

УДК 341.3

DOI <https://doi.org/10.24144/2307-3322.2025.91.5.12>

ШТУЧНИЙ ІНТЕЛЕКТ У СФЕРІ НАЦІОНАЛЬНОЇ БЕЗПЕКИ: ЕТИЧНІ ТА ПРАВОВІ АСПЕКТИ

Блавацька Н.М.,
кандидат технічних наук, доцент,
доцент НА СБУ
ORCID: 0000-0003-2247-8008

Рибка Д.М.,
доктор філософії,
старший викладач НА СБУ
ORCID: 0000-0002-8824-0089

Блавацька Н.М., Рибка Д.М. Штучний інтелект у сфері національної безпеки: етичні та правові аспекти.

Стрімкий розвиток технологій штучного інтелекту (ШІ) суттєво трансформує систему національної безпеки, змінюючи традиційні підходи до розвідки, кібероборони, стратегічного планування та ведення бойових дій. Інтеграція інтелектуальних алгоритмів у безпекову сферу забезпечує оперативну обробку великих масивів даних, точніше прогнозування загроз і підвищення ефективності управлінських рішень. Водночас, активне впровадження ШІ породжує комплекс етичних та правових викликів, пов'язаних із питаннями відповідальності, прозорості, пояснюваності алгоритмів, захисту прав людини та дотримання норм Міжнародного гуманітарного права.

Особливу увагу у статті приділено проблематиці автономних систем озброєння, які створюють ризик втрати людського контролю над застосуванням сили, а також потенційній упередженості алгоритмів, що може призвести до дискримінації та помилкових рішень у бойових умовах. Розглядаються етичні дилеми, пов'язані з «ефектом віддаленості» між оператором і наслідками його дій, що ставить під сумнів моральну відповідальність за ухвалені рішення.

З правової точки зору аналізуються питання відповідальності за шкоду, заподіяну автономними системами, та прогалини у Міжнародному праві щодо визначення суб'єктів такої відповідальності. Акцентується увага на тому, що принципи розрізнення та пропорційності, закладені в Міжнародному гуманітарному праві, є складними для відтворення в алгоритмічних системах.

У статті узагальнюється міжнародний досвід, зокрема ініціативи НАТО та ЄС, спрямовані на формування етичних стандартів і запровадження принципу значущого людського контролю. Автори обґрунтовують необхідність створення глобальних механізмів аудиту, прозорості та підзвітності систем ШІ, а також захисту прав людини від зловживань у сфері цифрової безпеки.

Додатково підкреслюється, що ефективне регулювання ШІ потребує міждисциплінарної співпраці держав, науковців і технологічних компаній для збереження глобальної стабільності та миру.

Дослідження підкреслює, що формування етико-правової архітектури використання ШІ у національній безпеці є ключовою умовою для забезпечення відповідального, гуманного та законного розвитку технологій майбутнього.

Ключові слова: штучний інтелект, національна безпека, етика ШІ, правове регулювання, автономні системи озброєння, прозорість алгоритмів, відповідальність, права людини.

Blavatska N.M., Rybka D.M. Artificial intelligence in national security: ethical and legal aspects.

The rapid development of artificial intelligence (AI) technologies is profoundly transforming the system of national security, reshaping traditional approaches to intelligence, cyber defense, strategic planning, and warfare. The integration of intelligent algorithms into the security domain enables real-time processing of vast data volumes, more accurate threat prediction, and greater efficiency in decision-

making. At the same time, the active implementation of AI gives rise to a complex range of ethical and legal challenges related to accountability, transparency, algorithmic explainability, human rights protection, and compliance with the norms of International Humanitarian Law.

Particular attention in the article is devoted to the issue of autonomous weapons systems, which pose a risk of losing human control over the use of force, as well as potential algorithmic bias that may lead to discrimination and erroneous decisions in combat situations. The ethical dilemmas associated with the “distance effect” between the operator and the consequences of their actions are analyzed, questioning moral responsibility for the decisions made.

From a legal perspective, the article examines questions of liability for harm caused by autonomous systems and the gaps in International Law regarding the identification of responsible actors. Emphasis is placed on the principles of distinction and proportionality enshrined in International Humanitarian Law, which remain difficult to replicate in algorithmic systems.

The study summarizes international experience, including NATO and EU initiatives aimed at establishing ethical standards and implementing the principle of meaningful human control. The authors justify the need to create global mechanisms for auditing, transparency, and accountability of AI systems, as well as for protecting human rights from abuses in the sphere of digital security.

Additionally, it is emphasized that effective AI regulation requires interdisciplinary cooperation among states, scholars, and technology companies to maintain global stability and peace.

