

УДК 327:004.8
DOI: 10.60022/3(4)-21S

Кубко Валентина Петрівна

кандидат філософських наук
доцент кафедри міжнародних відносин та права
Національний університет «Одеська політехніка», Україна

Kubko Valentyna

Candidate of Philosophical Sciences
Associate Professor of the Department of International Relations and Law
Odessa Polytechnic National University, Ukraine
ORCID: 0000-0003-0433-2013

КОНЦЕПТ ШТУЧНОГО ІНТЕЛЕКТУ В СУЧАСНІЙ ПАРАДИГМІ МІЖНАРОДНОЇ БЕЗПЕКИ

Анотація. Здійснено комплексне теоретичне дослідження штучного інтелекту (ШІ) як системного чинника трансформації сучасної парадигми міжнародної безпеки. Обґрунтовано перехід від традиційних моделей стратегічного аналізу до алгоритмічно керованих систем, що докорінно змінює онтологію прийняття рішень у воєнно-політичній сфері. Авторська увага зосереджена на аналізі епістемічних зміщень, зокрема на явищі «стиснення часу» та виникненні когнітивного феномену «automation bias», що веде до деградації критичної рефлексії оператора перед обличчям надшвидких цифрових рішень.

Окрему увагу приділено ролі генеративних мовних моделей як інструментів синтезу стратегічних сценаріїв та обробки неструктурованих даних. Виявлено специфічні загрози, пов'язані з використанням «стохастичних папуг», які здатні генерувати переконливі, але змістовно порожні аналітичні звіти, що створює ризик прийняття рішень на основі алгоритмічних галюцинацій. Досліджено когнітивні ризики методів навчання з підкріпленням (RLHF), які можуть консервувати упередження розробників та створювати ефект «алгоритмічного підтвердження», викривляючи об'єктивну картину безпекового середовища.

Крізь призму вітчизняного наукового дискурсу розглянуто досвід відсічі збройній агресії, де ШІ виступає домінуючим складником розвідки та ідентифікації цілей (зокрема в межах програм типу Project Maven). Доведено, що інтеграція ШІ в державне управління та оборонні альянси потребує збереження принципу «людина в контурі» як базової умови етичного контролю та стратегічної стабільності. Концептуалізовано ШІ не просто як технологічний інструмент, а як інтегральну складову нової архітектури глобального управління, що потребує негайного формування міжнародно-правових механізмів регулювання.

Ключові слова: штучний інтелект, міжнародна безпека, стратегічна стабільність, генеративні мовні моделі, алгоритмізація, когнітивна безпека, RLHF, стратегічне стримування.

THE CONCEPT OF ARTIFICIAL INTELLIGENCE IN THE CONTEMPORARY INTERNATIONAL SECURITY PARADIGM

Abstract. The article conducts a comprehensive theoretical study of artificial intelligence (AI) as a systemic factor in the transformation of the contemporary international security paradigm. The transition from traditional strategic analysis models to algorithmically driven systems is substantiated, fundamentally altering the ontology of decision-making in the military-political sphere. The author's attention is focused on the analysis of epistemic shifts within the security environment, specifically the phenomenon of "time compression" and the emergence of the cognitive phenomenon known as "automation bias," which leads to the degradation of critical reflection by human operators in the face of ultra-fast digital solutions.

Special emphasis is placed on the role of generative language models as instruments for synthesizing strategic scenarios and processing unstructured data. Specific threats associated with the use of "stochastic parrots" are identified, which are capable of generating extremely persuasive yet substantively hollow



analytical reports, creating the risk of decisions based on algorithmic hallucinations. The research delves into the cognitive risks of reinforcement learning from human feedback (RLHF) methods, which can institutionalize the biases of developers and create an "algorithmic confirmation" effect, thereby distorting the objective picture of the security environment.

Through the prism of Ukrainian scientific discourse, the experience of repelling armed aggression is examined, where AI emerges as a dominant component of intelligence and targeting (particularly within programs such as Project Maven). It is proven that the integration of AI into public administration and defense alliances requires the preservation of the "human-in-the-loop" principle as a fundamental condition for ethical control and strategic stability. The study conceptualizes AI not merely as a technological tool but as an integral component of a new global governance architecture that necessitates the immediate formation of international legal regulatory mechanisms. The paper argues that the future of international security depends on the ability of the global community to develop unified rules for "algorithmic behavior," where humans remain the subjects of ethical control, preventing the degradation of critical analysis in the digital age.

