity> (дата звернення: 16.03.2026).
7. Fortinet. AI in Cybersecurity Explained. URL: <https://www.fortinet.com/resources/cyberglossary/artificial-intelligence-in-cybersecurity> (дата звернення: 18.03.2026).
8. Check Point. AI and Cybersecurity: Benefits and Risks. URL: <https://www.checkpoint.com/cyber-hub/cyber-security/what-is-ai-in-cybersecurity/> (дата звернення: 15.03.2026).
9. Cloudflare. AI Security: Threats and Protection. URL: <https://www.cloudflare.com/learning/security/what-is-ai-security/> (дата звернення: 19.03.2026).