The research highlights that building an ethical and legal framework for the use of AI in national security is a key condition for ensuring responsible, humane, and lawful technological development in the future.

Key words: artificial intelligence, national security, AI ethics, legal regulation, autonomous weapons systems, algorithmic transparency, accountability, human rights.

Постановка проблеми. Сфера національної безпеки традиційно є каталізатором технологічних інновацій. Сучасний розвиток штучного інтелекту та машинного навчання відкриває нову еру в обороні та безпеці. Застосування ШІ охоплює широкий спектр: від аналізу величезних масивів розвідувальних даних та прогнозування загроз до управління безпілотними апаратами та розробки летальних автономних систем озброєння. Однак, переваги ШІ супроводжуються серйозними ризиками. Системи ШІ можуть функціонувати як «чорні скриньки», ускладнюючи встановлення причинно-наслідкових зв'язків та відповідальності за їхні рішення. Це особливо критично у контексті, де ціна помилки вимірюється людськими життями або ескалацією конфлікту. Отже, виникає нагальна потреба у всебічному дослідженні та регулюванні етичних і правових аспектів використання ШІ у цій високочутливій сфері.

Аналіз останніх досліджень і публікацій. Штучний інтелект стає одним із центральних об'єктів уваги сучасної наукової спільноти. Тема його застосування цікавить дослідників із різних сфер, від права й етики до інформаційних технологій і безпеки. У зв'язку з цим з'являється значний масив наукових публікацій, що аналізують як теоретичні основи, так і практичні наслідки інтеграції систем штучного інтелекту в соціальні та інституційні структури.

Так, Тарасюк В.М. провів ґрунтовне дослідження правового регулювання застосування штучного інтелекту в Україні, розглядаючи існуючі національні підходи та можливості їхньої гармонізації з міжнародними стандартами. Його робота спрямована на виявлення нормативних прогалин і пропозицій щодо адаптації внутрішнього законодавства з урахуванням кращих практик та документів міжнародних організацій. У своїх публікаціях Горелова В.Ю. зосереджується на етичних та правових перспективах використання ШІ в українському контексті. Вона аналізує проблеми відповідальності, прозорості та захисту прав людини при автоматизованих рішеннях, окреслюючи етичні рамки, які мають супроводжувати розробку й впровадження алгоритмічних рішень. Македон О.А. приділив увагу питанням застосування ШІ у юридичній практиці. Він досліджував етичні дилеми та правові наслідки використання аналітичних і автоматизованих інструментів для правозастосування, підкреслюючи необхідність гарантій справедливого доступу до правосуддя й механізмів контролю за алгоритмічними системами.

Явар Беті зосередився на проблемі «чорної скриньки» в системах ШІ, а саме на труднощах, пов'язаних із визначенням причинно-наслідкових зв'язків у їхніх рішеннях. Його дослідження висвітлює питання інтерпретованості та зрозумілості алгоритмічних рішень, що має критичне значення для встановлення відповідальності, забезпечення прозорості та довіри до автоматизованих систем.

Водночас, попри широкий спектр досліджень, тема правового регулювання та етичних аспектів застосування штучного інтелекту містить значну кількість невирішених питань особливо коли йдеться про її перетин зі сферою національної безпеки. Багатоаспектність і варіативність проблем, а також надзвичайно швидкий розвиток ІТ-технологій на базі ШІ створюють складний контекст, у якому законодавча та етична база часто не встигає за практичними потребами.

Метою даної статті є здійснити комплексний аналіз етичних і правових аспектів застосування технологій штучного інтелекту у сфері національної безпеки та оборони, з'ясувати основні ризики та виклики, що виникають у процесі їх впровадження, а також обґрунтувати необхідність формування ефективних регуляторних механізмів і етичних стандартів, спрямованих на забезпечення відповідального, безпечного та правомірного використання систем ШІ у цій галузі.

Виклад основного матеріалу. Інтеграція технологій штучного інтелекту докорінно змінює архітектуру національної безпеки, формуючи нову парадигму управління ризиками, реагування на загрози та ведення сучасних конфліктів. Використання ШІ охоплює кілька ключових напрямів.

Перший з них це розвідка та аналітика даних. Системи ШІ здатні в реальному часі обробляти величезні обсяги інформації з відкритих та технічних джерел, виявляючи приховані закономірності, потенційні загрози й ознаки дестабілізаційних дій. Алгоритми прогнозного аналізу дозволяють своєчасно ідентифікувати ризики та підвищують ефективність прийняття стратегічних рішень у сфері безпеки [1].