Keywords: artificial intelligence, international security, strategic stability, generative language models, algorithmization, cognitive security, RLHF, strategic deterrence.

Постановка проблеми. У сучасних умовах докорінної трансформації міжнародної системи безпеки штучний інтелект (ШІ) переходить із розряду допоміжних технологічних інструментів у статус засадничого фактора, що змінює архітектуру міжнародних відносин. Його системна інтеграція у військову, інформаційну та воєнно-політичну сфери ініціює перегляд класичних концептів безпеки, державного суверенітету та стратегічного балансу сил. На відміну від попередніх технологічних циклів, ШІ трансформує не лише швидкість та масштаби опрацювання даних, а й саму онтологію прийняття рішень, формуючи нову якість стратегічного мислення, де алгоритмічна логіка починає конкурувати з людською суб'єктивністю.

Особливої актуальності зазначена проблематика набуває в умовах гібридних загроз, де конвергенція кібероперацій, інформаційно-психологічних впливів та інтелектуальних систем аналізу даних створює нові виклики для міжнародної стабільності. Досвід відсічі повномасштабної агресії РФ проти України став глобальним каталізатором практичної апробації ШІ у сферах розвідки, ідентифікації цілей та стратегічних комунікацій, що дозволяє розглядати інтелектуальні системи як домінуючий складник сучасного безпекового середовища.

Водночас стрімке впровадження алгоритмічних моделей у процеси стратегічного планування актуалізує проблему когнітивної безпеки: здатність ШІ відтворювати людські упередження та створювати ефект «алгоритмічного підтвердження» загрожує деформацією об'єктивного аналізу ситуації. У цьому контексті виникає гостра наукова необхідність комплексного осмислення штучного інтелекту не просто як сукупності технологій, а як системного концепту в новій парадигмі міжнародної безпеки, що докорінно змінює стратегічну поведінку держав та механізми глобального управління.

Аналіз останніх досліджень і публікацій. Проблематика інтеграції штучного інтелекту (ШІ) в сучасну систему міжнародної безпеки є об'єктом прискільки уваги як зарубіжних, так і вітчизняних дослідників. Фундаментальне значення для теоретичного обґрунтування ролі штучного інтелекту мають праці Б. Ндзедзе та Т. Марвали [1], які здійснюють ревізію класичних парадигм міжнародних відносин (реалізму, лібералізму та конструктивізму) в умовах алгоритмізації глобальних процесів. Дослідники доводять, що ШІ стає новим «фактором сили», який змінює структуру міжнародної системи. Суттєві ризики, пов'язані з автономністю військових систем та трансформацією концепції «людина в контурі», детально проаналізовані у працях П. Шарра [2], який на прикладі програми *Project Maven* досліджує загрози надмірної автоматизації. Питання еволюції стратегічних відносин під впливом ШІ та можливість «революції у військових справах» розглядаються у роботах К. Пейна [3] та Дж. Джонсона [4], які акцентують увагу на зміні самої природи збройних конфліктів майбутнього.

Внесок у розуміння впливу ШІ на стратегічну та ядерну стабільність зробили фахівці Центру за нову американську безпеку (CNAS), зокрема М. Горовіч та А. Велез-Грін [5], які запропонували заходи зміцнення довіри для нівелювання ризиків алгоритмічних помилок. Етичні аспекти та проблему «вузького місця» у навчанні моделей, зокрема через методи RLHF (*Reinforcement Learning from Human Feedback* – навчання з підкріпленням на основі зворотного зв'язку від людини), що можуть призводити до розбіжностей у ціннісних орієнтирах між людиною та машиною, досліджує Б. Крістіан [6]. Водночас Е. Бендер [7] застерігає від надмірної довіри до великих мовних моделей, вказуючи на їхню здатність лише репродукувати дані без глибокого розуміння контексту безпекового середовища.

В українському науковому дискурсі питання військового застосування ШІ крізь призму міжнародного права досліджують І. Забара та О. Деркач [8]. Роль ШІ як глобального мегатренда та його

вплив на трансформацію політичних систем розкрито у працях О. Пархомчук [9], тоді як Н. Блавацька та Д. Рибка [10] зосереджуються на етико-правових колізіях використання інтелектуальних систем у національній безпеці. Окремий пласт досліджень присвячений інформаційній безпеці: О. Дубовський [11] аналізує суб'єктність ШІ в контексті інформаційного протистояння, а Р. Пантюшенко та Ю. Чайка [12] висвітлюють інновації у сфері захисту критичної інфраструктури. Питання адаптації досвіду США для України розглядає І. Лопатченко [13], загрози міжнародній стабільності – О. Фурсай [14], а генезис технології – М. Єфремов та ін. [15].

Водночас, попри наявність ґрунтовних напрацювань, питання комплексного теоретичного осмислення штучного інтелекту як інтегральної складової нової парадигми міжнародної безпеки залишається недостатньо розкритим. Зокрема, потребує поглибленого аналізу проблема алгоритмічних викривлень стратегічного аналізу та їхнього впливу на архітектуру прийняття рішень, де використання генеративних моделей може створювати нові вразливості для національної стійкості та глобальної стабільності.

Метою статті є концептуалізація штучного інтелекту як системного чинника в сучасній парадигмі міжнародної безпеки, аналіз його ролі у трансформації стратегічного середовища, а також оцінка впливу алгоритмізації прийняття рішень на характер міжнародних конфліктів.

Виклад основного матеріалу. Процес концептуальної трансформації сучасної архітектури безпеки під впливом штучного інтелекту (ШІ) виявляється насамперед у радикальній зміні епістемічних засад стратегічного аналізу. ШІ дедалі глибше інтегрується в когнітивні цикли управління, що дозволяє стверджувати про перехід від традиційних моделей «війни ресурсів» до інтелектуально містких моделей «війни алгоритмів». Як слушно зауважує Дж. Джонсон [4], впровадження інтелектуальних систем у військову діяльність не просто доповнює наявні спроможності, а формує якісно нові парадигми ведення збройної боротьби, де домінуючим фактором стає здатність до адаптивного реагування в умовах експоненційного зростання обсягів даних. На наш погляд, визначальним наслідком цієї трансформації є зміна часових параметрів ухвалення рішень, що спричиняє звуження горизонтів політичного маневру та вимагає від держав докорінного перегляду класичних підходів до забезпечення національної стійкості.

Згідно з концепцією Б. Ндзєндзе та Т. Марвали, інтеграція штучного інтелекту в оборонний сектор створює специфічну «алгоритмічну дилему безпеки»: держави змушені впроваджувати автономні системи не лише для отримання переваги, а й через страх відстати від опонентів, що автоматично знижує поріг початку конфлікту [1]. На наш погляд, це підтверджує тезу про те, що ШІ трансформує саму архітектуру міжнародного стримування, роблячи її більш динамічною та менш передбачуваною.

Окремим чинником цієї трансформації стає стрімке впровадження генеративних мовних моделей (*Large Language Models*), які виводять використання ШІ за межі суто обчислювальних операцій у сферу синтезу смислів та стратегічного прогнозування. Як зазначає Дж. Джонсон, здатність генеративних систем миттєво опрацьовувати та структурувати колосальні масиви неструктурованої інформації дозволяє моделювати багатоваріантні сценарії розвитку конфліктів, що докорінно змінює характер аналітичної підтримки дипломатії та воєнно-політичного планування [4]. Проте використання генеративного ШІ актуалізує проблему, яку Е. Бендер визначає через концепт «стохастичних папуг»: моделі здатні генерувати надзвичайно переконливі, але стратегічно порожні тексти, які лише репродукують статистичні закономірності даних без реального розуміння контексту міжнародної безпеки [7]. Це створює ілюзію «експертної об'єктивності», яка може ввести в оману політичне керівництво, провокуючи прийняття рішень на основі помилкових алгоритмічних висновків.

Досліджуючи вплив інтелектуальних систем на дипломатичну практику, Б. Ндзєндзе та Т. Марвала [1] акцентують увагу на тому, що ШІ (зокрема генеративні моделі) здатен радикально змінити процеси верифікації інформації у міжнародних відносинах. Це створює нові виклики для дипломатії, де здатність алгоритмів маніпулювати сприйняттям реальності стає інструментом реалізації національних інтересів у цифрову епоху.

У військовій сфері штучний інтелект забезпечує безпрецедентне підвищення ефективності управління обмеженими ресурсами, стратегічного планування складних багатокомпонентних операцій та ешелонованого кіберзахисту. Досліджуючи цей аспект, Р. Пантюшенко та Ю. Чайка обґрунтовують, що завдяки автоматизації процесів аналізу та можливості миттєвого реагування на загрози, ШІ дозволяє діяти зі швидкістю, що значно перевищує людські можливості [12, с. 206]. На наш погляд, це свідчить про формування нового типу алгоритмічно керованого протистояння, де перевага у швидкості обробки інформації стає визначальним фактором, нівелюючи чисельну перевагу супротивника.