Далі це кіберзахист і кібероборона. ШІ активно застосовується для автоматизованого моніторингу кіберпростору, виявлення аномалій, шкідливих вторгнень і реагування на атаки. Завдяки здатності навчатися на нових паттернах загроз, такі системи підвищують стійкість державних і військових мереж до кіберагресії, забезпечуючи оперативність і адаптивність захисних дій [2].

Ще слід згадати про автономні системи озброєння. Військові дрони, роботизовані платформи та напівавтономні бойові системи, що використовують алгоритми ШІ для навігації, розпізнавання цілей і прийняття рішень, стають невід'ємним елементом сучасних збройних сил [3]. Водночас, саме цей напрям створює найбільші етичні й стратегічні ризики – від можливості ескалації конфліктів до прийняття помилкових або невинуватих рішень через алгоритмічну упередженість. Особливе занепокоєння викликає поступове розмивання меж людського контролю над застосуванням сили та відповідальності за наслідки [4].

Таким чином, впровадження ШІ у безпековий сектор відкриває нові горизонти ефективності та аналітичних можливостей, але одночасно висуває безпрецедентні вимоги до етичного, правового та стратегічного регулювання, покликаного запобігти втраті людського контролю над технологічно детермінованими рішеннями.

Етичні дилеми є центральними для дискусії про ШІ у військовій сфері. Автономність ШІ-систем, зокрема летальних автономних озброєнь, ставить під загрозу принцип значущого людського контролю. Застосування автономних систем створює психологічний та моральний бар'єр між оператором та наслідками його дій, що може призвести до легшого прийняття рішень про застосування сили. Міжнародні документи та ініціативи пропонують, щоб рішення про застосування сили завжди перебувало під значним людським наглядом – людина має усвідомлювати наслідки, мати право втрутитися [5].

Деякі системи ШІ в воєнних або безпекових контекстах навчаються на даних, які містять історичні чи соціальні упередження – расові, гендерні, етнічні, територіальні. Якщо не запровадити постійний аудит і перевірку даних, це може призвести до дискримінації або неправильного профілювання цих систем. Якщо навчальні дані містять історичні, географічні чи соціальні упередження, система ШІ буде їх відтворювати, що може призвести до дискримінаційних або неправомірних рішень, наприклад, систематично неправильної ідентифікації цивільних об'єктів у певних регіонах [6].

Часто алгоритмічні моделі є «чорною скринькою», і неясно, чому прийняте саме таке рішення, тобто відсутня прозорість та пояснюваність. У сфері безпеки це може зруйнувати довіру, унеможливити оскарження неправомірних рішень, та ускладнює встановлення відповідальності [7]. Складні нейронні мережі часто не дозволяють зрозуміти логіку, за якою було прийнято рішення, наприклад, чому система ідентифікувала об'єкт як ворожу ціль. Це унеможливорює належне розслідування та встановлення відповідальності у разі помилки. Військові системи ШІ повинні бути пояснюваними, оскільки рішення на полі бою мають бути обґрунтовані.

Застосування ШІ у національній безпеці стикається з прогалинами та невизначеністю у чинному національному та міжнародному праві.

З правового боку ключові питання зосереджені на відповідальності. Хто має нести юридичну відповідальність за шкоду, заподіяну автономною системою – оператор, розробник, виробник, чи держава. У разі порушення Міжнародного гуманітарного права автономною системою виникає прогалина відповідальності [8]. Сама ШІ-система не може нести кримінальну чи юридичну відповідальність. Відсутність чіткої лінії відповідальності підриває фундаментальний принцип верховенства права і може призвести до безкарності за воєнні злочини.

Крім того дві фундаментальні норми Міжнародного гуманітарного права, а саме принцип розрізнення між комбатантами і цивільними та принцип пропорційності, тобто очікувана військова перевага має бути пропорційна очікуваним цивільним втратам, стають найбільшим викликом для автономних систем. Підлягає сумніву здатність машини, яка навчалася на базі даних, адекватно оцінити контекст і намір людини, щоб відрізнити, наприклад, озброєного комбатанта від цивільної особи, що тримає інструмент. До того ж рішення про пропорційність є глибоко контекстуальним і суб'єктивним морально-правовим судженням, яке традиційно вимагає людської оцінки. Передача цього судження машині є однією з найгостріших правових суперечностей.

У міжнародній практиці напрацьовано низку підходів, що демонструють прагнення до формування етичних і правових стандартів використання штучного інтелекту у сфері оборони. Так, у науковій статті *Ethical Principles for Artificial Intelligence in National Defence* (2021), визначено п'ять ключових засад відповідального застосування ШІ: виправданість його використання, справедливість і прозорість алгоритмічних систем, моральна відповідальність людини, забезпечення значущого людського контролю та надійність технологій [9].