Емпіричним підтвердженням цієї тези є український досвід застосування платформ ситуаційної обізнаності та систем обробки даних (зокрема рішень компанії Palantir та вітчизняної системи «Delta»). За оцінками експертів, використання ШІ для аналізу супутникових знімків та відеопотоків дозволило

скоротити операційний цикл «сенсор-постріл» (*sensor-to-shooter loop*) з традиційних 20–30 хвилин до декількох десятків секунд. Важлива ланка в цій системі – це мережа широкосмугового зв'язку, що забезпечується згори масивом із приблизно 2500 супутників Starlink на низькій навколоземній орбіті [16]. На наше переконання, це доводить, що ІІІ стає «помножувачем сили», який дозволяє меншій за ресурсами армії ефективно протистояти кількісно більшому агресору за рахунок домінування у швидкості обробки інформації.

Проте такий технологічний оптимізм має бути збалансований розумінням стратегічних ризиків. У цьому контексті К. Пейн акцентує увагу на тому, що розвиток технологій ІІІ деформує саму природу стратегічного середовища, вносячи елемент фатальної непередбачуваності у процеси оцінки загроз [3]. Це дає підстави стверджувати, що штучний інтелект поступово нівелює класичну теорію стримування: швидкість та певна «непрозорість» алгоритмічних рішень роблять поведінку суб'єктів міжнародної безпеки менш прозорою для опонентів, що підвищує ризики помилкової інтерпретації намірів. Особливого значення тут набуває концепт «стиснення часу», описаний М. Горовіцом [5], який вказує на те, що автоматизація аналітичних процесів веде до критичного дефіциту часу для людської рефлексії. Узагальнюючи ці підходи, можна дійти висновку, що «ущільнення» аналітичного циклу не лише підвищує оперативність, а й парадоксальним чином знижує якість стратегічного передбачення, роблячи систему вразливою до системних багів та маніпуляцій супротивника.

Емпіричним підтвердженням окреслених теоретичних трансформацій є сучасний досвід відсічі збройній агресії проти України, що стала унікальним майданчиком для апробації штучного інтелекту в умовах конвенційної війни високої інтенсивності. Нинішня ситуація демонструє перехід від стратегії виснаження ресурсами до стратегії технологічної переваги в інформаційно-аналітичному циклі. Використання інтелектуальних систем для автоматизації розвідки на основі відкритих джерел (OSINT), інтеграції систем ситуаційної обізнаності (на кшталт «Delta») демонструє перехід від стратегії виснаження ресурсами до стратегії технологічної переваги в інформаційному циклі. Тобто, український кейс доводить, що здатність алгоритмів до миттєвої селекції цілей та прогнозування дій ворога стає ключовим елементом національної безпеки, дозволяючи ефективно нівелювати асиметрію військового потенціалу. На наш погляд, Україна сьогодні формує прецедент етичного та ефективного використання ІІІ в екстремальних умовах, що має стати основою для майбутніх міжнародних стандартів. Цей досвід підтверджує, що навіть за максимальної автоматизації процесів, людський контроль залишається незамінним запобіжником проти алгоритмічних помилок у критичних точках прийняття рішень.

Окремою прикметою трансформації безпекового ландшафту стає інтенсивна інтеграція штучного інтелекту в державні механізми управління, що виводить технологію за межі суто мілітарної сфери. Досліджуючи досвід провідних демократій (зокрема США), І. Лопатченко доводить, що впровадження алгоритмічних моделей у державний апарат дозволяє не лише оптимізувати складні управлінські процеси, а й суттєво підвищити адаптивність політики до асиметричних викликів глобалізації [13, с. 248]. Водночас автор висуває критичне застереження: така інтеграція вимагає розробки надійних національних протоколів верифікації. Сліпа довіра до «машинного розуму» у питаннях національної безпеки може спровокувати феномен автоматизаційного упередження (*automation bias*), що загрожує стратегічними викривленнями та поступовою втратою політичного контролю над прийняттям доленосних рішень.

Концептуально важливим аспектом сучасного наукового дискурсу є переосмислення ролі штучного інтелекту як специфічного квазісуб'єкта безпеки. У розвідці О. Дубовського обґрунтовано тезу, що ІІІ сьогодні виступає активним елементом, здатним безпосередньо трансформувати процеси формування аналітичних висновків, створюючи якісно новий тип антропо-технологічної взаємодії [11]. Фактично, мова йде про виникнення гібридних систем, де когнітивні спроможності людини-оператора органічно доповнюються обчислювальною потужністю самонавчальних алгоритмів. На наш погляд, це формує стан «розподіленого інтелекту», за якого стратегічне рішення постає як результат складного синтезу людської інтуїції та машинної обробки ймовірностей. Проте, як зазначає Б. Крістіан, методи навчання RLHF можуть консервувати когнітивні упередження розробників, створюючи замкнені контури ухвалення рішень, де помилка лише посилюється алгоритмом [6].

Разом із технологічним прогресом, розвиток штучного інтелекту породжує низку системних соціально-економічних викликів, які прямо корелюють із питаннями внутрішньої безпеки держав. Глибока автоматизація виробничих та аналітичних процесів може спровокувати структурні деформації на ринку праці, що, за прогнозами вітчизняних фахівців (зокрема в контексті аналізу сутності технології), неминуче призведе до зростання технологічного безробіття та поглиблення соціальної нерівності [15]. Це створює додаткові ризики нестабільності у глобальному масштабі, оскільки цифрова прірва між країнами-лідерами та рештою світу лише збільшуватиметься, провокуючи нові конфлікти за контроль над технологічними ресурсами.

На міжнародному рівні штучний інтелект остаточно утвердився як центральний елемент стратегічної конкуренції між провідними державами світу. Його динамічний розвиток визначає технологічну перевагу акторів та радикально змінює глобальний баланс сил, що зумовлює формування нових військово-політичних альянсів.

У цьому контексті особливої ваги набуває діяльність міжнародних безпекових інституцій, зокрема НАТО, які трансформують свої стратегічні концепції відповідно до викликів цифрової епохи. Така системна інтеграція алгоритмічних рішень у контури колективної оборони повністю корелює з візією О. Пархомчук, яка визначає штучний інтелект як засадничий мегатренд глобального розвитку. Дослідниця обґрунтовує, що ШІ сьогодні виступає не лише технологічним інструментом, а й автономним чинником, який докорінно переформатовує світову політику, економіку та систему міжнародної безпеки [8, с. 192].

Масштаби цієї трансформації підтверджуються динамікою оборонних бюджетів. За даними запиту на бюджет Міністерства оборони США на 2025 фінансовий рік, на цифрові технології та ШІ було передбачено рекордні 14,3 млрд дол., що свідчить про перехід від експериментальних розробок до етапу масового впровадження. Своєю чергою, НАТО в межах ініціативи DIANA (*Акселератор оборонних інновацій для Північної Атлантики*) створило інноваційний фонд обсягом 1 млрд євро, де ШІ визначено пріоритетним напрямом у більшості грантових програм, що фактично закріплює статус алгоритмічних систем як нового стандарту міжнародної безпеки [17].

Водночас процес глобальної інтеграції ШІ має суперечливий характер. Як слушно наголошує О. Фурсай, перетворення алгоритмів на ключовий елемент міжнародних відносин неминуче актуалізує нові загрози для інформаційної стабільності. Дослідник акцентує увагу на тому, що здатність ШІ до генерації та поширення високоточного маніпулятивного контенту в промислових масштабах стає інструментом для проведення складних дезінформаційних операцій [14, с. 172]. Це створює ситуацію, за якої технологічний прогрес, забезпечуючи воєнну перевагу одних акторів, водночас нівелює епістемічний фундамент міжнародної довіри та глобальної безпеки в цілому. Кількісні показники підтверджують небезпеку цього тренду: за даними спільного звіту Microsoft Threat Intelligence та OpenAI (2024), державні хакерські угруповання вже активно використовують великі мовні моделі (LLM) для автоматизації фішингу та вдосконалення програмного коду атак. Більше того, статистика компанії Sumsb фіксує вибухове зростання кількості дипфейків, що поступово нівелює довіру до цифрового контенту на глобальному рівні [18].

Отже, майбутнє міжнародної безпеки залежатиме від здатності світової спільноти виробити єдині правила «алгоритмічної поведінки», де людина залишатиметься суб'єктом етичного контролю, запобігаючи деградації критичного аналізу перед обличчям швидких цифрових рішень.