Водночас у межах міжнародного політичного дискурсу важливе місце посідає *Political Declaration on Responsible Military Use of Artificial Intelligence and Autonomy*, яку підписали понад п'ятдесят держав. Цей документ наголошує на необхідності відповідального використання автономних і напіваавтономних систем, дотримання принципів міжнародного права та гарантування захисту прав людини навіть у контексті воєнних дій [10].

Додатковий внесок у розуміння етичних викликів робить дослідження *Ethical Considerations for the Military Use of Artificial Intelligence in Visual Reconnaissance* (2025). Автори аналізують практичні аспекти застосування ШІ у сфері візуальної розвідки у морському, повітряному та сухопутному просторах, акцентуючи увагу на принципах FATE: справедливості (Fairness), підзвітності (Accountability), прозорості (Transparency) та етичності (Ethics) [11].

Таким чином, міжнародний досвід засвідчує, що провідні наукові та політичні інституції рухаються в напрямі створення єдиної етичної рамки, яка поєднує технічні стандарти безпеки з фундаментальними цінностями гуманізму, права та відповідальності.

Застосування штучного інтелекту у сфері безпеки створює суттєві ризики для дотримання прав людини як у мирний час, так і в умовах внутрішніх загроз. Одним із найбільш дискусійних аспектів є проблема масового спостереження та приватності. Системи ШІ здатні інтегрувати величезні обсяги інформації з камер відеоспостереження, сенсорів, соціальних мереж і державних баз даних, формуючи детальні цифрові профілі громадян. Така практика, навіть якщо вона має на меті підвищення безпеки, може призвести до надмірного контролю за приватним життям осіб і створює передумови для зловживань.

Не менш небезпечним є автоматизоване профілювання, коли алгоритми класифікують людей на основі поведінкових або соціальних характеристик. Це може порушувати право на приватність і призводити до дискримінації, наприклад, через помилкове визначення “потенційних злочинців” або “неосяльких” громадян. Подібні технології здатні обмежувати свободу пересування, вираження поглядів і участі в громадському житті.

Іншим негативним наслідком стає так званий ефект охолодження, який полягає в тому, що усвідомлення постійного нагляду з боку держави або автоматизованих систем може стримувати людей від законної громадянської активності, участі у протестах чи висловлення критичних думок.

Окрему загрозу становить порушення права на належну правову процедуру. Якщо рішення, що впливають на свободу, репутацію чи навіть життя людини, наприклад, автоматичне внесення до списку потенційних, ухвалюються за допомогою алгоритмів, виникає питання про можливість їхнього оскарження. Без забезпечення пояснюваності ШІ і права особи знати причини прийняття рішень, принцип справедливого судового розгляду фактично втрачає сенс.

Таким чином, узагальнюючи викладене, можна стверджувати, що ефективно та етично виважене використання штучного інтелекту у сфері національної безпеки вимагає цілісного підхо-

ду, який поєднує міжнародно-правові, технічні та гуманітарні засади, що покликані захистити фундаментальні права людини та не допустити перетворення технологій на інструмент масового контролю.

Першочерговим завданням є закріплення на міжнародному рівні принципу значущого людського контролю над системами, що здатні приймати рішення про застосування сили. Це означає, що жодна автономна або напівавтономна технологія не повинна діяти без реальної можливості втручання людини. Такий контроль є не лише технічною вимогою, а й моральною гарантією дотримання людської гідності та відповідальності у воєнних діях. Збереження участі людини у критичних рішеннях має стати нормою міжнародного права, закріпленою в договорах і політичних деклараціях.

Не менш важливим напрямом є створення міжнародних стандартів для тестування, аудиту та забезпечення пояснюваності алгоритмів, що використовуються у сфері оборони. Алгоритмічні рішення мають бути прозорими, перевіреними на предмет технічної надійності, відсутності дискримінаційних упереджень та відповідності етичним критеріям. Це дозволить не лише підвищити довіру до таких систем, а й гарантувати, що рішення, прийняті за участю ШІ, можна відтворити, перевірити та, за потреби, оскаржити.

Крім того, необхідним є запровадження чітких юридичних режимів відповідальності, які розмежовують компетенцію і зобов'язання різних учасників життєвого циклу ШІ, від розробників і виробників до командирів та операторів. Подібна модель дає змогу уникнути прогалини відповідальності, коли жоден із суб'єктів не може бути притягнутий до відповідальності у разі порушення Міжнародного гуманітарного права або заподіяння шкоди цивільному населенню.