Висновки. Під час дослідження встановлено, що концепт штучного інтелекту в сучасній системі міжнародної безпеки остаточно подолав межі суто технологічного дискурсу. Сьогодні ШІ є фундаментальним стратегічним чинником, який формує нову парадигму глобальної стабільності. Він виступає інтегральною складовою безпекової архітектури, трансформуючи не лише інструментарій ведення конфліктів, а й саму філософію захисту національних інтересів. Ключовим викликом при цьому стає поступова ерозія класичних моделей стримування, спричинена тотальною алгоритмізацією стратегічних процесів. Світ спостерігає перехід від статичного балансу сил до динамічного «алгоритмічного протистояння». У цій новій реальності безпека держави прямо залежить від швидкості обробки великих масивів даних та стійкості когнітивних систем управління до маніпуляцій ззовні.

Доведено, що інтеграція генеративних мовних моделей і автономних систем прийняття рішень породжує якісно нові стратегічні ризики. Зокрема, автоматизаційне упередження та ймовірні алгоритмічні помилки у процесах ідентифікації цілей створюють пряму загрозу для міжнародної стабільності, суттєво підвищуючи ризик неконтрольованої ескалації. Безпековий вимір штучного інтелекту також чітко проявляється у масштабній трансформації інформаційного простору. Здатність алгоритмів продукувати дезінформацію в промислових обсягах руйнує епістемічний фундамент міжнародної довіри. Це перетворює когнітивну безпеку на пріоритетний напрям державної політики. Важливим підсумком роботи є теза, що ефективність безпекової архітектури в цифрову епоху критично залежить від збереження людської суб'єктності. Необхідно терміново розробляти міжнародні стандарти «алгоритмічної гігієни», спираючись на вже наявні прецеденти, зокрема Закон ЄС про штучний інтелект (EU AI Act) [19]. Хоча цей документ переважно орієнтований на цивільний сектор, його ризик-орієнтований підхід має стати юридичним фундаментом для етичного контролю над технологіями у сфері безпеки, унеможливаючи диктатуру машинних рішень над політичною волею.

Перспективи подальших наукових пошуків у контексті міжнародної безпеки логічно пов'язані з розробкою оновлених моделей стратегічної стабільності. Це особливо актуально в умовах асиметричного доступу різних акторів до передових технологій ШІ. Потребує глибшого вивчення питання трансформації

«державного суверенітету» у цифровому просторі, де межі безпеки розмиваються через активність автономних кіберсистем та генеративних моделей впливу. Пріоритетним вектором залишається дослідження правової відповідальності за рішення, які приймаються алгоритмами у військовій сфері. Також важливо створювати міжнародно-правові запобіжники проти використання ШІ як інструмента когнітивної агресії. Особливої уваги заслуговує аналіз досвіду інтеграції ШІ у стратегії НАТО та інших оборонних союзів як фактор зміцнення колективної безпеки. Подальші розвідки можуть бути зосереджені на побудові комплексних систем моніторингу алгоритмічних загроз та адаптації положень EU AI Act до специфіки сектору міжнародної безпеки та оборони. Це дозволить забезпечити стійкість світової системи безпеки перед обличчям непередбачуваного технологічного прогресу.

Література

1. Ndzendze B., Marwala T. Artificial Intelligence and International Relations Theories. *Singapore: Springer Nature Singapore*, 2023. 215 p. DOI: <https://doi.org/10.1007/978-981-19-4873-2>.
2. Scharre P. *Army of None: Autonomous Weapons and the Future of War*. New York: W. W. Norton & Company, 2018. URL: <https://www.paulscharre.com/army-of-none> (дата звернення 01.03.2026).
3. Payne K. Artificial Intelligence: A Revolution in Strategic Affairs? *Survival*. 2018. Vol. 60, Issue 5. P. 7-32. URL: <https://kclpure.kcl.ac.uk/portal/en/publications/artificial-intelligence-a-revolution-in-strategic-affairs/> (дата звернення 01.03.2026).
4. Johnson J. *Artificial Intelligence and the Future of Warfare: The USA, China, and strategic stability*. Manchester University Press. 2021. URL: <https://www.jstor.org/stable/jj.21996280> (дата звернення: 01.03.2026).
5. Horowitz M. C., Scharre P., Velez-Green A. *AI and International Stability: Risks and Confidence-Building Measures*. Center for a New American Security. 2020. URL: <https://www.cnas.org/publications/reports/ai-and-international-stability> (дата звернення: 02.03.2026).
6. Christian B. *The Alignment Problem: Machine Learning and Human Values*. New York: W. W. Norton & Company, 2020. URL: https://www.researchgate.net/publication/367576352_The_Alignment_Problem_Machine_Learning_and_Human_Values (дата звернення: 02.03.2026).
7. Bender E. M. et al. On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 2021. URL: <https://faculty.washington.edu/ebender/papers/Bender-NE-ExpAI.pdf> (дата звернення: 02.03.2026).
8. Забара І.М., Деркач О.О. Технології військового штучного інтелекту: сучасні міжнародно-правові напрямки наукових досліджень. *Науковий вісник Ужгородського національного університету*. Серія: Право. 2024. Вип. 86, ч. 5. С. 241–245. DOI: <https://doi.org/10.24144/2307-3322.2024.86.5.36>.
9. Пархомчук О.С. Штучний інтелект як мегатренд глобального розвитку. *Політикус*. 2023. № 4. С. 191–196. DOI: <https://doi.org/10.24195/2414-9616.2023-4.29>.
10. Блавацька Н.М., Рибка Д.М. Штучний інтелект у сфері національної безпеки: етичні та правові аспекти. *Науковий вісник Ужгородського національного університету*. Серія: Право. 2025. Вип. 91, ч. 5. С. 84–89. DOI: <https://doi.org/10.24144/2307-3322.2025.91.5.12>.
11. Дубовський О. Інформаційна безпека: суб'єктність і штучний інтелект. *Journal of Innovations and Sustainability*. 2024. Т. 8, № 2. Ст. 09. DOI: <https://doi.org/10.51599/is.2024.08.02.09>.
12. Пантюшенко Р., Чайка Ю. Штучний інтелект у сфері кібербезпеки: інновації, виклики та перспективи розвитку. *Military Science*. 2024. Т. 2, № 1. С. 200–207. DOI: <https://doi.org/10.62524/msj.2024.2.1.16>.
13. Лопатченко І.С. Застосування досвіду США у використанні штучного інтелекту в забезпеченні національної інформаційної безпеки України. *Державне будівництво*. 2024. № 1 (35). С. 247–257. DOI: <https://doi.org/10.26565/1992-2337-2024-1-18>.
14. Фурсай О.В. «Штучний інтелект» як виклик міжнародній інформаційній безпеці. *Регіональні студії*. 2023. № 35. С. 167–173. DOI: <https://doi.org/10.32782/2663-6170/2023.35.28>.
15. Єфремов М. Ф., Єфремов Ю. М. (2016). Штучний інтелект, історія та перспективи розвитку. *Вісник ЖДТУ. Серія «Технічні науки»*, (2(45), 123–126. DOI: [https://doi.org/10.26642/tn-2008-2\(45\)-123-126](https://doi.org/10.26642/tn-2008-2(45)-123-126).
16. Ignatius David. How the algorithm tipped the balance in Ukraine. December 19, 2022. URL: <https://www.washingtonpost.com/opinions/2022/12/19/palantir-algorithm-data-ukraine-war/> (дата звернення: 03.03.2026).
17. DIANA's Accelerator Programme. URL: <https://www.diana.nato.int/accelerator-programme.html> (дата звернення: 04.03.2026).
18. Microsoft Threat Intelligence. Staying ahead of threat actors in the age of AI. 2024. URL: <https://www.microsoft.com/en-us/security/blog/> (дата звернення: 04.03.2026).

19. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). URL: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj> (дата звернення: 04.03.2026).