Окрему увагу слід приділити захисту прав людини у цифровому середовищі. Використання ШІ в розвідці та внутрішній безпеці не повинно перетворюватися на інструмент масового спостереження чи неправомірного профілювання. Правове регулювання має передбачати обмеження щодо збору, зберігання та використання персональних даних, а також гарантувати право кожної особи на оскарження рішень, ухвалених автоматизованими системами.

Таким чином, формування етичної та правової архітектури використання штучного інтелекту у сфері безпеки є стратегічним завданням, від якого залежить не лише ефективність технологічного розвитку, а й збереження гуманістичних засад міжнародного правопорядку.

Висновок Застосування ШІ у сфері національної безпеки представляє собою подвійний потенціал: з одного боку – значні переваги у швидкості, точності та масштабі реагування, з іншого – ризики, які, якщо їх не контролювати, можуть порушити фундаментальні етичні принципи та міжнародне право.

Щоб забезпечити відповідальне й законне використання ШІ, необхідно поєднувати технологічний розвиток із чітким етичним та правовим регулюванням. Важливо, щоб держави брали участь у формуванні міжнародних норм та стандартів, які гарантують прозорість, справедливість, значущий людський контроль, і відповідальність за наслідки використання ШІ. Без таких кроків ШІ може підривати довіру до інститутів безпеки та спричиняти серйозні порушення права людини й міжнародного гуманітарного права.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ:

1. Штучний інтелект в OSINT. URL: <https://hackyoumom.com/kibervijna/stuchnyj-intelekt-v-osint-yak-pokrashhyty-rozsliduvannya-u-2024-rocz> (дата звернення: 02.10.2025).
2. Штучний інтелект та кібербезпека. URL: <https://www.bdo.ua/uk-ua/insights-2/information-materials/2024/corporate-cybersecurity-ai-role-in-data-protection> (дата звернення: 02.10.2025).
3. Гмиря В.П., Романовська Л.В., Фіцайло Т.М., Нікітченко А.О. Штучний інтелект і автономні системи озброєння: нова епоха військових операцій. *Збірник наукових праць Державного науково-дослідного інституту випробувань і сертифікації озброєння та військової техніки*. Черкаси, 2025. № 2(24) С. 7–14. URL: <https://dndivsovt.com/index.php/journal/article/view/563> (дата звернення: 02.10.2025).
4. Гбур З.В. Основи інформаційної безпеки держави в умовах війни. URL: <http://baltijapublishing.lv/omp/index.php/bp/catalog/view/237/6325/13361-1> (дата звернення: 02.10.2025).
5. Кутовий О.В., Бурма С.К. Автономні системи озброєнь і штучний інтелект як виклик міжнародному гуманітарному праву та правам людини. *Науковий вісник Ужгородського На-*

- ціонального Університету. Ужгород, 2025. С.185–194. URL: <https://visnyk-juris-uzhnu.com/wp-content/uploads/2025/09/24-4.pdf> (дата звернення: 02.10.2025).
6. Хазіка Саджид Упередження ШІ та культурні стереотипи: наслідки, обмеження та пом'якшення. URL: <https://www.unite.ai/uk/ai-bias-cultural-stereotypes-effects-limitations-mitigation/> (дата звернення: 02.10.2025).
 7. Проблема чорної скриньки ШІ: виклики та рішення для прозорого майбутнього. URL: <https://crypta.kyiv.ua/tehnologii-blokchejn/problema-chornoi-skrinki-shi-vikliki-ta-rishennja/> (дата звернення: 02.10.2025).
 8. Міжнародне гуманітарне право. URL: <https://unity.gov.ua/diyalnist/napryamky-proektiv-minreintegracziyi/mizhnarodne-gumanitarne-pravo/> (дата звернення: 02.10.2025).
 9. Mariarosaria Taddeo, David McNeish, Alexander Blanchard, Elizabeth Edgar Ethical Principles for Artificial Intelligence in National Defence. URL: <https://link.springer.com/article/10.1007/s13347-021-00482-3> (дата звернення: 02.10.2025).
 10. Political Declaration on Responsible Military Use of Artificial Intelligence and Autonomy. URL: <https://www.state.gov/bureau-of-arms-control-deterrence-and-stability/political-declaration-on-responsible-military-use-of-artificial-intelligence-and-autonomy> (дата звернення: 02.10.2025).
 11. Mathias Anneken, Nadia Burkart, Fabian Jeschke, Achim Kuwertz-Wolf, Almuth Müller, Arne Schumann, Michael Teutsch Ethical Considerations for the Military Use of Artificial Intelligence in Visual Reconnaissance. URL: <https://arxiv.org/abs/2502.03376> (дата звернення: 02.10.2025).