References

1. Ndzendze B., Marwala T. (2023). *Artificial Intelligence and International Relations Theories*. Singapore: Springer Nature Singapore. 215 p. DOI: <https://doi.org/10.1007/978-981-19-4873-2>.
2. Scharre P. (2018). *Army of None: Autonomous Weapons and the Future of War*. New York: W. W. Norton & Company. URL: <https://www.paulscharre.com/army-of-none> (accessed: March 01, 2026).
3. Payne K. (2018). Artificial Intelligence: A Revolution in Strategic Affairs? *Survival*. Vol. 60, Issue 5. P. 7–32. URL: <https://kclpure.kcl.ac.uk/portal/en/publications/artificial-intelligence-a-revolution-in-strategic-affairs/> (accessed: March 01, 2026).
4. Johnson J. (2021). *Artificial Intelligence and the Future of Warfare: The USA, China, and strategic stability*. Manchester University Press. URL: <https://www.jstor.org/stable/jj.21996280> (accessed: March 01, 2026).
5. Horowitz M. C., Scharre P., Velez-Green A. (2020). *AI and International Stability: Risks and Confidence-Building Measures*. Center for a New American Security. URL: <https://www.cnas.org/publications/reports/ai-and-international-stability> (accessed: March 02, 2026).
6. Christian B. (2020). *The Alignment Problem: Machine Learning and Human Values*. New York: W. W. Norton & Company. URL: https://www.researchgate.net/publication/367576352_The_Alignment_Problem_Machine_Learning_and_Human_Values (accessed: March 02, 2026).
7. Bender E. M. et al. (2021). On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*. URL: <https://faculty.washington.edu/ebender/papers/Bender-NE-ExpAI.pdf> (accessed: March 02, 2026).
8. Zabara I. M., Derkach O. O. (2024). Tekhnologii viiskovoho shtuchnoho intelektu: suchasni mizhnarodno-pravovi napriamky naukovykh doslidzhen [Technologies of military artificial intelligence: modern international legal directions of scientific research]. *Naukovyi visnyk Uzhhorodskoho natsionalnoho universytetu. Serii: Pravo. Vyp. 86, ch. 5*. S. 241–245. DOI: <https://doi.org/10.24144/2307-3322.2024.86.5.36>.
9. Parkhomchuk O. S. (2023). Shtuchnyi intelekt yak mehatrend hlobalnoho rozvytku [Artificial intelligence as a megatrend of global development]. *Politikus*. № 4. S. 191–196. DOI: <https://doi.org/10.24195/2414-9616.2023-4.29>.
10. Blavatska N. M., Rybka D. M. (2025). Shtuchnyi intelekt u sferi natsionalnoi bezpeky: etychni ta pravovi aspekty [Artificial intelligence in the field of national security: ethical and legal aspects]. *Naukovyi visnyk Uzhhorodskoho natsionalnoho universytetu. Serii: Pravo. Vyp. 91, ch. 5*. S. 84–89. DOI: <https://doi.org/10.24144/2307-3322.2025.91.5.12>.
11. Dubovskyi O. (2024). Informatsiina bezpeka: subiektivist i shtuchnyi intelekt [Information security: subjectivity and artificial intelligence]. *Journal of Innovations and Sustainability*. T. 8, № 2. St. 09. DOI: <https://doi.org/10.51599/is.2024.08.02.09>.
12. Pantiushenko R., Chaika Yu. (2024). Shtuchnyi intelekt u sferi kiberbezpeky: innovatsii, vyklyky ta perspektyvy rozvytku [Artificial intelligence in the field of cybersecurity: innovations, challenges, and development prospects]. *Military Science*. T. 2, № 1. S. 200–207. DOI: <https://doi.org/10.62524/msj.2024.2.1.16>.
13. Lopatchenko I. S. (2024). Zastosuvannia dosvidu SShA u vykorystanni shtuchnoho intelektu v zabezpechenni natsionalnoi informatsiinoi bezpeky Ukrainy [Applying the US experience in the use of artificial intelligence in ensuring national information security of Ukraine]. *Derzhavne budivnytstvo*. № 1 (35). S. 247–257. DOI: <https://doi.org/10.26565/1992-2337-2024-1-18>.
14. Fursai O. V. (2023). «Shtuchnyi intelekt» yak vyklyk mizhnarodnii informatsiinii bezpetsi [“Artificial intelligence” as a challenge to international information security]. *Rehionalni studii*. № 35. S. 167–173. DOI: <https://doi.org/10.32782/2663-6170/2023.35.28>.
15. Yefremov M. F., Yefremov Yu. M. (2016). Shtuchnyi intelekt, istoriia ta perspektyvy rozvytku [Artificial intelligence, history and development prospects]. *Visnyk ZhDTU. Serii «Tekhnichni nauky»*. (2(45), 123–126. DOI: [https://doi.org/10.26642/tn-2008-2\(45\)-123-126](https://doi.org/10.26642/tn-2008-2(45)-123-126).
16. Ignatius David. How the algorithm tipped the balance in Ukraine. December 19, 2022. URL: <https://www.washingtonpost.com/opinions/2022/12/19/palantir-algorithm-data-ukraine-war/> (accessed: March 03, 2026).
17. DIANA's Accelerator Programme. URL: <https://www.diana.nato.int/accelerator-programme.html> (accessed: March 04, 2026).
18. Microsoft Threat Intelligence. (2024). Staying ahead of threat actors in the age of AI. URL: <https://www.microsoft.com/en-us/security/threat-intelligence>

www.microsoft.com/en-us/security/blog/ (accessed: March 04, 2026).

19. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). URL: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj> (accessed: March 04, 2026).

Отримано: 06.03.2026

Прийнято до публікації: 06.04.2026

Опубліковано: 15.04.2026